

Following the Autoregressive Nature of LLM Embeddings via Compression and Alignment

Jingcheng Deng^{1,2 *}, Zhongtao Jiang^{3*}, Liang Pang^{1†}, Zihao Wei^{1,2},
Liwei Chen³, Kun Xu, Yang Song, Huawei Shen^{1,2}, Xueqi Cheng^{1,2}

¹Key Laboratory of AI Safety, Institute of Computing Technology,
Chinese Academy of Sciences

² University of Chinese Academy of Sciences

³ Kuaishou Technology

{dengjingcheng23s, pangliang}@ict.ac.cn

Abstract

A new trend uses LLMs as dense text encoders via contrastive learning. However, since LLM embeddings predict the probability distribution of the next token, they are inherently generative and distributive, conflicting with contrastive learning, which requires embeddings to capture full-text semantics and align via cosine similarity. This discrepancy hinders the full utilization of LLMs’ pre-training capabilities, resulting in inefficient learning. In response to this issue, we propose AutoRegEmbed, a new contrastive learning method built on embedding conditional probability distributions, which integrates two core tasks: information compression and conditional distribution alignment. The information compression task encodes text into the embedding space, ensuring that the embedding vectors capture global semantics. The conditional distribution alignment task focuses on aligning text embeddings with positive samples embeddings by leveraging the conditional distribution of embeddings while simultaneously reducing the likelihood of generating negative samples from text embeddings, thereby achieving embedding alignment and uniformity. Experimental results demonstrate that our method significantly outperforms traditional contrastive learning approaches and achieves performance comparable to state-of-the-art models when using the same amount of data. Our code is available at <https://github.com/TrustedLLM/AutoRegEmbed>

1 Introduction

Text embeddings, which represent the semantic content of natural language text as vectors, are extensively utilized in domains such as information retrieval (Xia et al., 2015), semantic similarity assessment, retrieval-augmented generation (RAG) (Khandelwal et al., 2020; Shi et al., 2024; Deng et al., 2023; Ding et al., 2024; Xu et al., 2024b,

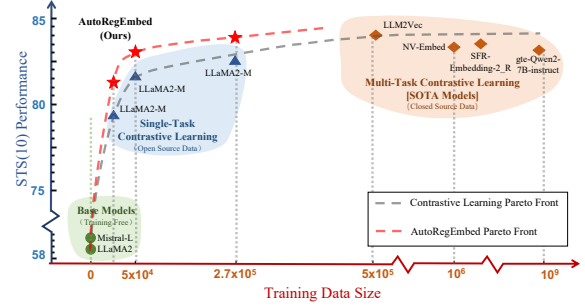


Figure 1: Comparison of pareto front between AutoRegEmbed and other methods. The horizontal axis represents the number of training samples, while the vertical axis indicates the average performance across 10 STS datasets. The upper left corner represents the region with the highest learning efficiency.

2025), LLMs-based agents (Chen et al., 2024a; Xu et al., 2024a) and data attribution (Beigi et al., 2024). Traditional text embedding models typically employ transformer-based architectures with encoder-only designs, including examples like Bert (Devlin et al., 2019), DeBERTa (He et al., 2021) and MPNet (Song et al., 2020), and are trained using contrastive learning.

After extensive pre-training on a large-scale corpus, LLMs have outperformed previous encoder-only small models (Deng et al., 2022) and demonstrated strong adaptability across diverse downstream tasks (Zhao and Zhang, 2024; Wang et al., 2024b; Zhao et al., 2025; Duan et al., 2025; Liu et al., 2025). Recently, contrastive learning has been directly applied to decoder-only LLMs, which are trained to generate embedding vectors based on task-specific instructions, enabling adaptability to various embedding scenarios (Lee et al., 2024; BehnamGhader et al., 2024). Despite initial advancements, training a high-performance 7B-scale text embedding model using this approach remains highly resource-intensive. It typically requires millions of triplets (Wang et al., 2024a; Li et al., 2024b, 2023) and substantial computational power, includ-

* Equal Contributions

† Corresponding Author

ing thousands of hours on an A100 80GB GPUs (Muennighoff et al., 2024; Ma et al., 2024), even with the application of Parameter-Efficient Fine-Tuning (PEFT) (Hu et al., 2022; Dao, 2024). The high resource consumption might reasonably be attributed to the inability of the *discriminative* contrastive learning method to fully harness the capabilities of *generative* LLMs (Li et al., 2024a). Firstly, the constraint of unidirectional attention in LLMs leads to the aggregation of information in the hidden state of the output layer corresponding to the final token. However, as LLMs are optimized for next-token prediction, this hidden state can only represent the semantics of the next token (local) rather than the semantics of the input text itself (global). Consequently, employing this hidden state directly in contrastive learning requires additional training time and computational cost to transition from a localized to a more global semantic representation. Secondly, the hidden state in LLMs is used to generate the probability distribution of the next token, whereas contrastive learning optimizes the cosine distance between the hidden states of different texts. This divergence in optimization objectives introduces additional training costs. This raises an important question: *Is it feasible to develop a method that follows the autoregressive nature while generating high-quality text embeddings and significantly reducing resource requirements?*

We formalize three key requirements to address this problem. Firstly, embeddings should capture global semantics rather than focusing solely on next-token semantics. Secondly, they must follow alignment and uniformity principles (Wang and Isola, 2020). Finally, the transformation from the original embedding to one that meets these criteria should follow an autoregressive nature. To this end, we propose AutoRegEmbed, which encompasses two tasks: information compression and conditional distribution alignment.

The information compression task is inspired by the concept of context compression (Chevalier et al., 2023; Ge et al., 2024; Mu et al., 2023), which addresses the limitations of context window length and the high computational cost faced by LLMs when processing long texts. Specifically, we encode the context and instructions into a set of compressed variables, which are then passed to a decoder with the same architecture but frozen parameters, forcing it to reconstruct the corresponding target. By restricting the decoder to rely solely

on the compressed variables—without access to the original context or instructions—we introduce an information bottleneck. This ensures that the compressed variables effectively capture the global semantics of the instructions and context.

The conditional distribution alignment task draws inspiration from traditional contrastive learning and LLM alignment techniques (Wang et al., 2024c). We begin by treating the compressed vectors as embeddings of their corresponding inputs. Then, we adopt the structure of the InfoNCE (van den Oord et al., 2018) loss function, but redefine the similarity metric. Simply put, we align the distance between the conditional probability distributions of text and positive sample embeddings while increasing the likelihood of text embeddings generating positive samples and decreasing the likelihood of generating negative samples. This approach promotes the alignment and uniformity of compressed variables while maintaining the autoregressive nature.

Experimental results demonstrate that AutoRegEmbed outperforms traditional contrastive learning methods while utilizing the same computational resources, making it a highly efficient and scalable solution. Remarkably, even with a limited number of training samples, AutoRegEmbed achieves performance on par with the current state-of-the-art (SOTA) models, showcasing its superior ability to learn robust and generalizable representations from scarce data. As shown in Figure 1, the Pareto frontier of AutoRegEmbed consistently outperforms traditional contrastive learning methods, demonstrating a more optimal trade-off between computational efficiency and performance. This indicates that AutoRegEmbed achieves superior representation learning while maintaining a balanced resource utilization.

2 Related Works

Text embedding is a technique that maps text data into a numerical vector space, capturing both semantic and contextual features of the text. Research in this area can be divided into three categories based on the underlying model: Early models, LLMs with fine-tuning, and LLMs without fine-tuning.

Early Models Early approaches include SentenceBERT (Reimers and Gurevych, 2019) (supervised) and SimCSE (Gao et al., 2021) (unsupervised), which leverage contrastive learning to

generate high-quality text embeddings using small encoder-only models. Inspired by instruction fine-tuning, recent research (Su et al., 2023; Asai et al., 2023) has shifted toward using text paired with instructions to enhance the generalization and transferability of text embeddings in complex scenarios. At this stage, several studies have investigated the use of generative tasks to enhance text embeddings. For example, coCondenser (Gao and Callan, 2022) proposes a pre-training strategy that improves the expressiveness of the [CLS] token using unsupervised corpora. PaSeR (Wu and Zhao, 2022) and RetroMAE (Xiao et al., 2022) adopt encoder-decoder architectures: PaSeR emphasizes the reconstruction of key phrases, while RetroMAE employs a masked-token objective to recover the original sentence. However, these approaches do not incorporate contrastive supervision signals into their generative objectives. Moreover, they are evaluated only on smaller-scale models, making them fundamentally different from our method in both design philosophy and scaling strategy.

LLMs with Fine-Tuning Many studies have focused on transforming LLMs into text embedding models through contrastive learning fine-tuning. RepLLaMA (Ma et al., 2024), for example, follows the DPR (Karpukhin et al., 2020) pipeline, using the hidden state of the last token generated by LLaMA as a text embedding vector and applying contrastive learning fine-tuning. Recognizing that the unidirectional attention mechanism in LLMs may limit text embedding quality, LLM2Vec (BehnamGhader et al., 2024) introduces a bidirectional attention mechanism combined with average pooling to enhance embedding quality. NV-Embed (Lee et al., 2024) takes this further by incorporating an additional Latent Attention Layer to generate pooled embeddings. bge-en-icl (Li et al., 2024b) suggests that retaining the original framework of LLMs and leveraging in-context learning is the optimal approach for generating text embeddings. Some studies (Wang et al., 2024a) even use synthetic data generated by LLMs, rather than real-world data, for fine-tuning and achieve competitive performance on the MTEB leaderboard (Muenighoff et al., 2023). However, these approaches often overlook the fundamental differences between language modeling and contrastive learning, failing to fully leverage the potential of LLMs. More closely related to our work is Llama2Vec (Li et al., 2024a), which proposes two pretext tasks to en-

able unsupervised adaptation of LLMs, followed by contrastive learning fine-tuning to achieve better performance. In contrast, our approach achieves strong results without any need for traditional contrastive learning fine-tuning (cosine-based), as our task fully exploits the inherent potential of LLMs.

LLMs without Fine-Tuning Several studies have explored methods to transform LLMs into text encoders without fine-tuning. (Liu et al., 2024) proposed using possible trajectory distributions as text representations, achieving effectiveness but at a high computational cost. (Springer et al., 2024) introduced echo embeddings by repeatedly feeding text into autoregressive models, addressing architectural limitations but doubling computational requirements. Other methods focus on prompt adjustments to produce meaningful embeddings. PromptEOL (Jiang et al., 2024) introduced a One-Word Limitation prompt to improve embedding performance, while MetaEOL (Lei et al., 2024) extended this idea by using eight different prompt types to generate multi-view embeddings. GenEOL (Thirukovalluru and Dhingra, 2024) leveraged LLMs to create various sentence transformations that retain their meaning, aggregating the resulting embeddings to enhance the overall sentence representation. Meanwhile, PromptReps (Zhuang et al., 2024) developed a hybrid document retrieval framework leveraging prompts to address challenges in information retrieval tasks. Despite these innovations, these approaches either perform poorly or require multiple inferences to achieve good results. By contrast, our method surpasses these methods with minimal training costs.

3 Method

In this section, we first introduce the preliminary information about the task of text embedding with instructions. We then discuss the information compression, which transitions LLM embeddings from local semantics to global semantics, followed by the conditional distribution alignment, which optimizes the conditional probability distribution of embeddings to ensure alignment and uniformity.

3.1 Preliminary

Text embeddings with instructions can adapt to various downstream tasks. Formally, given a large collection $D = \{d_1, d_2, \dots, d_N\}$ containing N documents, as well as a text q and an instruction t , the embedding $e_{q,t} = E(q, t)$ generated from q

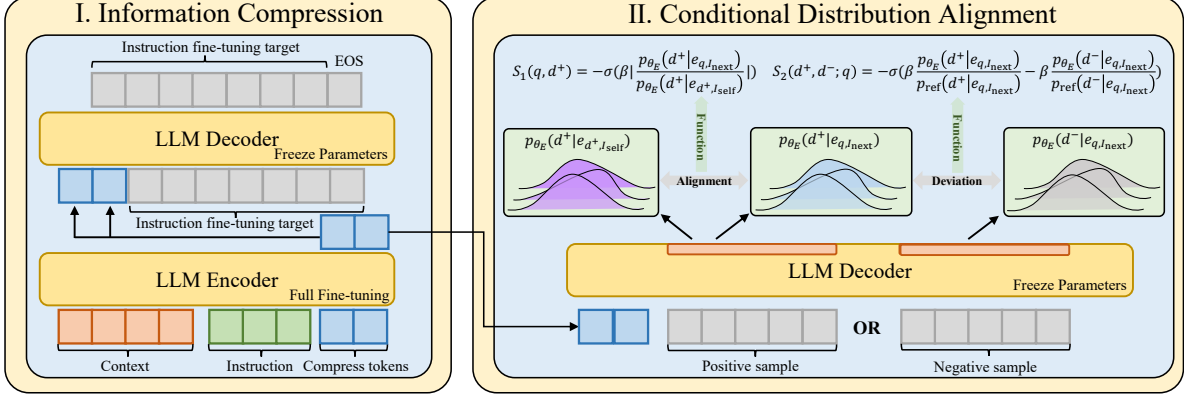


Figure 2: Overall framework of AutoRegEmbed. Firstly, we perform the information compression task to inject key information from the context and instruction into the compressed tokens. Then, we optimize the conditional probability distribution of these tokens to align the distributions of $e_{q,I_{next}}$ and $e_{d^+,I_{self}}$ as closely as possible through $S_1(q, d^+)$, while increasing the probability of $e_{q,I_{next}}$ generating positive samples and reducing the probability of $e_{q,I_{next}}$ generating negative samples through $S_2(d^+, d^-; q)$. Encoder and decoder share a structure.

and t can match documents $d \in D$ that are relevant to q , according to t , where E represents the text encoder. Thus, by simply changing the instruction t , the relevance measure can be adapted to different downstream tasks. For example, for dense retrieval tasks, the instruction might be “*find documents that can answer this question*,” while for semantic similarity tasks, the instruction could be “*find sentences that are semantically similar to this text*”. Numerous studies have explored various embedding techniques and instruction diversities. Our goal is to identify a simple yet effective way to enable LLMs to generate high-quality embeddings directly from autoregressive framework.

3.2 Information Compression: from Discriminative to Generative Embeddings

In this section, we first explain the motivation for transitioning from discriminative embeddings to generative embeddings, followed by a formal definition of the information compression task.

In decoder-based LLMs, embeddings are typically generated by extracting the hidden state of the final token in the input sequence. However, this approach primarily captures the semantics of the first output token rather than encoding the global semantics of the entire input. Various pooling techniques, such as average pooling and attention pooling, have been explored to mitigate this limitation, yet they introduce their own challenges. The average pooling method, which computes the mean of all token hidden states, does not necessarily encapsulate global semantics but instead serves as a mechanism for “convexity preservation” (Li et al.,

2020). Conversely, attention pooling modifies the attention mechanism or introduces additional parameters, thereby altering the original architecture of LLMs. Such modifications deviate from the model’s pre-training design and can lead to unintended consequences, as prior studies (Li et al., 2024b) indicate that maintaining the original LLM framework often yields optimal performance. To enable LLMs to generate embeddings that represent global semantics, we introduce an information compression task. This task compels LLMs to reconstruct the original target using a compressed embedding derived from the input text. Given that this compressed embedding models the conditional probability distribution of the target, we designate it as the generative embedding to contrast it with the discriminative embedding produced by conventional pooling approaches.

The information compression task is inspired by the concept of context compression. Specifically, we append k compressed tokens $c = (c_1, \dots, c_k)$, where $k \ll n + m$, to the text $q = (q_1, \dots, q_n)$ and instruction $t = (t_1, \dots, t_m)$, with n and m representing their respective token lengths. This combined (q, t, c) is then fed into an encoder E to generate the embedding $e_c = (e_{c_1}, \dots, e_{c_k})$. As mentioned earlier, we expect the embedding e_c to capture the global semantics of the text q and the instruction t . To achieve this, we input e_c into a frozen decoder D , which shares the same architecture, and force it to generate the most relevant document d . The optimization objective for this

task can be expressed as:

$$\begin{aligned}\mathcal{L}_{\text{IC}} &= \max_{e_{c_1}, \dots, e_{c_k}} P(d|e_{c_1}, \dots, e_{c_k}; \theta_D) \\ &= \max_{\Theta_E} P(d|c_1 \dots c_k, t_1 \dots t_m, q_1 \dots q_n; \theta_E, \theta_D),\end{aligned}$$

where θ_E and θ_D denote the parameters of E and D , respectively.

3.3 Conditional Distribution Alignment: from Data-Point to Distribution Perspective

After addressing the global semantic representation issue of the embedding vector, we also require the embedding vector to meet the criteria of alignment and uniformity. In general, we optimize these two properties asymptotically using a contrastive loss, such as InfoNCE,

$$\begin{aligned}\mathcal{L}_{\text{InfoNCE}}(f; \tau) &= \\ \mathbb{E}[-\log \frac{e^{f(q)^T f(d^+)/\tau}}{e^{f(q)^T f(d^+)/\tau} + \sum_i e^{f(q)^T f(d_i^-)/\tau}}],\end{aligned}\quad (1)$$

where τ denotes the temperature parameter and d_i^- represents the i -th negative sample. Clearly, Equation 1 differs fundamentally from the generative pre-training task, as it optimizes the cosine distance between sample embeddings, aligning data points in the embedding space rather than modeling the next-token probability distribution, which is central to pre-training. So, using this loss function to optimize an LLM may not fully unlock its potential.

To address this, we propose the Conditional Distribution Alignment task to minimize this discrepancy as much as possible. The concept is straightforward: Instead of using the cosine distance between embeddings, we assess similarity based on the conditional probability distribution corresponding to each embedding. Simply put, we extend point alignment to distribution alignment. Formally, the decoder L_D is a well-trained autoregressive language model with the following conditional probability distribution:

$$p(d|e_c) = \prod_{t=1}^T p(d_t|d_{<t}, e),$$

where $e_c = (e_{c_1}, \dots, e_{c_k})$ is the embedding variables, $d = (d_1, d_2, \dots, d_T)$ represents the generated sentence, and $d_{<t}$ denotes the part of the sentence before time step t . Intuitively, the similarity between corresponding samples q and d can be

measured by computing the distance between the conditional probability distributions of their embeddings, e_q and e_d :

$$S(q, d) = \frac{1}{T} \sum_{t=1}^T D(p(d_t|d_{<t}, e_q), p(d_t|d_{<t}, e_d)),$$

where $D(\cdot, \cdot)$ is any function that measures the divergence between two probability distributions. Given that the embedding distribution is instruction-dependent (see Section 3.1), we can adopt multiple approaches to measure similarity between q and d^+/d^- . Our empirical approach uses a basic alignment strategy: aligning the probability distributions of q and d via distinct instructions (I_{next} and I_{self} ; see Appendix D). This increases the probability of q generating d^+ and decreases that of generating d^- . As shown in Appendix E, this simple alignment outperforms more complex alternatives. Building on the above insights and incorporating the structure of InfoNCE, we empirically derive the final loss function:

$$\begin{aligned}\mathcal{L}_{\text{CDA}} &= \mathbb{E}[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(d^+, d_i^-; q)/\tau}}], \\ S_1(q, d^+) &= -\sigma(\beta |\log \frac{p_{\theta_E}(d^+|e_q, I_{\text{next}})}{p_{\theta_E}(d^+|e_{d^+, I_{\text{self}}})}|), \\ S_2(d^+, d_i^-; q) &= -\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_q, I_{\text{next}})}{p_{\text{ref}}(d^+|e_q, I_{\text{next}})} \\ &\quad - \beta \log \frac{p_{\theta_E}(d_i^-|e_q, I_{\text{next}})}{p_{\text{ref}}(d_i^-|e_q, I_{\text{next}})}),\end{aligned}\quad (2)$$

where τ and β are temperature parameters, and p_{θ_E} represents the initial model. p_{ref} denotes the log probability of the model generating the target before conditional distribution alignment. We use the Sigmoid function $\sigma(\cdot)$ (see Appendix C for analysis) to normalize the similarity measured from the conditional probability distribution to the range $[0, 1]$ to prevent overflow during the exponential operation. S_1 represents the similarity function between text q and the positive sample d^+ . We define it by measuring the absolute value of the difference in the logarithmic probability of their corresponding embeddings, e_q, I_{next} and $e_{d^+, I_{\text{self}}}$, generating the positive sample d^+ . To minimize this difference, we apply the absolute value function. In addition, we then add a negative sign to ensure that the value of S_1 increases as the similarity between q and d^+ increases. S_2 calculates the difference between the logarithmic probabilities of generating

positive and negative samples for text q , similar to DPO (Rafailov et al., 2023). We amplify this difference to boost the probability of embedding $e_{q,I_{\text{next}}}$ generating positive samples and decrease the probability of generating negative samples. We normalize the probabilities by dividing them by the corresponding values from the initial model to account for the length discrepancy between positive and negative samples. Note that during inference, the model still computes similarity between $e_{q,I_{\text{next}}}$ and $e_{d,I_{\text{self}}}$ using inner product matching, enabling seamless integration with existing embedding systems.

4 Experiments

4.1 Experimental Settings

Evaluations Previous studies (Gao et al., 2021; Li et al., 2020) highlight that a key goal of text embedding is to cluster semantically similar sentences. Following this approach, we use the MTEB (Muennighoff et al., 2023) evaluation framework to evaluate AutoRegEmbed on ten semantic text similarity datasets, including STS12 (Agirre et al., 2012), STS13 (Agirre et al., 2013), STS14 (Agirre et al., 2014), STS15 (Agirre et al., 2015), STS16 (Agirre et al., 2016), STS17 (Cer et al., 2017), STS22 (Chen et al., 2022), STS-B, BIOSSES and SICK-R. Each pair of text in the STS dataset is labeled with a similarity score ranging from 0 to 5 or 0 to 4, indicating their semantic similarity. The evaluation metric is the Spearman correlation between the similarity scores predicted by the model and the scores annotated by humans. In addition, to evaluate retrieval performance, we assess the model on MS MARCO (Nguyen et al., 2016), NFCorpus (Boteva et al., 2016), and SCIDOCS (Cohan et al., 2020) using nDCG@10 as the evaluation metric.

Training In the information compression stage, we use the training set of the instruction fine-tuning dataset PWC (Ge et al., 2024), which includes a diverse range of instruction types, as the training data. The original dataset contains 241,564 (context, instruction, target) samples. To reduce redundancy caused by repeated contexts, we remove duplicates, resulting in the PWC-Unique dataset with 16,382 samples as the final training data. In the conditional distribution alignment stage, we use the NLI part of the MEDI (Wang et al., 2024a) and BGE (Chen et al., 2024b) datasets as training data. The former contains 50,000 samples, while the latter consists of 274,951 samples. Each sample includes

an anchor, a positive sample, and a negative sample. Unless otherwise specified, the AutoRegEmbed results presented in the experiment section are based on the NLI part of the MEDI. For the retrieval task, we train our model on the MS MARCO training set. Because of the asymmetric nature of the retrieval task, the gap between the cosine distance-based evaluation metric and the conditional probability-based optimization objective becomes pronounced. To mitigate this, we perform an additional epoch of contrastive fine-tuning to better align with the evaluation process. Based on prior studies and empirical observations, the quality of negative samples plays a crucial role in training effectiveness. To enhance negative sampling, we adopt NV-Embed for hard negative mining. Specifically, we randomly select 7 hard negatives from the ranked list positions 30 to 210, treating these as challenging negative samples during training.

Baselines We categorize the baselines into three groups: (1) models without contrast training, including base models with various embedding methods using the same instructions as AutoRegEmbed and prompt-adjusted embedded models, including Echo (Springer et al., 2024), PromptEOL (Jiang et al., 2024), MetaEOL (Lei et al., 2024), and GenEOL (Thirukovalluru and Dhingra, 2024); (2) unsupervised contrast training models, primarily LLM2Vec (BehnamGhader et al., 2024) with different base models; and (3) supervised contrast training models, which consist of NV-Embed (Lee et al., 2024), SFR-Embedding-2_R (Meng et al., 2024), gte-Qwen2-7B-instruct (Li et al., 2023), LLM2Vec (BehnamGhader et al., 2024), and fair baselines.

4.2 Main Results

Table 1 summarizes the results of various baselines and AutoRegEmbed on ten STS datasets, along with the training data required for each method.

AutoRegEmbed vs. Without Contrastive Training Models without contrastive training are divided into two categories. The first is our own fair baseline model, which performs significantly worse than AutoRegEmbed, with an average performance 20% lower. This highlights the difficulty of untrained LLMs in directly generating high-quality embeddings. While some methods enhance the base model’s embeddings through prompt optimization, their improvements remain limited—even on a 13B-parameter model—and come with significant additional reasoning costs.

Method	Params	BIOSSES	SICK-R	STS12	STS13	STS14	STS15	STS16	STS17	STS22	STS-B	Avg.	Vol.
<i>Without Contrastive Training</i>													
LLaMA2-L	7B	63.29	65.10	45.26	70.83	56.69	62.48	63.27	49.76	-7.76	60.43	60.58 ₍₇₎ /56.91 ₍₁₀₎	0
LLaMA2-M	7B	65.96	60.01	44.76	64.13	48.66	62.33	63.16	64.35	27.59	53.50	56.65 ₍₇₎ /58.67 ₍₁₀₎	0
Mistral-v0.1-L	7B	54.40	67.40	48.54	64.27	54.89	65.05	62.12	48.22	13.71	63.05	60.76 ₍₇₎ /56.20 ₍₁₀₎	0
Mistral-v0.1-M	7B	67.46	62.42	50.11	66.45	52.60	61.93	65.02	71.28	29.79	54.19	58.96 ₍₇₎ /61.13 ₍₁₀₎	0
Echo-LLaMA2	7B	-	64.39	52.40	72.40	61.24	72.67	73.51	-	-	65.73	66.05 ₍₇₎ /-	0
Echo-LLaMA2	13B	-	70.27	59.36	79.01	69.75	79.86	76.75	-	-	71.31	72.33 ₍₇₎ /-	0
PromptEOL-LLaMA2	7B	-	69.64	58.81	77.01	66.34	73.22	73.56	-	-	71.66	70.03 ₍₇₎ /-	0
PromptEOL-Mistral	7B	-	69.47	63.08	78.58	69.40	77.92	79.01	-	-	75.77	73.32 ₍₇₎ /-	0
PromptEOL-LLaMA3	8B	-	60.88	68.94	78.57	68.18	76.75	77.16	-	-	72.83	71.90 ₍₇₎ /-	0
PromptEOL-LLaMA2	13B	-	68.23	56.19	76.42	65.42	72.73	75.21	-	-	67.96	68.83 ₍₇₎ /-	0
MetaEOL-LLaMA2	7B	-	74.86	64.16	81.61	73.09	81.11	78.94	-	-	77.96	75.96 ₍₇₎ /-	0
MetaEOL-Mistral	7B	-	75.13	64.05	82.35	71.57	81.36	79.85	-	-	78.29	76.09 ₍₇₎ /-	0
GenEOL-LLaMA2-Mistral	7B	-	78.08	70.24	83.43	78.03	81.79	80.65	-	-	80.46	78.95 ₍₇₎ /-	0
GenEOL-LLaMA2-ChatGPT	7B	-	78.71	70.78	83.28	77.75	82.10	80.45	-	-	79.83	78.99 ₍₇₎ /-	0
<i>Unsupervised Contrastive Training</i>													
LLM2Vec-LLaMA2 [♣]	7B	82.41	71.77	65.39	79.26	72.98	82.72	81.02	86.70	63.47	78.32	75.92 ₍₇₎ /76.41 ₍₁₀₎	~160,000
LLM2Vec-Mistral [♣]	7B	83.29	75.55	67.65	83.90	76.97	83.80	81.91	85.58	65.93	80.42	78.60 ₍₇₎ /78.50 ₍₁₀₎	~160,000
<i>Supervised Contrastive Training</i>													
NV-Embed [♣]	7.73B	85.59	82.80	76.22	86.30	82.09	87.24	84.77	87.42	69.85	86.14	83.65 ₍₇₎ /82.84 ₍₁₀₎	1,054,000
SFR-Embedding-2_R [♣]	7B	87.60	77.01	75.67	82.40	79.93	85.82	84.50	88.93	67.10	83.60	81.28 ₍₇₎ /81.26 ₍₁₀₎	~1,751,000
gte-Qwen2-7B-instruct [♣]	7.49B	81.37	79.16	79.53	88.97	83.87	88.48	86.49	88.75	67.16	86.81	84.76 ₍₇₎ /83.06 ₍₁₀₎	~791,000,000
LLM2Vec-LLaMA2 [♣]	7B	82.13	83.01	78.85	86.84	84.04	88.72	86.79	90.63	67.55	88.72	85.28 ₍₇₎ /83.73 ₍₁₀₎	544,000
LLM2Vec-Mistral [♣]	7B	85.24	83.70	78.80	86.37	84.04	88.99	87.22	90.19	67.68	<u>88.65</u>	85.40 ₍₇₎ /84.01 ₍₁₀₎	544,000
LLaMA2-L	7B	77.58	77.85	73.72	84.04	79.82	85.03	84.78	87.53	26.87	86.18	81.63 ₍₇₎ /76.34 ₍₁₀₎	50,000
LLaMA2-inbatch-L	7B	78.81	82.76	77.70	85.01	81.82	88.30	86.12	90.53	20.70	87.94	84.24 ₍₇₎ /77.97 ₍₁₀₎	50,000
LLaMA2-M	7B	75.65	78.92	74.12	84.17	80.00	85.63	83.28	85.65	65.09	86.27	81.77 ₍₇₎ /79.88 ₍₁₀₎	50,000
LLaMA2-inbatch-M	7B	78.09	83.17	77.10	82.82	80.53	87.40	84.43	90.02	64.59	87.18	83.23 ₍₇₎ /81.53 ₍₁₀₎	50,000
LLaMA2-inbatch-M	7B	77.43	82.26	77.95	84.90	82.06	87.22	86.43	88.22	66.42	86.12	83.85 ₍₇₎ /81.90 ₍₁₀₎	274,951
<i>Information Compression and Conditional Distribution Alignment</i>													
AutoRegEmbed-LLaMA2	7B	84.65	81.46	79.98	86.35	83.33	89.21	86.91	87.67	65.90	86.98	84.89 ₍₇₎ /83.24 ₍₁₀₎	50,000(16,382)
AutoRegEmbed-Mistral	7B	86.84	80.32	78.92	86.18	83.29	88.98	86.75	88.77	64.53	87.24	84.53 ₍₇₎ /83.18 ₍₁₀₎	50,000(16,382)
AutoRegEmbed-LLaMA2	7B	85.62	<u>83.87</u>	<u>79.60</u>	<u>87.36</u>	84.29	89.43	87.72	89.46	<u>67.78</u>	87.96	85.75 ₍₇₎ /84.31 ₍₁₀₎	274,951(16,382)
AutoRegEmbed-Mistral	7B	<u>87.48</u>	83.90	79.56	<u>87.64</u>	<u>84.11</u>	89.58	<u>87.46</u>	<u>89.87</u>	67.77	88.48	85.82 ₍₇₎ / 84.59 ₍₁₀₎	274,951(16,382)

Table 1: Results on STS tasks (Spearman correlation scaled by 100x). The parentheses in the **Avg.** column indicate the number of datasets used to compute the average. **Vol.** denotes the number of training triplets, while the numbers in brackets indicate the instruction fine-tuning data used by AutoRegEmbed during the information compression stage. The symbol “~” denotes an estimated value. “ ” represents our own fair baselines, and we apply a grid search to ensure optimal performance. “-L” and “-M” denote the hidden state of the last token and the average pooling of all token hidden states, respectively. The symbol [♣] indicates that not all data are open source. Bold indicates the best result, and underline indicates the second-best (suboptimal) result.

AutoRegEmbed vs. Unsupervised Contrastive Training LLM2Vec enhances existing LLMs using an unsupervised contrastive learning approach similar to SimCSE, leading to significant performance gains. Compared to the base model, the unsupervised version of LLM2Vec improves performance by over 15%. Although it utilizes almost 160,000 data samples, its performance remains 4.74% lower than AutoRegEmbed, demonstrating its lower efficiency.

AutoRegEmbed vs. Supervised Contrastive Training Supervised contrastive learning is the mainstream approach for building high-quality embedding models. We first compared SOTA methods that once ranked on the MTEB leaderboard, including NV-Embed, SFR-Embedding-2_R, gte-Qwen2-7B-instruct, and LLM2Vec. In terms of performance, AutoRegEmbed surpasses all leading methods across 10 STS datasets, outperforming LLM2Vec by a margin of 0.58. From a data effi-

ciency perspective, AutoRegEmbed achieves performance comparable to the previous SOTA models with just 66,382 training samples, whereas the latter requires tens of millions of triplets to reach peak performance. Additionally, previous SOTA models employ multi-task learning (e.g., retrieval and clustering), whose impact on STS performance remains unclear. To ensure a fair comparison, we use single-task contrastive learning as a baseline. Unlike traditional contrastive learning, AutoRegEmbed does not rely on in-batch negative samples. So we add two baselines to single-task contrastive learning that also exclude the in-batch negative sample strategy. As shown in Table 1, even under identical training data, AutoRegEmbed outperforms four different single-task contrastive learning, further validating its effectiveness.

4.3 Performance on retrieval tasks

Table 2 summarizes the performance of various baseline methods and AutoRegEmbed across three retrieval datasets. On MS MARCO, AutoRegEmbed outperforms most prior state-of-the-art (SOTA) models and ranks second only to gte-Qwen2-7B-instruct. It also achieves competitive results on the other two datasets, NFcorpus and SCIDOCS. This outcome is expected, as AutoRegEmbed is trained solely on MS MARCO, whereas previous SOTA models are trained on large-scale datasets spanning diverse distributions, which contributes to their strong generalization on NFcorpus and SCIDOCS.

Notably, AutoRegEmbed consistently outperforms LLaMA2-inbatch-M, which is trained on the same data, across all three datasets. This further confirms the effectiveness of our method in retrieval tasks, even without large-scale multi-domain training.

Method	MS MARCO	NFcorpus	SCIDOCS
LLM2Vec(Unsupervised)	18.81	26.81	10.00
LLM2Vec(Supervised)	41.45	40.33	21.05
SFR-Embedding-2_R	42.18	41.34	24.69
gte-Qwen2-7B-instruct	45.98	<u>40.60</u>	<u>23.48</u>
LLaMA2-inbatch-M	41.67	34.19	16.15
AutoRegEmbed	<u>42.49</u>	38.16	19.79

Table 2: Evaluation results on MS MARCO, NFcorpus, and SCIDOCS.

4.4 Ablation Study

To verify the effectiveness of AutoRegEmbed, we conducted an ablation study. First, we removed Conditional Distribution Alignment to evaluate its impact on model performance. Second, since Equation 2 was derived empirically in our previous work, we tested different variants of this equation to confirm that it remains the optimal choice. Different variants of Equation 2 include **Log_sigmoid**, which maps similarity to a logarithmic scale for integration with the exponential function e , as well as **KL divergence** and **JS divergence**, which quantify the distance between the conditional probability distributions of positive and negative sample embeddings in distinct ways.

Table 3 presents the ablation results. The experiments on different tasks indicate that Conditional Distribution Alignment improves performance by 9.17%, while Information Compression contributes a 16.99% improvement, demonstrating the effectiveness of both tasks. Additionally, experiments

Method	Avg.
AutoRegEmbed-LLaMA2	83.24 ₍₁₀₎
<i>Tasks</i>	
w/o Conditional Distribution Alignment	73.90 ₍₁₀₎
LLaMA2-L (Without Training)	56.91 ₍₁₀₎
<i>Equation 2</i>	
Log_sigmoid	82.93 ₍₁₀₎
KL divergence	79.82 ₍₁₀₎
JS divergence	79.02 ₍₁₀₎

Table 3: Ablation experiments of AutoRegEmbed. We conduct ablation and contrast experiments on various tasks and Equation 2 to demonstrate the effectiveness of AutoRegEmbed.

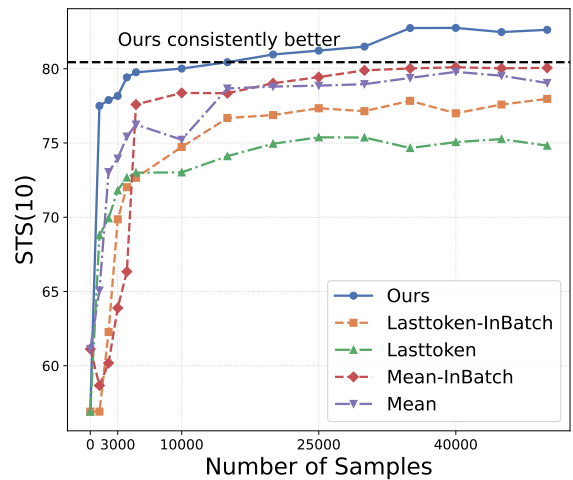


Figure 3: We evaluate the learning efficiency of our method against traditional contrastive learning on 10 STS datasets, comparing their performance under the same number of samples. Further details are provided in Appendix A.

on variants of Equation 2 reveal that, although using a logarithmic scale for similarity and employing KL or JS divergence to measure distribution distance are more intuitive approaches, they do not surpass the performance of the original loss function in Equation 2. Thus, Equation 2 can be regarded as a more effective loss function. The specific equations and more detailed analysis are provided in Appendix B.

4.5 Learning Efficiency

To verify that AutoRegEmbed is better suited for LLMs, we compare its performance with four contrastive learning baselines under the same training data. Figure 3 shows that as the training data increases, the performance of both AutoRegEmbed and other contrastive learning methods improves,

but AutoRegEmbed exhibits the fastest growth. Notably, with just 15,000 samples, AutoRegEmbed already surpasses the maximum performance of other contrastive learning models. The results indicate that AutoRegEmbed significantly outperforms the baseline models in learning efficiency.

5 Conclusions

To address the limitation that traditional contrastive learning does not adhere to the autoregressive nature of LLMs, we propose AutoRegEmbed—a novel contrastive learning method based on embedded conditional probability distributions. AutoRegEmbed ensures that LLM-generated embeddings capture global semantics while maintaining alignment and uniformity through information compression and conditional distribution alignment tasks. AutoRegEmbed achieves comparable performance to SOTA models with fewer training samples and superior learning efficiency.

6 Limitations

The primary advantage of AutoRegEmbed lies in its ability to effectively harness the power of large language models (LLMs) to construct robust and high-quality text embeddings. However, it is important to acknowledge several limitations of our approach.

AutoRegEmbed does not possess inherent mechanisms to filter or detect malicious or harmful content in the data it processes. While the model is capable of generating embeddings from a wide range of text inputs, it lacks the ability to evaluate the ethical or safety implications of the data. This makes it vulnerable to issues related to biased, offensive, or otherwise problematic content present in the training corpus. In cases where the training data contains harmful or discriminatory material, the embeddings generated by AutoRegEmbed may inadvertently carry forward these biases, potentially leading to unintended and undesirable outcomes when applied to real-world tasks.

To mitigate this risk, we recommend that users of AutoRegEmbed ensure that the training data is carefully curated, and ideally, filtered for harmful content. Additionally, users should be cautious when applying AutoRegEmbed to sensitive domains, where the generation of unsafe or biased embeddings could have significant consequences.

Acknowledgments

This work was supported by the Strategic Priority Research Program of the CAS under Grants No.XDB0680302, the National Natural Science Foundation of China (NSFC) under Grants No. 62276248, the Key Research and Development Program of Xinjiang Uyghur Autonomous Region Grant No. 2024B03026, the Beijing Nova Program under Grants No. 20250484765, and the Youth Innovation Promotion Association CAS under Grants No. 2023111.

References

- Eneko Agirre, Carmen Banea, Claire Cardie, Daniel M. Cer, Mona T. 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@NAACL-HLT 2015, Denver, Colorado, USA, June 4-5, 2015*, pages 252–263. The Association for Computer Linguistics.
- Eneko Agirre, Carmen Banea, Claire Cardie, Daniel M. Cer, Mona T. Diab, Aitor Gonzalez-Agirre, Weiwei Guo, Rada Mihalcea, German Rigau, and Janyce Wiebe. 2014. [Semeval-2014 task 10: Multilingual semantic textual similarity](#). In *Proceedings of the 8th International Workshop on Semantic Evaluation, SemEval@COLING 2014, Dublin, Ireland, August 23-24, 2014*, pages 81–91. The Association for Computer Linguistics.
- Eneko Agirre, Carmen Banea, Daniel M. Cer, Mona T. Diab, Aitor Gonzalez-Agirre, Rada Mihalcea, German Rigau, and Janyce Wiebe. 2016. [Semeval-2016 task 1: Semantic textual similarity, monolingual and cross-lingual evaluation](#). In *Proceedings of the 10th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2016, San Diego, CA, USA, June 16-17, 2016*, pages 497–511. The Association for Computer Linguistics.
- Eneko Agirre, Daniel M. Cer, Mona T. Diab, and Aitor Gonzalez-Agirre. 2012. [Semeval-2012 task 6: A pilot on semantic textual similarity](#). In *Proceedings of the 6th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2012, Montréal, Canada, June 7-8, 2012*, pages 385–393. The Association for Computer Linguistics.
- Eneko Agirre, Daniel M. Cer, Mona T. Diab, Aitor Gonzalez-Agirre, and Weiwei Guo. 2013. [*sem 2013 shared task: Semantic textual similarity](#). In *Proceedings of the Second Joint Conference on Lexical and Computational Semantics, *SEM 2013, June 13-14, 2013, Atlanta, Georgia, USA*, pages 32–43. Association for Computational Linguistics.

- Akari Asai, Timo Schick, Patrick S. H. Lewis, Xilun Chen, Gautier Izacard, Sebastian Riedel, Hannaneh Hajishirzi, and Wen-tau Yih. 2023. [Task-aware retrieval with instructions](#). In *Findings of the Association for Computational Linguistics: ACL 2023, Toronto, Canada, July 9-14, 2023*, pages 3650–3675. Association for Computational Linguistics.
- Parishad BehnamGhader, Vaibhav Adlakha, Marius Mosbach, Dzmitry Bahdanau, Nicolas Chapados, and Siva Reddy. 2024. [Llm2vec: Large language models are secretly powerful text encoders](#). *CoRR*, abs/2404.05961.
- Alimohammad Beigi, Zhen Tan, Nivedh Mudiam, Canyu Chen, Kai Shu, and Huan Liu. 2024. [Model attribution in llm-generated disinformation: A domain generalization approach with supervised contrastive learning](#). In *11th IEEE International Conference on Data Science and Advanced Analytics, DSAA 2024, San Diego, CA, USA, October 6-10, 2024*, pages 1–10. IEEE.
- Vera Boteva, Demian Gholipour Ghalandari, Artem Sokolov, and Stefan Riezler. 2016. [A full-text learning to rank dataset for medical information retrieval](#). In *Advances in Information Retrieval - 38th European Conference on IR Research, ECIR 2016, Padua, Italy, March 20-23, 2016. Proceedings*, volume 9626 of *Lecture Notes in Computer Science*, pages 716–722. Springer.
- Daniel M. Cer, Mona T. Diab, Eneko Agirre, Iñigo Lopez-Gazpio, and Lucia Specia. 2017. [Semeval-2017 task 1: Semantic textual similarity multilingual and crosslingual focused evaluation](#). In *Proceedings of the 11th International Workshop on Semantic Evaluation, SemEval@ACL 2017, Vancouver, Canada, August 3-4, 2017*, pages 1–14. Association for Computational Linguistics.
- Guangyao Chen, Siwei Dong, Yu Shu, Ge Zhang, Jaward Sesay, Börje Karlsson, Jie Fu, and Yemin Shi. 2024a. [Autoagents: A framework for automatic agent generation](#). In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence, IJCAI 2024, Jeju, South Korea, August 3-9, 2024*, pages 22–30. ijcai.org.
- Jianlv Chen, Shitao Xiao, Peitian Zhang, Kun Luo, Defu Lian, and Zheng Liu. 2024b. [BGE m3-embedding: Multi-lingual, multi-functionality, multi-granularity text embeddings through self-knowledge distillation](#). *CoRR*, abs/2402.03216.
- Xi Chen, Ali Zeynali, Chico Q. Camargo, Fabian Flöck, Devin Gaffney, Przemyslaw A. Grabowicz, Scott Hale, David Jurgens, and Mattia Samory. 2022. [Semeval-2022 task 8: Multilingual news article similarity](#). In *Proceedings of the 16th International Workshop on Semantic Evaluation, SemEval@NAACL 2022, Seattle, Washington, United States, July 14-15, 2022*, pages 1094–1106. Association for Computational Linguistics.
- Alexis Chevalier, Alexander Wettig, Anirudh Ajith, and Danqi Chen. 2023. [Adapting language models to compress contexts](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing, EMNLP 2023, Singapore, December 6-10, 2023*, pages 3829–3846. Association for Computational Linguistics.
- Arman Cohan, Sergey Feldman, Iz Beltagy, Doug Downey, and Daniel S. Weld. 2020. [SPECTER: document-level representation learning using citation-informed transformers](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, ACL 2020, Online, July 5-10, 2020*, pages 2270–2282. Association for Computational Linguistics.
- Tri Dao. 2024. [Flashattention-2: Faster attention with better parallelism and work partitioning](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Jingcheng Deng, Hengwei Dai, Xuewei Guo, Yuanchen Ju, and Wei Peng. 2022. [IRRGN: an implicit relational reasoning graph network for multi-turn response selection](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022, Abu Dhabi, United Arab Emirates, December 7-11, 2022*, pages 8529–8541. Association for Computational Linguistics.
- Jingcheng Deng, Liang Pang, Huawei Shen, and Xueqi Cheng. 2023. [Regavae: A retrieval-augmented gaussian mixture variational auto-encoder for language modeling](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023, Singapore, December 6-10, 2023*, pages 2500–2510. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL-HLT 2019, Minneapolis, MN, USA, June 2-7, 2019, Volume 1 (Long and Short Papers)*, pages 4171–4186. Association for Computational Linguistics.
- Hanxing Ding, Liang Pang, Zihao Wei, Huawei Shen, and Xueqi Cheng. 2024. [Retrieve only when it needs: Adaptive retrieval augmentation for hallucination mitigation in large language models](#). *CoRR*, abs/2402.10612.
- Zenghao Duan, Wenbin Duan, Zhiyi Yin, Yinghan Shen, Shaoling Jing, Jie Zhang, Huawei Shen, and Xueqi Cheng. 2025. [Related knowledge perturbation matters: Rethinking multiple pieces of knowledge editing in same-subject](#). In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 2: Short Papers)*,

- pages 363–373, Albuquerque, New Mexico. Association for Computational Linguistics.
- Luyu Gao and Jamie Callan. 2022. [Unsupervised corpus aware language model pre-training for dense passage retrieval](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2022, Dublin, Ireland, May 22-27, 2022, pages 2843–2853. Association for Computational Linguistics.
- Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021. [Simcse: Simple contrastive learning of sentence embeddings](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP 2021, Virtual Event / Punta Cana, Dominican Republic, 7-11 November, 2021*, pages 6894–6910. Association for Computational Linguistics.
- Tao Ge, Jing Hu, Lei Wang, Xun Wang, Si-Qing Chen, and Furu Wei. 2024. [In-context autoencoder for context compression in a large language model](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Pengcheng He, Xiaodong Liu, Jianfeng Gao, and Weizhu Chen. 2021. [Deberta: decoding-enhanced bert with disentangled attention](#). In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. [Lora: Low-rank adaptation of large language models](#). In *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net.
- Ting Jiang, Shaohan Huang, Zhongzhi Luan, Deqing Wang, and Fuzhen Zhuang. 2024. [Scaling sentence embeddings with large language models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2024, Miami, Florida, USA, November 12-16, 2024*, pages 3182–3196. Association for Computational Linguistics.
- Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick S. H. Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. [Dense passage retrieval for open-domain question answering](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing, EMNLP 2020, Online, November 16-20, 2020*, pages 6769–6781. Association for Computational Linguistics.
- Urvashi Khandelwal, Omer Levy, Dan Jurafsky, Luke Zettlemoyer, and Mike Lewis. 2020. [Generalization through memorization: Nearest neighbor language models](#). In *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26-30, 2020*. OpenReview.net.
- Chankyu Lee, Rajarshi Roy, Mengyao Xu, Jonathan Raiman, Mohammad Shoyebi, Bryan Catanzaro, and Wei Ping. 2024. [Nv-embed: Improved techniques for training llms as generalist embedding models](#). *CoRR*, abs/2405.17428.
- Yibin Lei, Di Wu, Tianyi Zhou, Tao Shen, Yu Cao, Chongyang Tao, and Andrew Yates. 2024. [Meta-task prompting elicits embeddings from large language models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2024, Bangkok, Thailand, August 11-16, 2024, pages 10141–10157. Association for Computational Linguistics.
- Bohan Li, Hao Zhou, Junxian He, Mingxuan Wang, Yiming Yang, and Lei Li. 2020. [On the sentence embeddings from pre-trained language models](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing, EMNLP 2020, Online, November 16-20, 2020*, pages 9119–9130. Association for Computational Linguistics.
- Chaofan Li, Zheng Liu, Shitao Xiao, Yingxia Shao, and Defu Lian. 2024a. [Llama2vec: Unsupervised adaptation of large language models for dense retrieval](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2024, Bangkok, Thailand, August 11-16, 2024, pages 3490–3500. Association for Computational Linguistics.
- Chaofan Li, Minghao Qin, Shitao Xiao, Jianlyu Chen, Kun Luo, Yingxia Shao, Defu Lian, and Zheng Liu. 2024b. [Making text embedders few-shot learners](#). *CoRR*, abs/2409.15700.
- Zehan Li, Xin Zhang, Yanzhao Zhang, Dingkun Long, Pengjun Xie, and Meishan Zhang. 2023. [Towards general text embeddings with multi-stage contrastive learning](#). *CoRR*, abs/2308.03281.
- Sheng Liu, Qiang Sheng, Danding Wang, Yang Li, Guang Yang, and Juan Cao. 2025. [Forewarned is forearmed: Pre-synthesizing jailbreak-like instructions to enhance llm safety guardrail to potential attacks](#). *Preprint*, arXiv:2508.20038.
- Tian Yu Liu, Matthew Trager, Alessandro Achille, Pramuditha Perera, Luca Zancato, and Stefano Soatto. 2024. [Meaning representations from trajectories in autoregressive models](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Xueguang Ma, Liang Wang, Nan Yang, Furu Wei, and Jimmy Lin. 2024. [Fine-tuning llama for multi-stage text retrieval](#). In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2024, Washington DC, USA, July 14-18, 2024*, pages 2421–2425. ACM.
- Rui Meng, Ye Liu, Shafiq Rayhan Joty, Caiming Xiong, Yingbo Zhou, and Semih Yavuz. 2024. [Sfr-embedding-2: Advanced text embedding with multi-stage training](#).

- Jesse Mu, Xiang Li, and Noah D. Goodman. 2023. [Learning to compress prompts with gist tokens](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.
- Niklas Muennighoff, Hongjin Su, Liang Wang, Nan Yang, Furu Wei, Tao Yu, Amanpreet Singh, and Douwe Kiela. 2024. [Generative representational instruction tuning](#). *CoRR*, abs/2402.09906.
- Niklas Muennighoff, Nouamane Tazi, Loïc Magne, and Nils Reimers. 2023. [MTEB: massive text embedding benchmark](#). In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics, EACL 2023, Dubrovnik, Croatia, May 2-6, 2023*, pages 2006–2029. Association for Computational Linguistics.
- Tri Nguyen, Mir Rosenberg, Xia Song, Jianfeng Gao, Saurabh Tiwary, Rangan Majumder, and Li Deng. 2016. [MS MARCO: A human generated machine reading comprehension dataset](#). In *Proceedings of the Workshop on Cognitive Computation: Integrating neural and symbolic approaches 2016 co-located with the 30th Annual Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain, December 9, 2016*, volume 1773 of *CEUR Workshop Proceedings*. CEUR-WS.org.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D. Manning, Stefano Ermon, and Chelsea Finn. 2023. [Direct preference optimization: Your language model is secretly a reward model](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.
- Nils Reimers and Iryna Gurevych. 2019. [Sentence-bert: Sentence embeddings using siamese bert-networks](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, EMNLP-IJCNLP 2019, Hong Kong, China, November 3-7, 2019*, pages 3980–3990. Association for Computational Linguistics.
- Weijia Shi, Sewon Min, Michihiro Yasunaga, Minjoon Seo, Richard James, Mike Lewis, Luke Zettlemoyer, and Wen-tau Yih. 2024. [REPLUG: retrieval-augmented black-box language models](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers), NAACL 2024, Mexico City, Mexico, June 16-21, 2024*, pages 8371–8384. Association for Computational Linguistics.
- Kaitao Song, Xu Tan, Tao Qin, Jianfeng Lu, and Tie-Yan Liu. 2020. [Mpnet: Masked and permuted pre-training for language understanding](#). In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*.
- Jacob Mitchell Springer, Suhas Kotha, Daniel Fried, Graham Neubig, and Aditi Raghunathan. 2024. [Repetition improves language model embeddings](#). *CoRR*, abs/2402.15449.
- Hongjin Su, Weijia Shi, Jungo Kasai, Yizhong Wang, Yushi Hu, Mari Ostendorf, Wen-tau Yih, Noah A. Smith, Luke Zettlemoyer, and Tao Yu. 2023. [One embedder, any task: Instruction-finetuned text embeddings](#). In *Findings of the Association for Computational Linguistics: ACL 2023, Toronto, Canada, July 9-14, 2023*, pages 1102–1121. Association for Computational Linguistics.
- Raghuveer Thirukovalluru and Bhuwan Dhingra. 2024. [Geneol: Harnessing the generative power of llms for training-free sentence embeddings](#). *CoRR*, abs/2410.14635.
- Aäron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. [Representation learning with contrastive predictive coding](#). *CoRR*, abs/1807.03748.
- Liang Wang, Nan Yang, Xiaolong Huang, Linjun Yang, Rangan Majumder, and Furu Wei. 2024a. [Improving text embeddings with large language models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2024, Bangkok, Thailand, August 11-16, 2024*, pages 11897–11916. Association for Computational Linguistics.
- Tongzhou Wang and Phillip Isola. 2020. [Understanding contrastive representation learning through alignment and uniformity on the hypersphere](#). In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13-18 July 2020, Virtual Event, volume 119 of Proceedings of Machine Learning Research*, pages 9929–9939. PMLR.
- Yining Wang, Jinman Zhao, and Yuri Lawryshyn. 2024b. [GPT-signal: Generative AI for semi-automated feature engineering in the alpha research process](#). In *Proceedings of the Eighth Financial Technology and Natural Language Processing and the 1st Agent AI for Scenario Planning*, pages 42–53, Jeju, South Korea. -.
- Zhichao Wang, Bin Bi, Shiva Kumar Pentiyala, Kiran Ramnath, Sougata Chaudhuri, Shubham Mehrotra, Zixu James Zhu, Xiang-Bo Mao, Sitaram Asur, and Na Claire Cheng. 2024c. [A comprehensive survey of LLM alignment techniques: Rlhf, rlai, ppo, DPO and more](#). *CoRR*, abs/2407.16216.
- Bohong Wu and Hai Zhao. 2022. [Sentence representation learning with generative objective rather than contrastive objective](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022, Abu Dhabi, United Arab Emirates, December 7-11, 2022*, pages 3356–3368. Association for Computational Linguistics.

Long Xia, Jun Xu, Yanyan Lan, Jiafeng Guo, and Xueqi Cheng. 2015. [Learning maximal marginal relevance model via directly optimizing diversity evaluation measures](#). In *Proceedings of the 38th International ACM SIGIR Conference on Research and Development in Information Retrieval, Santiago, Chile, August 9-13, 2015*, pages 113–122. ACM.

Shitao Xiao, Zheng Liu, Yingxia Shao, and Zhao Cao. 2022. [Retromae: Pre-training retrieval-oriented language models via masked auto-encoder](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022, Abu Dhabi, United Arab Emirates, December 7-11, 2022*, pages 538–548. Association for Computational Linguistics.

Shicheng Xu, Liang Pang, Huawei Shen, and Xueqi Cheng. 2025. [A theory for token-level harmonization in retrieval-augmented generation](#). In *The Thirteenth International Conference on Learning Representations, ICLR 2025, Singapore, April 24-28, 2025*. OpenReview.net.

Shicheng Xu, Liang Pang, Huawei Shen, Xueqi Cheng, and Tat-Seng Chua. 2024a. [Search-in-the-chain: Interactively enhancing large language models with search for knowledge-intensive tasks](#). In *Proceedings of the ACM on Web Conference 2024, WWW 2024, Singapore, May 13-17, 2024*, pages 1362–1373. ACM.

Shicheng Xu, Liang Pang, Mo Yu, Fandong Meng, Huawei Shen, Xueqi Cheng, and Jie Zhou. 2024b. [Unsupervised information refinement training of large language models for retrieval-augmented generation](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2024, Bangkok, Thailand, August 11-16, 2024*, pages 133–145. Association for Computational Linguistics.

Jinman Zhao, Zifan Qian, Linbo Cao, Yining Wang, Yitian Ding, Yulan Hu, Zeyu Zhang, and Zeyong Jin. 2025. [Role-play paradox in large language models: Reasoning performance gains and ethical dilemmas](#). *Preprint*, arXiv:2409.13979.

Jinman Zhao and Xueyan Zhang. 2024. [Large language model is not a \(multilingual\) compositional relation reasoner](#). In *First Conference on Language Modeling*.

Shengyao Zhuang, Xueguang Ma, Bevan Koopman, Jimmy Lin, and Guido Zuccon. 2024. [Promptreps: Prompting large language models to generate dense and sparse representations for zero-shot document retrieval](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, EMNLP 2024, Miami, FL, USA, November 12-16, 2024*, pages 4375–4391. Association for Computational Linguistics.

A Implementation details

AutoRegEmbed For the information compression task, we set the learning rate to $2e-5$, the batch size to 32, and train for 2 epoch. To represent the semantics of the input, we use 5 compressed tokens. For the conditional distribution alignment task, the learning rate is set to $5e-6$, with a batch size of 32 and 4 epochs. The temperature parameters τ and β are set to 0.05 and 0.1. For the above two tasks, we set the maximum token length of context, instruction, and target to 512. Furthermore, we employ the bfloat16 format, enable FlashAttention 2, and train on four A100-80G GPUs with DeepSpeed and Zero-2. The information compression task takes 20 minutes, while the conditional distribution alignment task, involving 50,000 samples, takes approximately 1 hour. In theory, our approach eliminates the need for in-batch negative sampling, which substantially reduces memory usage. In traditional contrastive learning, a batch size of 64 requires $64 * 64 = 4096$ pairwise computations. In contrast, our method requires only 64, reducing the number of pairwise comparisons per step by a factor of 64 and lowering memory consumption accordingly.

Fair Comparative Learning Baselines We train our own fair contrastive learning baseline based on the standard InfoNCE loss, with some code available in the FlagEmbedding repository¹. For baselines utilizing the in-batch negative sample strategy (LLaMA2-inbatch-L and LLaMA2-inbatch-M), we experimented with batch sizes of 128, 256, 512, and 1024, determining that 512 yields the best performance. Additionally, we ensure that gradients are propagated across different devices. For baselines that do not use the in-batch negative sample strategy, we set the batch size to 32, maintaining consistency with AutoRegEmbed. Regarding the learning rate, we tested $1e-5$, $5e-5$, $1e-4$, and $2e-4$, finding that $1e-4$ delivers the best results. All training data is consistent with AutoRegEmbed. We train the fair contrastive learning baseline using DeepSpeed and Zero-2 on four A100-80G GPUs in 1 hour.

B Variants of Equation 2

This section explores various possible modifications and extensions of Equation 2.

¹<https://github.com/FlagOpen/FlagEmbedding>

Log_sigmoid Given that most loss functions are logarithmic in nature, we can modify the similarity function in Equation 2 by replacing the sigmoid with a Log-Sigmoid function, resulting in a more interpretable formulation:

$$\begin{aligned}\mathcal{L}_{\text{CDA}} &= \mathbb{E}\left[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(d^+, d_i^-; q)/\tau}}\right], \\ S_1(q, d^+) &= -\log \sigma(\beta \left| \log \frac{p_{\theta_E}(d^+ | e_{q, I_{\text{next}}})}{p_{\theta_E}(d^+ | e_{d^+, I_{\text{self}}})} \right|), \\ S_2(d^+, d_i^-; q) &= -\log \sigma(\beta \log \frac{p_{\theta_E}(d^+ | e_{q, I_{\text{next}}})}{p_{\text{ref}}(d^+ | e_{q, I_{\text{next}}})} \\ &\quad - \beta \log \frac{p_{\theta_E}(d_i^- | e_{q, I_{\text{next}}})}{p_{\text{ref}}(d_i^- | e_{q, I_{\text{next}}})}).\end{aligned}$$

KL divergence We also experimented with replacing the difference in log probabilities with the KL divergence between the conditional probability distributions:

$$\begin{aligned}\mathcal{L}_{\text{CDA}} &= \mathbb{E}\left[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(q, d_i^-)/\tau}}\right], \\ S_1(q, d^+) &= -\sigma\left(\frac{1}{T} \sum_{t=1}^T \text{KL}(p_{\theta_E}(d_t^+ | d_{<t}^+, e_{d^+, I_{\text{self}}}), \right. \\ &\quad \left. p_{\theta_E}(d_t^+ | d_{<t}^+, e_{q, I_{\text{next}}}))\right), \\ S_2(q, d_i^-) &= -\sigma\left(\frac{1}{T} \sum_{t=1}^T \text{KL}(p_{\theta_E}(d_t^+ | d_{i, <t}^-, e_{d_i^-, I_{\text{self}}}), \right. \\ &\quad \left. p_{\theta_E}(d_t^+ | d_{<t}^+, e_{q, I_{\text{next}}}))\right).\end{aligned}$$

JS divergence In addition to KL divergence, we also employed JS divergence as a measure of distribution distance:

$$\begin{aligned}\mathcal{L}_{\text{CDA}} &= \mathbb{E}\left[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(q, d_i^-)/\tau}}\right], \\ S_1(q, d^+) &= -\sigma\left(\frac{1}{T} \sum_{t=1}^T \text{JS}(p_{\theta_E}(d_t^+ | d_{<t}^+, e_{d^+, I_{\text{self}}}), \right. \\ &\quad \left. p_{\theta_E}(d_t^+ | d_{<t}^+, e_{q, I_{\text{next}}}))\right), \\ S_2(q, d_i^-) &= -\sigma\left(\frac{1}{T} \sum_{t=1}^T \text{JS}(p_{\theta_E}(d_t^+ | d_{i, <t}^-, e_{d_i^-, I_{\text{self}}}), \right. \\ &\quad \left. p_{\theta_E}(d_t^+ | d_{<t}^+, e_{q, I_{\text{next}}}))\right).\end{aligned}$$

Why do KL divergence and JS divergence fail to achieve good performance in this setting? We hypothesize that the effectiveness of using log-odds ratios stems from its ability to directly supervise the generation probability of a specific label token, thus providing a strong and targeted learning

signal. In contrast, KL and JS divergence operate over the entire output distribution across the vocabulary. While these measures are theoretically well-founded for capturing distributional differences, they often suffer from instability in practice—especially in models like LLaMA2, which have large vocabularies (30,000 tokens). This instability weakens the gradient signal and ultimately degrades performance.

C Selection Analysis of the σ

To stabilize exponential operations, we apply the σ function to map distribution-based distances to a bounded range. Sigmoid and Tanh are typical choices. Table 4 presents their comparative performance under uniform settings.

Overall, under identical data and training settings, the Tanh variant slightly outperforms Sigmoid on 7 STS datasets, while Sigmoid shows a marginal advantage across all 10 datasets. Given the minimal performance gap, either function is a reasonable choice.

D Explanation of I_{next} and I_{self}

Embedding tasks can be broadly categorized into sentence-to-sentence (STS) and sentence-to-document (retrieval), each associated with distinct instructions, I_{next} and I_{self} . Based on prior studies (Jiang et al., 2024), the instructions used in this work are summarized in Table 5. Due to the symmetric nature of STS, both instructions are the same.

E Analysis of Alignment Strategies for Conditional Probability Distributions

The motivation behind our naive alignment strategy is as follows: I_{next} aims for q to generate d , while I_{self} guides d to generate itself. Therefore, we align the probability distribution of q generating d^+ with that of d^+ generating itself. Intuitively, when the conditional distributions produced by their embeddings $e_{q, I_{\text{next}}}$ and $e_{d^+, I_{\text{self}}}$ are closely matched, their embeddings are likely to be similar as well. Additionally, inspired by Direct Preference Optimization (DPO), we introduce negative sample feedback by encouraging $e_{q, I_{\text{next}}}$ to increase the distance to negative samples while decreasing the distance to positive ones.

However, the use of instructions enables more sophisticated strategies for aligning the probability distributions between q and d . To explore this, we

Method	BIOSSES	SICK-R	STS12	STS13	STS14	STS15	STS16	STS17	STS22	STS-B	Avg.
AutoRegEmbed-LLaMA2(Tanh)	83.97	80.75	80.58	86.92	83.19	88.98	86.96	85.80	65.91	85.98	84.77 ₍₇₎ / 82.90 ₍₁₀₎
AutoRegEmbed-LLaMA2(Sigmoid)	85.50	79.07	79.57	86.90	83.28	88.45	86.57	88.61	66.16	86.59	84.35 ₍₇₎ / 83.07 ₍₁₀₎

Table 4: Performance comparison of different σ functions for AutoRegEmbed under the same settings.

Retrieval Task	
I_{next}	Use one word to represent the query in a retrieval task. The word is: “
I_{self}	Use one word to represent the passage in a retrieval task. The word is: “
STS Task	
I_{next}	This sentence means in one word: “
I_{self}	This sentence means in one word: “

Table 5: Instructions used for Retrieval and STS tasks.

conducted a comprehensive and detailed analysis. We begin by introducing several alignment strategy configurations.

Strategy 1 In this strategy, we align the conditional probability distribution of generating q rather than d^+ . Specifically, we adjust the instructions to treat q as the anchor, aligning the probability distribution between q and d^+ . This is implemented by simply replacing the corresponding prompts, as detailed in Equation 3. For the STS task, where symmetry holds, I_{self} and I_{prev} are identical. For the asymmetric retrieval task, an example of I_{prev} could be: “Use a word to express the question that this article can answer:”.

$$\begin{aligned}
S_1(q, d^+) &= -\sigma(\beta \log \frac{p_{\theta_E}(q|e_{q, I_{\text{self}}})}{p_{\theta_E}(q|e_{d^+, I_{\text{prev}}})}) \\
S_2(d^+, d_i^-; q) &= -\sigma(\beta \log \frac{p_{\theta_E}(q|e_{d^+, I_{\text{prev}}})}{p_{\text{ref}}(q|e_{d^+, I_{\text{prev}}})}) \\
&\quad - \beta \log \frac{p_{\theta_E}(q|e_{d_i^-, I_{\text{prev}}})}{p_{\text{ref}}(q|e_{d_i^-, I_{\text{prev}}})}.
\end{aligned} \quad (3)$$

Strategy 2 We treat both q and d^+ as anchors and compute a weighted similarity between their

respective conditional distributions.

$$\begin{aligned}
S_1(q, d^+) &= -\frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(q|e_{q, I_{\text{self}}})}{p_{\theta_E}(q|e_{d^+, I_{\text{prev}}})}) \\
&\quad - \frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_{q, I_{\text{next}}})}{p_{\theta_E}(d^+|e_{d^+, I_{\text{self}}})}) \\
S_2(d^+, d_i^-; q) &= -\frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_{q, I_{\text{next}}})}{p_{\text{ref}}(d^+|e_{q, I_{\text{next}}})}) \\
&\quad - \beta \log \frac{p_{\theta_E}(d_i^-|e_{q, I_{\text{next}}})}{p_{\text{ref}}(d_i^-|e_{q, I_{\text{next}}})} \\
&\quad - \frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(q|e_{d^+, I_{\text{prev}}})}{p_{\text{ref}}(q|e_{d^+, I_{\text{prev}}})}) \\
&\quad - \beta \log \frac{p_{\theta_E}(q|e_{d_i^-, I_{\text{prev}}})}{p_{\text{ref}}(q|e_{d_i^-, I_{\text{prev}}})}.
\end{aligned} \quad (4)$$

Strategy 3 We can flexibly expand the range of negative samples by varying instructions—for example, by increasing the probability of positive samples generating themselves, while decreasing the probability of them generating negative samples.

$$\begin{aligned}
\mathcal{L}_{\text{CDA}} &= \mathbb{E}[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(d^+, d_i^-; q)/\tau}}], \\
S_1(q, d^+) &= -\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_{q, I_{\text{next}}})}{p_{\theta_E}(d^+|e_{d^+, I_{\text{self}}})}) \\
S_2(d^+, d_i^-; q) &= -\frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_{q, I_{\text{next}}})}{p_{\text{ref}}(d^+|e_{q, I_{\text{next}}})}) \\
&\quad - \beta \log \frac{p_{\theta_E}(d_i^-|e_{q, I_{\text{next}}})}{p_{\text{ref}}(d_i^-|e_{q, I_{\text{next}}})} \\
&\quad - \frac{1}{2}\sigma(\beta \log \frac{p_{\theta_E}(d^+|e_{d^+, I_{\text{self}}})}{p_{\text{ref}}(d^+|e_{d^+, I_{\text{self}}})}) \\
&\quad - \beta \log \frac{p_{\theta_E}(d_i^-|e_{d^+, I_{\text{self}}})}{p_{\text{ref}}(d_i^-|e_{d^+, I_{\text{self}}})}.
\end{aligned} \quad (5)$$

Strategy 4 Similar to Strategy 3, we can also increase the probability of negative samples generating themselves, while reducing the probability of them generating positive samples, as shown in Equation 6.

Method	Params	BIOSSES	SICK-R	STS12	STS13	STS14	STS15	STS16	STS17	STS22	STS-B	Avg.
Strategy 1	7B	86.31	81.97	81.05	86.27	82.80	88.51	86.31	86.43	64.23	86.35	84.75 ₍₇₎ / 83.02 ₍₁₀₎
Strategy 2	7B	83.34	81.87	79.82	84.11	81.97	89.11	86.39	88.79	66.11	86.90	84.32 ₍₇₎ / 82.84 ₍₁₀₎
Strategy 3	7B	86.03	83.20	78.35	81.54	81.27	87.80	87.04	89.16	64.55	86.28	83.64 ₍₇₎ / 82.52 ₍₁₀₎
Strategy 4	7B	84.10	80.74	78.09	83.79	81.77	88.01	86.14	89.64	64.14	87.16	83.67 ₍₇₎ / 82.36 ₍₁₀₎
AutoRegEmbed	7B	85.50	79.07	79.57	86.90	83.28	88.45	86.57	88.61	66.16	86.59	84.35 ₍₇₎ / 83.07 ₍₁₀₎

Table 6: Performance comparison of AutoRegEmbed under different alignment strategies on STS benchmarks.

Method	Params	BIOSSES	SICK-R	STS12	STS13	STS14	STS15	STS16	STS17	STS22	STS-B	Avg.
AutoRegEmbed($\tau=0.1, \beta=0.1$)	7B	85.50	79.07	79.57	86.90	83.28	88.45	86.57	88.61	66.16	86.59	84.35 ₍₇₎ / 83.07 ₍₁₀₎
AutoRegEmbed($\tau=0.02, \beta=0.1$)	7B	84.90	79.85	78.58	85.54	84.64	88.71	86.88	87.57	65.61	86.02	84.32 ₍₇₎ / 82.83 ₍₁₀₎
AutoRegEmbed($\tau=0.05, \beta=0.1$)	7B	84.65	81.46	79.98	86.35	83.33	89.21	86.91	87.67	65.90	86.98	84.89 ₍₇₎ / 83.24 ₍₁₀₎
AutoRegEmbed($\tau=0.2, \beta=0.1$)	7B	83.85	81.72	79.77	86.73	83.19	88.41	86.53	87.69	66.22	86.37	84.67 ₍₇₎ / 83.05 ₍₁₀₎
AutoRegEmbed($\tau=1.0, \beta=0.1$)	7B	81.23	80.57	77.63	83.90	81.92	87.08	85.75	88.18	63.91	85.95	83.26 ₍₇₎ / 81.61 ₍₁₀₎
AutoRegEmbed($\tau=0.1, \beta=0.2$)	7B	84.06	81.49	80.32	87.15	83.49	88.67	86.85	87.22	66.45	86.56	84.93 ₍₇₎ / 83.23 ₍₁₀₎
AutoRegEmbed($\tau=0.1, \beta=0.3$)	7B	84.50	79.56	79.45	86.62	83.55	87.78	86.01	89.77	65.63	86.13	84.16 ₍₇₎ / 82.90 ₍₁₀₎
AutoRegEmbed($\tau=0.1, \beta=0.4$)	7B	84.27	81.04	78.92	85.76	82.44	88.54	86.05	87.48	65.99	86.08	84.12 ₍₇₎ / 82.65 ₍₁₀₎
AutoRegEmbed($\tau=0.05, \beta=0.2$)	7B	84.31	79.65	79.59	84.16	81.95	89.53	87.54	89.37	66.78	87.48	84.27 ₍₇₎ / 83.04 ₍₁₀₎

Table 7: Ablation of temperature τ and alignment weight β on STS benchmarks. The Avg. column shows results over 7 and 10 datasets.

Table 6 presents the performance of the LLaMA2-7B model trained on the STS split of the MEDI dataset under various alignment strategies. Both τ and β are fixed at 0.1.

$$\begin{aligned}
\mathcal{L}_{\text{CDA}} &= \mathbb{E} \left[-\log \frac{e^{S_1(q, d^+)/\tau}}{e^{S_1(q, d^+)/\tau} + \sum_i e^{S_2(d^+, d_i^-; q)/\tau}} \right], \\
S_1(q, d^+) &= -\sigma(\beta \left| \log \frac{p_{\theta_E}(d^+ | e_{q, I_{\text{next}}})}{p_{\theta_E}(d^+ | e_{d^+, I_{\text{self}}})} \right|), \\
S_2(d^+, d_i^-; q) &= -\frac{1}{2} \sigma(\beta \log \frac{p_{\theta_E}(d^+ | e_{q, I_{\text{next}}})}{p_{\text{ref}}(d^+ | e_{q, I_{\text{next}}})} \\
&\quad - \beta \log \frac{p_{\theta_E}(d_i^- | e_{q, I_{\text{next}}})}{p_{\text{ref}}(d_i^- | e_{q, I_{\text{next}}})} \\
&\quad - \frac{1}{2} \sigma(\beta \log \frac{p_{\theta_E}(d_i^- | e_{d_i^-, I_{\text{self}}})}{p_{\text{ref}}(d_i^- | e_{d_i^-, I_{\text{self}}})} \\
&\quad - \beta \log \frac{p_{\theta_E}(d^+ | e_{d_i^-, I_{\text{self}}})}{p_{\text{ref}}(d^+ | e_{d_i^-, I_{\text{self}}})})
\end{aligned} \tag{6}$$

The experimental results reveal a consistent trend: the simpler the modification to the original loss function (Equation 2), the better the performance. Specifically, **Strategy 1** alters the generation target from d^+ to q . Given the inherent symmetry between the two, this change leads to no significant performance difference. **Strategy 2** computes a weighted average of the generation probabilities for q and d^+ ; however, despite its increased computational cost, it does not improve performance. **Strategies 3 and 4** introduce more complex alignment designs involving positive and negative sam-

ples, but these result in inferior performance.

F Analysis of Temperature Coefficients τ and β

The reasons for introducing these two temperature coefficients are as follows:

- We introduce the temperature coefficient τ to align our loss function more closely with the InfoNCE objective. Intuitively, τ controls the strength of separation between positive and negative samples—the lower the value, the sharper the contrast enforced by the loss.
- The temperature coefficient β is inspired by the DPO (Direct Preference Optimization) loss, which uses a similar scaling factor to modulate the influence of preference differences—also expressed as log probability ratios. In our setting, β controls the sensitivity of the loss to differences between distributions.

To empirically assess the impact of these hyperparameters, we conducted a grid search using the LLaMA2-7B model and the STS split from the MEDI dataset. For τ , we adopted values commonly used in contrastive learning: 0.02, 0.05, 0.1, 0.2, 1.0. For β , we selected values based on prior DPO studies: 0.1, 0.2, 0.3, 0.4. After systematic evaluation, we identified the best-performing (τ, β) pair.

To empirically assess the impact of these hyperparameters, we conducted a grid search us-

ing the LLaMA2-7B model and the STS split from the MEDI dataset. For τ , we adopted values commonly used in contrastive learning: 0.02, 0.05, 0.1, 0.2, 1.0. For β , we selected values based on prior DPO studies: 0.1, 0.2, 0.3, 0.4. After systematic evaluation, we observe that the model performs better when $\tau = 0.05$ or $\beta = 0.2$. Furthermore, except for a few extreme cases (e.g., $\tau = 1.0$ or $\beta = 0.4$), the model’s performance remains relatively stable across different settings, which demonstrates the robustness of our method. When $\tau = 0.05$ and $\beta = 0.2$ are applied simultaneously, the model does not exhibit improved performance. This suggests that the effects of these two temperature parameters may not be orthogonal.