

One Model to Train them All: Hierarchical Self-Distillation for Enhanced Layer Representations in Embeddings

Anonymous ACL submission

Abstract

Deploying language models often requires handling model size vs. performance trade-offs to satisfy downstream latency constraints while preserving the model’s usefulness. Model distillation is commonly employed to reduce model size while maintaining acceptable performance. However, distillation can be inefficient since it involves multiple training steps. In this work, we introduce MODULARSTARENCODER, a modular multi-exit encoder with 1B parameters, useful for multiple tasks within the scope of code retrieval. MODULARSTARENCODER is trained with a novel self-distillation mechanism that significantly improves lower-layer representations—allowing different portions of the model to be used while still maintaining a good trade-off in terms of performance. Our architecture focuses on enhancing text-to-code and code-to-code search by systematically capturing syntactic and semantic structures across multiple levels of representation. Specific encoder layers are targeted as exit heads, allowing higher layers to guide earlier layers during training. This self-distillation effect improves intermediate representations, increasing retrieval recall at no extra training cost. In addition to the multi-exit scheme, our approach integrates a repository-level contextual loss that maximally utilizes the training context window, further enhancing the learned representations. We also release a new dataset constructed via code translation, seamlessly expanding traditional text-to-code benchmarks with code-to-code pairs across diverse programming languages. Experimental results highlight the benefits of self-distillation through multi-exit supervision.

1 Introduction

Large language models (LLMs) have significantly impacted the field of natural language processing, demonstrating remarkable performance across various applications (Niu et al., 2023). However, the

amount of computation required to operate state-of-the-art models poses significant challenges for the large-scale deployment of these models.

To mitigate these challenges, the research community has explored several model strategies to reduce the operational cost of LLMs without sacrificing their effectiveness. A prominent technique in model compression is quantization (Jacob et al., 2017; Lin et al., 2023; Egiazarian et al., 2024), which involves the reduction of numerical precision in the model’s parameters. Quantization effectively decreases memory requirements and enhances inference speed, facilitating the deployment of large language models in resource-constrained environments. Concurrently, knowledge distillation has emerged as a powerful technique whereby a smaller “student” model is trained to emulate the behavior of a larger “teacher” model, as evidenced by works such as DISTILBERT (Sanh et al., 2019) and TINYBERT (Jiao et al., 2019). Additionally, pruning methods selectively eliminate less influential weights or neurons, further reducing model complexity and aiming to preserve performance (Han et al., 2015).

Recent efforts have increasingly focused on developing efficient architectures requiring fewer parameters. Model families such as LLaMA (Dubey et al., 2024), Qwen (Hui et al., 2024), Mistral (Jiang et al., 2023), and SmolLM (Allal et al., 2025) exemplify a paradigm shift towards smaller, more accessible architectures. These model families are deployed at various resolutions—ranging from lightweight variants optimized for heavily resource-constrained environments to larger versions that retain competitive performance.

In parallel, advancements in dynamic inference strategies have introduced mechanisms that further optimize computational efficiency. Techniques like multi-exit networks enable early predictions at intermediate layers, reducing unnecessary computations. For instance, early-exit architectures

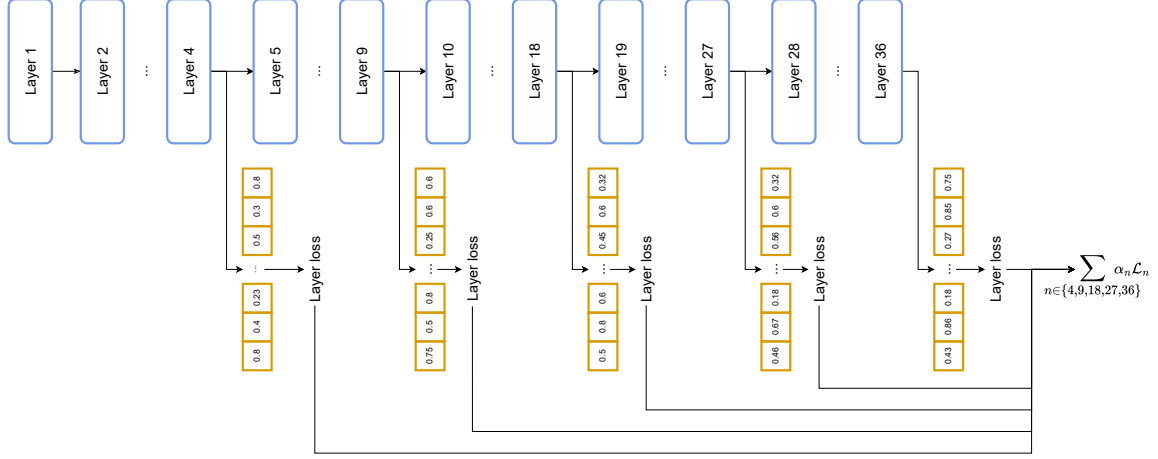


Figure 1: Overview of our multi-exit self-distillation encoder, shown here with exit heads at selected layers (e.g., Layers 4, 9, 18, 27, and 36). Each exit head predicts an output embedding and adds a “layer loss,” contribution weighted by a coefficient α_i , summed into the overall objective $\mathcal{L}$.

such as BranchyNet (Teerapittayanon et al., 2017) dynamically balance computation and accuracy by allowing predictions before full model execution. Similarly, Matryoshka representation learning (Kusupati et al., 2022) extends this idea to embeddings, introducing a loss function that yields multi-granular representations. This approach allows downstream tasks to adjust computational complexity by pruning embedding dimensionality, further contributing to efficient model deployment.

Building on these principles, we propose MODULARSTARENCODER, a modular multi-exit encoder architecture that integrates a novel intra-model self-distillation mechanism. In our design, specific intermediate layers are supervised by both the primary task loss and auxiliary distillation losses on specific exit heads, encouraging lower layers to learn better representations by mimicking the outputs of higher layers. We apply a shared embedding head comprising a masked language modeling head and an in-context classification head across a chosen subset of layers. We then fine-tuned the model with different projection heads for each exit point. We reached state-of-the-art results on multiple retrieval tasks (such as code-to-code and text-to-code), fine-tuning *one* single modular model that can be sliced depending on the end-user computational constraints.

Our contributions are as follows:

- We introduce a self-distillation framework that enables training multiple model resolutions within a unified layer stack, reducing redundancy and improving scalability. We believe

this approach can significantly affect LLM training pipelines that depend on multiple model distillations.

- We train and plan to release MODULARSTARENCODER, which consists of a pre-trained and fine-tuned encoder: The former is a modular pre-trained encoder with up to 1 billion parameters and four exit points, allowing users to perform multiple exit fine-tuning depending on downstream tasks. The latter is a fine-tuned encoder for various retrieval tasks, enabling the end user to select the model size that meets their memory and computational constraints.
- We plan to release SYNTHCODE2CODE2NL a new dataset constructed via code translation, expanding popular text-to-code datasets across diverse programming languages with code-to-code pairs. SYNTHCODE2CODE2NL comprises 1.071.367 triplets of natural language-code-code.

2 Methodology

2.1 Dataset

In the pre-training phase, we leveraged The Stack V2 dataset (Lozhkov et al., 2024), a large open-source code repository. Since many samples exceed typical context lengths, we split any tokenized sample that surpassed 1536 tokens — 75% of our 2048 token limit — while preserving each repository’s original structure. This strategy ensures that our model sees manageable code segments.

Table 1: SYNTHCODE2CODE2NL details: Average character count and sample size per language for the CodeSearchNet dataset and the synthesized portion obtained through translation.

Language	CSN samples	CSN avg. char	Synth. samples	Synth. avg. char
English	1.071.367	180	-	-
PHP	280.706	514	116.967	579
Python	274.454	474	117.374	518
Go	234.089	350	124.125	541
Java	282.118	505	116.098	707
C++	-	-	141.956	938
Ruby	-	-	158.494	456
C	-	-	136.365	1029
JavaScript	-	-	159.988	557

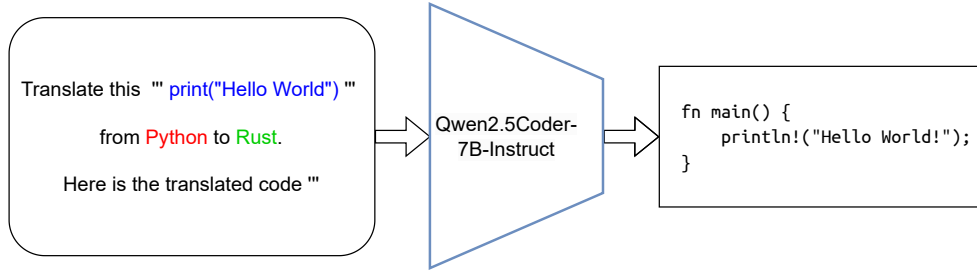


Figure 2: Prompt provided to Qwen2.5-Coder-7B-Instruct for translating a given code snippet (`print("Hello World")`) in the example) from a source programming language (Python) to a target one (Rust).

We created SYNTHCODE2CODE2NL, a dataset for fine-tuning that supports text-to-code and code-to-code search. Using the popular CODESEARCHNET (Husain et al., 2019) as a seed dataset and selecting popular programming languages (Python, Java, Go, and PHP), we augmented it by transpiling available code snippets onto other languages.

During the preprocessing phase of SYNTHCODE2CODE2NL, we deduplicated the dataset using both the CodeSearchNet code column and the synthesized code column. After a manual inspection, we discovered that both columns contained code snippets that differed only in identifiers or function arguments, as several tasks were semantically identical but paraphrased with different parameter requirements (e.g., two identical paraphrased tasks were asking for opening a socket on a different port). During the data near deduplication phase, we relied on Locality Sensitive Hashing (LSH) with a Jaccard similarity threshold of 0.7 and 256 permutations, analyzing character-level 5-grams.

To generate semantically similar code snippets for code-to-code search, we translated each snippet into a different language randomly sampled from Go, Ruby, Python, Java, C++, PHP, C, JavaScript. We prompted the QWEN2.5-CODER-

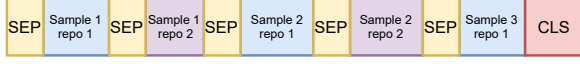
7B-INSTRUCT model with the source code, the name of the source language, and the name of the target language (see fig. 2). During code translation, we choose the token with the highest probability as output (greedy search) to prevent semantic discrepancies.

This process yielded pairs of code snippets in distinct languages tied to the same natural language description. As a result, every sample in the fine-tuning dataset includes a natural language description and two code snippets from distinct languages. SYNTHCODE2CODE2NL contains 1.071.367 samples where in the first code column we directly processed code snippets from CodeSearchNet, including Python, Java, PHP, and Go. The third column, artificially synthesized via code translation, includes Go, Ruby, JavaScript, Python, C++, PHP, C, and Java code snippets. Motivated by preliminary experiments indicating that near-deduplication on the code columns improved model performance, we then near-deduplicated the code columns of the datasets. Table 1 shows the average number of characters per language in SYNTHCODE2CODE2NL, we emphasize that synthesized data is significantly longer than human-written code. Appendix A provides examples of code translation.

Positive example



Negative example



```
def build_examples(repo_1, repo_2):
    is_negative = random(0, 1)
    input = empty_string

    while not empty(repo_1):
        input += <sep_token>

        # Positive case
        pos_sample = sample_random_snippet(repo_1)
        if len(input) + len(pos_sample) > context_length - 1:
            break
        input += pos_sample

        if is_negative and not empty(repo_2):
            # Negative (snippets from different repos)
            input += <sep_token>
            neg_sample = sample_random_snippet(repo_2)
            if len(input) + len(neg_sample) > context_length - 1:
                break
            input += neg_sample
        input += <cls_token>

    return input
```

Figure 3: On the left side the illustration of the in-context loss framework, where samples from different repositories are concatenated. Positive examples share the same repository context, while negative examples come from different repositories. On the right side, in-context loss framework pseudocode.

2.2 Architecture

We built MODULARSTARENCODER on top of STARCORDER-2 (Lozhkov et al., 2024), reducing its size from 15B to 1B parameters in bfloat16. Our architecture consists of 36 hidden layers, each with 16 attention heads and 4 key-value heads, using Grouped Query Attention (GQA) (Ainslie et al., 2023). The model employs Rotary Positional Encoding (RoPE) (Su et al., 2021) with a base period $\theta = 10^{-6}$ and features a hidden dimensionality of 1024 with an intermediate size of 12,288.

To enhance efficiency, we replaced the causal self-attention layers with bidirectional self-attention, similar to recent work that “modernized” BERT (Warner et al., 2024). Unlike STARCORDER-2, which uses sliding window attention, we opted for full attention to ensure greater modularity, avoiding the receptive field constraints of sliding window mechanisms (Zhu et al., 2021). Additionally, we extended the maximum input length to 2048 tokens, accommodating longer code snippets compared to prior code encoders such as STARENCORDER (Li et al., 2023).

Finally, our implementation integrates FLASHATTENTION V2 (Dao, 2023) for faster inference. Table 2 summarizes the architectural details.

2.3 Pre-training

We pre-trained MODULARSTARENCODER with a batch size of 3.99M tokens for 245,000 training steps, processing 1T tokens. The pre-training and fine-tuning were conducted on 512 NVIDIA Ampere (64GB) GPUs using the Leonardo super-computer (Turisini et al., 2023), requiring 450,000 GPU working hours.

To enable both token-level and snippet-level embeddings after pre-training, we employed a multi-objective pre-training strategy that combined two losses, as detailed in section 2.3.1 and section 2.3.2. The pre-training was performed on THESTACKV2, whose context length analysis revealed an average of ≈ 630 tokens per code snippet. As described in section 2.3.1, we concatenated multiple snippets to facilitate our multi-loss methodology, allowing our in-context classification loss to expand the average context window to ≈ 1300 tokens, reaching the maximum context length 20

We optimized the model using the AdamW optimizer with β_1 set to 0.9, β_2 to 0.95, ϵ to $1e-6$, and a weight decay of $1e-1$. The learning rate was initialized at $6.24e-4$ and scheduled using a multi-step learning rate decay (Bi et al., 2024) with 4,000 warmup steps. The learning rate was reduced at 120,000, 185,000, 220,000, 230,000, and 240,000 training steps, applying a decay factor of 0.36, and from step 185,000 onward, further reduced by factors of 0.1, 0.031, 0.01, and 0.001. Table 2 summarizes the hyperparameters for architecture, pre-training, and fine-tuning.

2.3.1 Masked Language Modeling and In Context Classification

The training objectives of BERT (Feng et al., 2020), specifically Masked Language Modeling (MLM) and Next Sentence Prediction (NSP), have become a de facto standard. However, The NSP loss constrains the context window length to the sentence length, leading to too many padding tokens and redundant computation (Zeng et al., 2022), and has been shown to not yield significant benefits after fine-tuning (Warner et al., 2024; Aroca-Ouellette

Table 2: Hyperparameters for Architecture, Pre-training, and Fine-tuning

Architecture	
Hyperparameter	Value
Model size	1B parameters
Precision	bfloat16
Hidden layers	36
Attention heads	16
Hidden dimensionality	1024
Positional encoding	RoPE ($\theta = 10^{-6}$)
Context length	2048
Attention mechanism	Grouped-Query Attention
Attention pattern	Bi-directional
Pre-training	
Batch size	3.99M tokens
Pretraining steps	245.000
Pretraining Tokens	1T
Loss function	MLM + In-Context loss
Multi-layer loss	yes
Optimizer	AdamW
Weight decay	1e-1
Initial learning rate	6.24e-4
Learning rate schedule	Multi-step
Warmup steps	4000
Fine-tuning	
Dataset size	635.404 samples
Fine-tuning steps	20.000
Loss function	CLIP loss
Multi-layer loss	yes
Batch size	2048
Learning rate	1.0e-5
Temperature parameter	10.0
Hardware (Pre-training + fine-tuning)	
GPUs	512 NVIDIA Ampere (64GB)
Overall Training hours	450.000

and Rudzicz, 2020). Given that the average number of tokens per data sample in Stack v2 is 630, a large context window of 2048 results in substantial padding, making long-context training inefficient. While Wang et al. (2023) demonstrated the advantages of training LLMs with multiple objectives, we revisited the NSP loss and introduced an in-context classification (ICC) objective. We hypothesize that predicting whether multiple code snippets belong to the same context (in our case, the same repository) can enhance semantic search performance while allowing efficient concatenation of multiple code fragments. Our final training objective is the summation of two losses: (1) MLM loss and (2) ICC loss: $\mathcal{L} = \mathcal{L}_{MLM} + \mathcal{L}_{ICC}$.

In $\mathcal{L}_{MLM}$, a certain percentage of tokens are randomly masked and predicted using a classification head. Following Zhang et al. (2024), we adopt a 15% masking rate with the standard 80-10-10 token replacement strategy (Feng et al., 2020). The secondary objective, $\mathcal{L}_{ICC}$, determines

whether randomly concatenated inputs (separated by a `< SEP >` token) originate from the same repository (see fig. 3). Each concatenated sample has a 50% probability of containing source code from different repositories. This approach increases input density—reducing padding by expanding the average input length from 630 to 1300 tokens—and potentially enhances cross-language understanding. Since repositories are inherently modular and often contain files written in multiple languages, learning from repository-level context may improve inter-language generalization.

2.3.2 Multi-layer Loss

To achieve layer-wise modularity in transformer architectures, we apply the previously introduced loss (section 2.3.1) across a selected set of layers, sharing classification heads (masked language modeling and in-context classification) while incorporating a positional embedding of the layer index. The total loss is computed as the sum of individual layer losses, weighted by a factor α to prioritize deeper layers: $\mathcal{L} = \sum_{i \in \iota} \mathcal{L}_i \cdot \alpha$ where $\alpha = i/|I|$ and $I = \{1, \dots, 36\}$ represents all layers, and the selected subset $\iota = \{4, 9, 18, 27, 36\}$ defines the layers where the loss is applied. The selected subset was chosen to enable four model variants equally spaced in depth (9, 18, 27, 36) along with an additional “tiny” version (4) to see the model performance in a lower number of parameters set. This approach allows for flexible model deployment, enabling adaptive layer pruning while maintaining performance trade-offs.

2.4 Fine-tuning

Following (Su et al., 2023), we fine-tune a single model for both text-to-code and code-to-code retrieval using instruction prompting. The optimization objective combines CLIP loss (Radford et al., 2021) with a multi-layer loss (details in 2.3.2).

To enhance representation learning, we replace the standard single-head projection with five distinct projection heads, applied at different exit points of the pre-trained model (layers 4, 9, 18, 27, and 36). We used a batch of 2048 elements, ensuring that text-to-code and code-to-code were equally distributed across the batch.

We performed data augmentation by randomly replacing frequently occurring words (appearing more than twice and having at least three characters) with random strings. We applied the augmentation exclusively to code snippets in 30% of cases,

Table 3: Performance of different models on text-to-code with CodeSearchNet using **codeXGLUE**. We reported the results presented in codet5plus, unixcoder and modernBERT (Wang et al., 2023; Guo et al., 2022; Warner et al., 2024).

Model	CodeSearchNet							
	Ruby	JS	Go	Python	Java	PHP	avg. MRR	avg. NDCG
MODULARSTARENCODER	74.1	74.0	82.5	92.5	78.7	84.5	81.0	84.2
Codet5+ 770M	78.0	71.3	92.7	75.8	76.2	70.1	77.4	-
OpenAI text-embedding-3-large	84.7	85.3	95.9	99.8	90.1	95.6	91.9	93.3
Unixcoder	74.0	68.4	91.5	72.0	72.6	67.6	74.4	-
ModernBERT-large	-	-	-	-	-	-	-	59.5

leaving natural language descriptions unchanged. After conducting a grid search, we selected $1.0e-5$ as the learning rate, maintained throughout the fine-tuning process, and set the temperature parameter at 10.0.

2.5 Evaluation

We evaluated MODULARSTARENCODER on both text-to-code and code-to-code retrieval tasks using the CODEXGLUE benchmark (Lu et al., 2021). For text-to-code retrieval, we employed the CODESEARCHNET dataset, where the goal is to retrieve the most relevant code snippet given a natural language query. Specifically, the query corresponds to a documentation comment, and the model is tasked with ranking the correct code snippet among 999 distractor snippets (Husain et al., 2019). This setup assesses the model’s ability to learn meaningful cross-modal representations between code and natural language.

For code-to-code retrieval, we relied on two datasets from CODEXGLUE: the Code Translation (CT) benchmark and POJ-104. The Code Translation dataset consists of semantically equivalent code snippets in different programming languages, and we framed the task as cross-language code retrieval rather than translation. In this setting, given a Java code snippet as a query, the model retrieves the corresponding C# implementation, testing its capability to capture cross-lingual semantic similarities between functionally equivalent programs.

In contrast, the POJ-104 dataset contains C/C++ programs that solve the same problem but with different implementations, making it a suitable benchmark for assessing retrieval of semantically similar code across diverse implementations. This setup evaluates the model’s capacity to generalize across structural variations while preserving semantic equivalence.

Table 4: Performance of different models on Code Translation (CT) and POJ104 for code-to-code search with codeXGLUE dataset.

Model	CT	POJ104
	MRR	mAP
MODULARSTARENCODER	98.9	56.5
Codet5+ 110M-embedding	98.4	24.5
OpenAI text-embedding-3-large	98.8	82.9
Unixcoder	97.6	41.0
ModernBERT-large	93.1	27.3

3 Results and Discussion

3.1 Benchmarks

Table 3 presents the results for CodeSearchNet (t2c) task in terms of Mean Reciprocal Rank (MRR) for each single language, average NDCG and average MRR. Results for Unixcoder, ModernBERT, and CodeT5+ are reported from the original papers (Guo et al., 2022; Warner et al., 2024; Wang et al., 2023). On CODESEARCHNET, MODULARSTARENCODER achieves an MRR of 81.0, outperforming CODET5+ (Wang et al., 2023) (770M), UNIXCODER (Guo et al., 2022), and MODERNBERT-LARGE (Warner et al., 2024). Although MODERNBERT-LARGE reports only average NDCG results (59.5). The only encoder that surpasses MODULARSTARENCODER is OpenAI’s text-embedding-3-large.

Table 4 presents results from both POJ104 and CT dataset reported respectively in MRR for code translation and mean average precision for POJ104. For Code Translation (CT), we report MRR on Java-to-C# retrieval. MODULARSTARENCODER reaches the best performance among the tests. We decided to replicate the experiments in a zero-shot setting for code-to-code tasks, as our model does not integrate POJ104 and code translation datasets in the training set.

Referring to Table 4, on the POJ104 dataset in

Table 5: Performance comparison of MODULARSTARENCODER model layers and baseline fine-tuned models on the CodeSearchNet benchmark. The table displays the overall retrieval performance measured by Mean Reciprocal Rank (MRR). We refer to the model fine-tuned with multiple exit points simultaneously as *self-distilled*; the others are fine-tuned singularly for each exit point.

Model	Size	CodeSearchNet						
		Ruby	Javascript	Go	Python	Java	PHP	avg. MRR
Layer-4	110M	59.5	61.3	72.1	86.2	68.2	75.5	70.5
Layer-4 (self-distilled)		62.2	64.7	74.8	88.1	71.4	78.0	73.2
Layer-9	250M	64.9	65.7	74.3	87.3	72.0	78.8	73.8
Layer-9 (self-distilled)		67.6	69.4	78.9	90.2	75.5	82.3	77.3
Layer-18	500M	73.8	73.5	82.4	92.1	78.4	84.0	80.7
Layer-18 (self-distilled)		74.1	74.0	82.5	92.5	78.7	84.5	81.0
Layer-27	750M	72.3	71.8	80.8	90.8	76.9	82.3	79.1
Layer-27 (self-distilled)		73.2	73.3	81.7	92.1	77.8	83.8	80.3
Layer-36	1B	72.3	72.9	80.7	91.5	77.1	82.9	79.5
Layer-36 (self-distilled)		73.5	72.6	80.5	91.4	76.9	82.7	79.6

zero-shot, MODULARSTARENCODER achieves an mAP of 0.57, which is state-of-art between open-sourced models, even though it is significantly behind OpenAI text-embedding-3-large. However, direct comparison with OpenAI text-embedding-3-large remains challenging because it is closed-source, and details such as model size, training methodology, or potential data contamination are undisclosed.

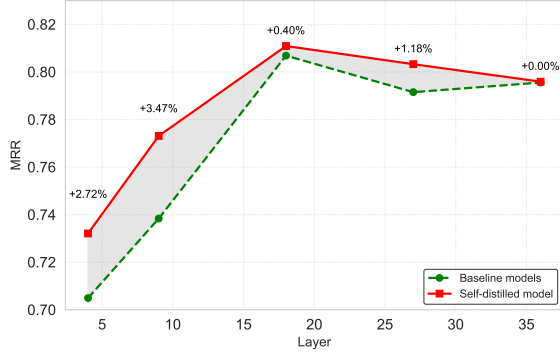
3.2 Ablation Study

We conducted an ablation study by fine-tuning singularly each exit point (also starting from MODULARSTARENCODER, pre-trained) and pruning the subsequent layers (e.g., for the baseline on layer 18, we retain only the first 18 layers and fine-tune the model using just one projection head on that layer). The multi-exit model consistently outperforms the single-exit baseline, indicating that *lower-level layers benefit from training signals propagated from deeper layers*. This behavior is highlighted in Table 5, where our model, indicated by *self-distilled*, outperforms all the single exit baselines consistently. This finding underscores a promising new direction in self-distillation for large-scale code and text models, enabling high performance even in more compact configurations. Moreover, Figure 4 illustrates that MODULARSTARENCODER maintains robust performance from layers 18 to 36, allowing users to scale down the network to match their memory, computational, or latency constraints while preserving strong retrieval accuracy.

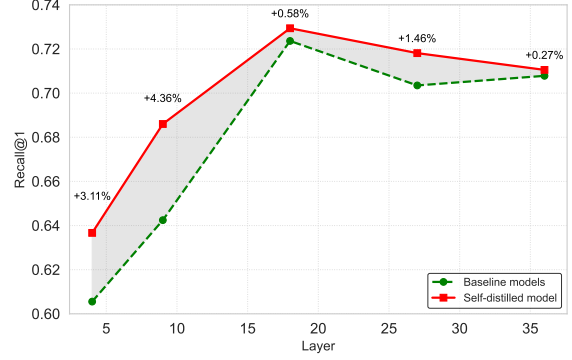
4 Related work

Since the introduction of ELMo (Peters et al., 2018), deep contextual information has enhanced generating embeddings for textual retrieval or classification, reaching state-of-the-art results in several tasks. BERT (Feng et al., 2020) followed those findings, adapting the Transformer architecture (Vaswani et al., 2017) to enable a bi-directional representation with two different training objectives, namely the masked language modeling and the next sentence prediction losses. (Lan et al., 2019; Liu et al., 2019) adapted the BERT architecture to obtain an enhanced pre-trained model by removing or modifying the NSP, focusing on pre-training data or hyperparameters optimization. More recently, modernBERT (Warner et al., 2024) tied the gap between modern decoders (Jiang et al., 2023; Hui et al., 2024; Dubey et al., 2024; Touvron et al., 2023; Lozhkov et al., 2024) advancements that rely upon models with an increased number of parameters, trained upon more tokens, and being capable of handling longer contextual information.

In code representation, large language models must be adapted by training them on a curated corpus focused on software and by leveraging code’s syntactic and semantic structures, which differ significantly from natural language. Feng et al. (2020) adapted the BERT architecture to produce semantically meaningful embeddings for source code, resulting in codeBERT. This was accomplished by including more source code in the training set and focusing on a training loss that can leverage bimodal (natural language and code) contextual in-



(a) MRR



(b) Recall@1

Figure 4: Performance Comparison Across Layers: The graph illustrates the MRR and the Recall@1 for different layers, comparing baseline models and a self-distilled model.

formation (Clark et al., 2020). GraphCodeBERT enhanced codeBERT (Feng et al., 2020) representations by incorporating data flow graphs, capturing dependencies between variables and operations, and improving tasks like code summarization and clone detection. UniXcoder (Guo et al., 2022) extended this by introducing a unified encoder-decoder framework, integrating abstract syntax trees (ASTs) and data flow information. Wang et al. (2023) expanded these findings with codet5plus, stressing how multiple losses that leverage code semantics impact the model pertaining. The work incorporated text-code contrastive learning, text-code Matching, and text-code causal LM for better code understanding and generation.

When trying to achieve better performance, research has shifted toward models with a high number of parameters. While this trend appears effective from a performance perspective, end users may face computational or memory limitations as LLMs vary from millions to billions of parameters. Sanh et al. (2019) pioneered the introduction of knowledge distillation, using a “teacher” model that guides a smaller model to emulate its behavior. This methodology has been widely adopted and improved upon recently (DeepSeek-AI et al., 2025; Hui et al., 2024), becoming a standard for obtaining high-performing smaller LLMs.

Our work differs from previous work by adapting a modern architecture (Lozhkov et al., 2024) to a code encoder-only based model and introducing a novel ‘self-distillation’ mechanism. We replace the next sentence prediction loss with an in-context classification focused on the repository level and expand the context to 2048 tokens. Our novel self-distillation mechanism improves low-level lay-

ers, resulting in a modular transformer architecture without additional teacher models or further data for distillation.

5 Conclusion

In this work, we introduced MODULARSTARENCODER, a modular multi-exit encoder architecture designed to improve efficiency and scalability in code retrieval tasks. By integrating an intra-model self-distillation mechanism, our approach enables multiple resolution models to be trained within a unified layer stack, reducing redundancy while maintaining high retrieval performance. Our evaluation on CODEXGLUE demonstrates that MODULARSTARENCODER achieves state-of-the-art results among open-source models, outperforming prior baselines across text-to-code and code-to-code retrieval tasks. Ablations further highlighted the benefits of self-distillation, showing that lower layers gain representational strength from deeper layers, leading to superior performance compared to single-exit models.

Beyond performance gains, MODULARSTARENCODER offers practical benefits by providing multiple exit points, allowing users to balance computational efficiency and accuracy based on resource constraints. The results suggest that self-distillation provides a promising direction for efficient large-scale encoders, reducing deployment costs without sacrificing effectiveness.

Finally, we plan to release in open-access our SYNTHCODE2CODE2NL and both pre-trained and fine-tuned MODULARSTARENCODER models.

Limitations

Due to our dependence on multiple GPUs, we encountered significant computational constraints. Parameter grid searches with smaller and embryonic models were the only ways to extrapolate the best hyperparameter setup. The best hyperparameters for smaller models can differ from those for larger ones; thus, we faced a limitation in finding an optimal training setup. Ablating both the in-context classification and the multi-layer loss in a real scenario was impossible as we depended on smaller models to understand their performances. Therefore, computational resources pose a significant constraint in this work, and we want to emphasize how this factor undermines the possibility of replicating the experiments.

Here, we highlight potential threats to the validity of the research process, focusing on both external and internal factors.

External validity When synthesizing the SYNTHCODE2CODE2NL code, we rely on code translation; we understand that synthesized data adheres to stylistic writing patterns distinct from those of humans. We tested the model’s performance on standard benchmarks. However, the impact of utilizing code snippets as synthetic data in training large language models for generalization over human text-to-code and code-to-code search is still not fully understood.

Internal validity The ablation study focused on fine-tuning the model with and without multi-layer loss. However, this comparison does not account for how the model behaves when starting from a model not pre-trained on multi-layer loss. Although our experiments present promising results, further inspection is necessary to better understand this phenomenon.

References

Joshua Ainslie, James Lee-Thorp, Michiel de Jong, Yury Zemlyanskiy, Federico Lebrón, and Sumit Sanghai. 2023. *GQA: Training Generalized Multi-Query Transformer Models from Multi-Head Checkpoints*. *arXiv e-prints*, arXiv:2305.13245.

Loubna Ben Allal, Anton Lozhkov, Elie Bakouch, Gabriel Martín Blázquez, Guilherme Penedo, Lewis Tunstall, Andrés Marafioti, Hynek Kydlíček, Agustín Piqueres Lajarín, Vaibhav Srivastav, et al. 2025. Smollm2: When smol goes big—data-centric training of a small language model. *arXiv preprint arXiv:2502.02737*.

Stéphane Aroca-Ouellette and Frank Rudzicz. 2020. On losses for modern language models. *arXiv preprint arXiv:2010.01694*.

Xiao Bi, Deli Chen, Guanting Chen, Shanhuang Chen, Damai Dai, Chengqi Deng, Honghui Ding, Kai Dong, Qiusi Du, Zhe Fu, Huazuo Gao, Kaige Gao, Wenjun Gao, Ruiqi Ge, Kang Guan, Daya Guo, Jianzhong Guo, Guangbo Hao, Zhewen Hao, Ying He, Wenjie Hu, Panpan Huang, Erhang Li, Guowei Li, Jiashi Li, Yao Li, Y. K. Li, Wenfeng Liang, Fangyun Lin, Alex X. Liu, Bo Liu, Wen Liu, Xiaodong Liu, Xin Liu, Yiyuan Liu, Haoyu Lu, Shanghao Lu, Fuli Luo, Shirong Ma, Xiaotao Nie, Tian Pei, Yishi Piao, Junjie Qiu, Hui Qu, Tongzheng Ren, Zehui Ren, Chong Ruan, Zhangli Sha, Zhihong Shao, Junxiao Song, Xuecheng Su, Jingxiang Sun, Yaofeng Sun, Minghui Tang, Bingxuan Wang, Peiyi Wang, Shiyu Wang, Yaohui Wang, Yongji Wang, Tong Wu, Y. Wu, Xin Xie, Zhenda Xie, Ziwei Xie, Yiliang Xiong, Hanwei Xu, R. X. Xu, Yanhong Xu, Dejian Yang, Yuxiang You, Shuiping Yu, Xingkai Yu, B. Zhang, Haowei Zhang, Lecong Zhang, Liyue Zhang, Mingchuan Zhang, Minghua Zhang, Wentao Zhang, Yichao Zhang, Chenggang Zhao, Yao Zhao, Shangyan Zhou, Shunfeng Zhou, Qihao Zhu, and Yuheng Zou. 2024. *Deepseek LLM: scaling open-source language models with longtermism*. *CoRR*, abs/2401.02954.

Kevin Clark, Minh-Thang Luong, Quoc V. Le, and Christopher D. Manning. 2020. *ELECTRA: pre-training text encoders as discriminators rather than generators*. In *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26-30, 2020*. OpenReview.net.

Tri Dao. 2023. *FlashAttention-2: Faster Attention with Better Parallelism and Work Partitioning*. *arXiv e-prints*, arXiv:2307.08691.

DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojuan Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiusi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng

652	Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing	Elias Frantar, Artem Babenko, and Dan Alistarh.	714
653	Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun,	2024. Extreme Compression of Large Language	715
654	T. Wang, Wangding Zeng, Wanjia Zhao, Wen Liu,	Models via Additive Quantization . <i>arXiv e-prints</i> ,	716
655	Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao	arXiv:2401.06118.	717
656	Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan		
657	Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin	Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan,	718
658	Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li,	Xiaocheng Feng, Ming Gong, Linjun Shou, Bing	719
659	Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin,	Qin, Ting Liu, Daxin Jiang, and Ming Zhou.	720
660	Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxi-	2020. CodeBERT: A Pre-Trained Model for Pro-	721
661	ang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang,	gramming and Natural Languages . <i>arXiv e-prints</i> ,	722
662	Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang	arXiv:2002.08155.	723
663	Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng		
664	Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi,	Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xi-	724
665	Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang,	aocheng Feng, Ming Gong, Linjun Shou, Bing Qin,	725
666	Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo,	Ting Liu, Daxin Jiang, and Ming Zhou. 2020. Code-	726
667	Yuan Ou, Yudian Wang, Yue Gong, Yuheng Zou, Yu-	bert: A pre-trained model for programming and nat-	727
668	jia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You,	ural languages . In <i>Findings of the Association for</i>	728
669	Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu,	<i>Computational Linguistics: EMNLP 2020, Online</i>	729
670	Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu,	<i>Event, 16-20 November 2020</i> , volume EMNLP 2020	730
671	Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan,	of <i>Findings of ACL</i> , pages 1536–1547. Association	731
672	Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean	for Computational Linguistics.	732
673	Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao,		
674	Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zi-	Daya Guo, Shuai Lu, Nan Duan, Yanlin Wang, Ming	733
675	jia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song,	Zhou, and Jian Yin. 2022. Unixcoder: Unified cross-	734
676	Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu	modal pre-training for code representation . In <i>Pro-</i>	735
677	Zhang, and Zhen Zhang. 2025. Deepseek-r1: Incen-	<i>ceedings of the 60th Annual Meeting of the Associa-</i>	736
678	tivizing reasoning capability in llms via reinforce-	<i>tion for Computational Linguistics (Volume 1: Long</i>	737
679	ment learning . <i>Preprint</i> , arXiv:2501.12948.	<i>Papers)</i> , ACL 2022, Dublin, Ireland, May 22-27,	738
		2022, pages 7212–7225. Association for Computa-	739
		tional Linguistics.	740
680	Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey,		
681	Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman,	Song Han, Jeff Pool, John Tran, and William J. Dally.	741
682	Akhil Mathur, Alan Schelten, Amy Yang, Angela	2015. Learning both Weights and Connections	742
683	Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang,	for Efficient Neural Networks . <i>arXiv e-prints</i> ,	743
684	Archi Mitra, Archie Sravankumar, Artem Korenev,	arXiv:1506.02626.	744
685	Arthur Hinsvark, Arun Rao, Aston Zhang, Aurélien		
686	Rodriguez, Austen Gregerson, Ava Spataru, Bap-	Binyuan Hui, Jian Yang, Zeyu Cui, Jiayi Yang, Day-	745
687	tiste Rozière, Bethany Biron, Binh Tang, Bobbie	iheng Liu, Lei Zhang, Tianyu Liu, Jiajun Zhang,	746
688	Chern, Charlotte Caucheteux, Chaya Nayak, Chloe	Bowen Yu, Kai Dang, An Yang, Rui Men, Fei Huang,	747
689	Bi, Chris Marra, Chris McConnell, Christian Keller,	Xingzhang Ren, Xuancheng Ren, Jingren Zhou, and	748
690	Christophe Touret, Chunyang Wu, Corinne Wong,	Junyang Lin. 2024. Qwen2.5-coder technical report .	749
691	Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-	<i>CoRR</i> , abs/2409.12186.	750
692	lonsius, Daniel Song, Danielle Pintz, Danny Livshits,		
693	David Esiobu, Dhruv Choudhary, Dhruv Mahajan,	Hamel Husain, Ho-Hsiang Wu, Tiferet Gazit, Miltiadis	751
694	Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes,	Allamanis, and Marc Brockschmidt. 2019. Code-	752
695	Egor Lakomkin, Ehab AlBadawy, Elina Lobanova,	searchnet challenge: Evaluating the state of semantic	753
696	Emily Dinan, Eric Michael Smith, Filip Radenovic,	code search . <i>CoRR</i> , abs/1909.09436.	754
697	Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georg-		
698	ia Lewis Anderson, Graeme Nail, Grégoire Mialon,	Benoit Jacob, Skirmantas Kligys, Bo Chen, Meng-	755
699	Guan Pang, Guillem Cucurell, Hailey Nguyen, Han-	long Zhu, Matthew Tang, Andrew Howard, Hartwig	756
700	nah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov,	Adam, and Dmitry Kalenichenko. 2017. Quantiza-	757
701	Imanol Arrieta Ibarra, Isabel M. Kloumann, Ishan	tion and Training of Neural Networks for Efficient	758
702	Misra, Ivan Evtimov, Jade Copet, Jaewon Lee, Jan	Integer-Arithmetic-Only Inference . <i>arXiv e-prints</i> ,	759
703	Geffert, Jana Vranes, Jason Park, Jay Mahadeokar,	arXiv:1712.05877.	760
704	Jeet Shah, Jelier van der Linde, Jennifer Billock,		
705	Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi,	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	761
706	Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu,	sch, Chris Bamford, Devendra Singh Chaplot, Diego	762
707	Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	763
708	Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia,	laume Lample, Lucile Saulnier, Léo Renard Lavaud,	764
709	Kalyan Vasuden Alwala, Kartikeya Upasani, Kate	Marie-Anne Lachaux, Pierre Stock, Teven Le Scao,	765
710	Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, and	Thibaut Lavril, Thomas Wang, Timothée Lacroix,	766
711	et al. 2024. The llama 3 herd of models . <i>CoRR</i> ,	and William El Sayed. 2023. Mistral 7B . <i>arXiv</i>	767
712	abs/2407.21783.	<i>e-prints</i> , arXiv:2310.06825.	768
713	Vage Egiazarian, Andrei Panferov, Denis Kuznedelev,	Xiaoqi Jiao, Yichun Yin, Lifeng Shang, Xin Jiang, Xiao	769
		Chen, Linlin Li, Fang Wang, and Qun Liu. 2019.	770

771	TinyBERT: Distilling BERT for Natural Language Understanding . <i>arXiv e-prints</i> , arXiv:1909.10351.	830
772		831
773	Aditya Kusupati, Gantavya Bhatt, Aniket Rege, Matthew Wallingford, Aditya Sinha, Vivek Ramanujan, William Howard-Snyder, Kaifeng Chen, Sham Kakade, Prateek Jain, and Ali Farhadi. 2022. Ma-tryoshka representation learning . In <i>Advances in Neural Information Processing Systems</i> , volume 35, pages 30233–30249. Curran Associates, Inc.	832
774		833
775		834
776		835
777		836
778		837
779		838
780	Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. 2019. ALBERT: A Lite BERT for Self-supervised Learning of Language Representations . <i>arXiv e-prints</i> , arXiv:1909.11942.	839
781		840
782		841
783		842
784		843
785	Raymond Li, Loubna Ben Allal, Yangtian Zi, Niklas Muennighoff, Denis Kocetkov, Chenghao Mou, Marc Marone, Christopher Akiki, Jia Li, Jenny Chim, Qian Liu, Evgenii Zheltonozhskii, Terry Yue Zhuo, Thomas Wang, Olivier Dehaene, Mishig Davaadorj, Joel Lamy-Poirier, João Monteiro, Oleh Shliazhko, Nicolas Gontier, Nicholas Meade, Armel Zebaze, Ming-Ho Yee, Logesh Kumar Umapathi, Jian Zhu, Benjamin Lipkin, Muhtasham Oblokulov, Zhiruo Wang, Rudra Murthy, Jason Stillerman, Siva Sankalp Patel, Dmitry Abulkhanov, Marco Zocca, Manan Dey, Zhihan Zhang, Nour Fahmy, Urvashi Bhat-tacharyya, Wenhao Yu, Swayam Singh, Sasha Luc-cioni, Paulo Villegas, Maxim Kunakov, Fedor Zh-danov, Manuel Romero, Tony Lee, Nadav Timor, Jennifer Ding, Claire Schlesinger, Hailey Schoelkopf, Jan Ebert, Tri Dao, Mayank Mishra, Alex Gu, Jen-nifer Robinson, Carolyn Jane Anderson, Brendan Dolan-Gavitt, Danish Contractor, Siva Reddy, Daniel Fried, Dzmitry Bahdanau, Yacine Jernite, Carlos Muñoz Ferrandis, Sean Hughes, Thomas Wolf, Ar-jun Guha, Leandro von Werra, and Harm de Vries. 2023. StarCoder: may the source be with you! <i>arXiv e-prints</i> , arXiv:2305.06161.	844
786		845
787		846
788		847
789		848
790		849
791		850
792		851
793		852
794		853
795		854
796		855
797		856
798		857
799		858
800		859
801		860
802		861
803		862
804		863
805		864
806		865
807		866
808		867
809	Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang, Wei-Ming Chen, Wei-Chen Wang, Guangxuan Xiao, Xingyu Dang, Chuang Gan, and Song Han. 2023. AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration . <i>arXiv e-prints</i> , arXiv:2306.00978.	868
810		869
811		870
812		871
813		872
814		873
815	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-dar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A robustly optimized BERT pretraining approach . <i>CoRR</i> , abs/1907.11692.	874
816		875
817		876
818		877
819		878
820	Anton Lozhkov, Raymond Li, Loubna Ben Allal, Fed-erico Cassano, Joel Lamy-Poirier, Nouamane Tazi, Ao Tang, Dmytro Pykhtar, Jiawei Liu, Yuxiang Wei, Tianyang Liu, Max Tian, Denis Kocetkov, Arthur Zucker, Younes Belkada, Zijian Wang, Qian Liu, Dmitry Abulkhanov, Indraneil Paul, Zhuang Li, Wen-Ding Li, Megan Risdal, Jia Li, Jian Zhu, Terry Yue Zhuo, Evgenii Zheltonozhskii, Nii Osae Osae Dade, Wenhao Yu, Lucas Krauß, Naman Jain, Yixuan Su, Xuanli He, Manan Dey, Edoardo Abati, Yekun Chai,	879
821		880
822		881
823		882
824		883
825		884
826		885
827		886
828		887
829		888
		889
	Niklas Muennighoff, Xiangru Tang, Muhtasham Oblokulov, Christopher Akiki, Marc Marone, Cheng-hao Mou, Mayank Mishra, Alex Gu, Binyuan Hui, Tri Dao, Armel Zebaze, Olivier Dehaene, Nicolas Patry, Canwen Xu, Julian McAuley, Han Hu, Torsten Scholak, Sebastien Paquet, Jennifer Robinson, Car-olyn Jane Anderson, Nicolas Chapados, Mostofa Pat-wary, Nima Tajbakhsh, Yacine Jernite, Carlos Muñoz Ferrandis, Lingming Zhang, Sean Hughes, Thomas Wolf, Arjun Guha, Leandro von Werra, and Harm de Vries. 2024. StarCoder 2 and The Stack v2: The Next Generation . <i>arXiv e-prints</i> , arXiv:2402.19173.	889
		890
	Shuai Lu, Daya Guo, Shuo Ren, Junjie Huang, Alexey Svyatkovskiy, Ambrosio Blanco, Colin B. Clement, Dawn Drain, Daxin Jiang, Duyu Tang, Ge Li, Li-dong Zhou, Linjun Shou, Long Zhou, Michele Tu-fano, Ming Gong, Ming Zhou, Nan Duan, Neel Sun-daresan, Shao Kun Deng, Shengyu Fu, and Shujie Liu. 2021. Codexglue: A machine learning bench-mark dataset for code understanding and generation . In <i>Proceedings of the Neural Information Process-ing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual</i> .	891
		892
	Changan Niu, Chuanyi Li, Vincent Ng, Dongxiao Chen, Jidong Ge, and Bin Luo. 2023. An empirical compar-ison of pre-trained models of source code. In <i>2023 IEEE/ACM 45th International Conference on Soft-ware Engineering (ICSE)</i> , pages 2136–2148. IEEE.	893
		894
	Matthew E. Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. 2018. Deep contextualized word repre-sentations . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Tech-nologies, Volume 1 (Long Papers)</i> , pages 2227–2237, New Orleans, Louisiana. Association for Computa-tional Linguistics.	895
		896
	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sas-try, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. 2021. Learn-ing transferable visual models from natural language supervision . In <i>Proceedings of the 38th International Conference on Machine Learning, ICML 2021, 18-24 July 2021, Virtual Event</i> , volume 139 of <i>Proceedings of Machine Learning Research</i> , pages 8748–8763. PMLR.	897
		898
	Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. 2019. DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter . <i>arXiv e-prints</i> , arXiv:1910.01108.	899
		900
	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 em-beddings . In <i>Findings of the Association for Com-putational Linguistics: ACL 2023, Toronto, Canada, July 9-14, 2023</i> , pages 1102–1121. Association for Computational Linguistics.	901

890	Jianlin Su, Yu Lu, Shengfeng Pan, Ahmed Murtadha,	Dejiao Zhang, Wasi Ahmad, Ming Tan, Hantian Ding,	949
891	Bo Wen, and Yunfeng Liu. 2021. RoFormer: En-	Ramesh Nallapati, Dan Roth, Xiaofei Ma, and Bing	950
892	hanced Transformer with Rotary Position Embedding.	Xiang. 2024. Code Representation Learning At Scale.	951
893	<i>arXiv e-prints</i> , arXiv:2104.09864.	<i>arXiv e-prints</i> , arXiv:2402.01935.	952
894	Surat Teerapittayanon, Bradley McDanel, and H. T.	Chen Zhu, Wei Ping, Chaowei Xiao, Mohammad	953
895	Kung. 2017. BranchyNet: Fast Inference via Early	Shoeybi, Tom Goldstein, Anima Anandkumar, and	954
896	Exiting from Deep Neural Networks. <i>arXiv e-prints</i> ,	Bryan Catanzaro. 2021. Long-short transformer: Ef-	955
897	arXiv:1709.01686.	ficient transformers for language and vision. In <i>Ad-</i>	956
898	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	<i>advances in Neural Information Processing Systems 34:</i>	957
899	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	<i>Annual Conference on Neural Information Process-</i>	958
900	Bashlykov, Soumya Batra, Prajwal Bhargava, Shruti	<i>ing Systems 2021, NeurIPS 2021, December 6-14,</i>	959
901	Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton-	2021, virtual, pages 17723–17736.	960
902	Ferrer, Moya Chen, Guillem Cucurull, David Esiobu,		
903	Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller,	A Synthetic dataset	961
904	Cynthia Gao, Vedanuj Goswami, Naman Goyal, An-		
905	thony Hartshorn, Saghar Hosseini, Rui Hou, Hakan	SYNTHCODE2CODE2NL is a fine-tuning dataset	962
906	Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa,	designed for text-to-code and code-to-code search,	963
907	Isabel Kloumann, Artem Korenev, Punit Singh Koura,	built by augmenting CODESEARCHNET (Husain	964
908	Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Di-	et al., 2019) with transpiled code snippets across	965
909	ana Liskovich, Yinghai Lu, Yuning Mao, Xavier Mar-	multiple languages (Python, Java, Go, PHP, Ruby,	966
910	tinet, Todor Mihaylov, Pushkar Mishra, Igor Moly-	C++, C, JavaScript). The dataset underwent a pre-	967
911	bog, Yixin Nie, Andrew Poulton, Jeremy Reizen-	processing phase, including deduplication based on	968
912	stein, Rashi Rungta, Kalyan Saladi, Alan Schelten,	the original and synthesized code columns. Near-	969
913	Ruan Silva, Eric Michael Smith, Ranjan Subrama-	deduplication was performed using Locality Sensi-	970
914	nian, Xiaoqing Ellen Tan, Binh Tang, Ross Tay-	tive Hashing (LSH) with a Jaccard similarity thresh-	971
915	lor, Adina Williams, Jian Xiang Kuan, Puxin Xu,	old of 0.7 over character-level 5-grams to remove	972
916	Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan,	semantically identical snippets differing only in	973
917	Melanie Kambadur, Sharan Narang, Aurélien Ro-	identifiers or function arguments.	974
918	driguez, Robert Stojnic, Sergey Edunov, and Thomas		
919	Scialom. 2023. Llama 2: Open foundation and fine-	For code-to-code search, we translated each snip-	975
920	tuned chat models. <i>CoRR</i> , abs/2307.09288.	pet into a randomly sampled target language us-	976
921	Matteo Turisini, Giorgio Amati, and Mirko Cestari.	using the QWEN2.5-CODER-7B-INSTRUCT model	977
922	2023. LEONARDO: A pan-european pre-exascale	with greedy search to ensure consistency. Each	978
923	supercomputer for HPC and AI applications. <i>CoRR</i> ,	dataset entry consists of a natural language descrip-	979
924	abs/2307.16885.	tion and two code snippets in different languages.	980
925	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob	SYNTHCODE2CODE2NL contains 1,071,367 sam-	981
926	Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz	ples, with original code from CODESEARCHNET	982
927	Kaiser, and Illia Polosukhin. 2017. Attention Is All	(Python, Java, PHP, Go) and translated code (Go,	983
928	You Need. <i>arXiv e-prints</i> , arXiv:1706.03762.	Ruby, JavaScript, Python, C++, PHP, C, Java). In	984
929	Yue Wang, Hung Le, Akhilesh Gotmare, Nghi D. Q. Bui,	Figure 5, In Figure 6 and Figure 7 some examples	985
930	Junnan Li, and Steven C. H. Hoi. 2023. Codet5+:	of code translation are shown.	986
931	Open code large language models for code under-		
932	standing and generation. In <i>Proceedings of the 2023</i>		
933	<i>Conference on Empirical Methods in Natural Lan-</i>		
934	<i>guage Processing, EMNLP 2023, Singapore, Decem-</i>		
935	<i>ber 6-10, 2023</i> , pages 1069–1088. Association for		
936	Computational Linguistics.		
937	Benjamin Warner, Antoine Chaffin, Benjamin Clavié,		
938	Orion Weller, Oskar Hallström, Said Taghadouini,		
939	Alexis Gallagher, Raja Biswas, Faisal Ladhak, Tom		
940	Aarsen, Nathan Cooper, Griffin Adams, Jeremy		
941	Howard, and Iacopo Poli. 2024. Smarter, better,		
942	faster, longer: A modern bidirectional encoder for		
943	fast, memory efficient, and long context finetuning		
944	and inference. <i>CoRR</i> , abs/2412.13663.		
945	Jinle Zeng, Min Li, Zhihua Wu, Jiaqi Liu, Yuang Liu,		
946	Dianhai Yu, and Yanjun Ma. 2022. Boosting dis-		
947	tributed training performance of the unpadded BERT		
948	model. <i>CoRR</i> , abs/2208.08124.		

GO

```
// readNextChunk creates an io.LimitedReader for the next chunk of data,
// and returns io.EOF if the last chunk has been read.
func readNextChunk(r wire.SyncScanner) (io.Reader, error) {
    status, err := r.ReadStatus("read-chunk")
    if err != nil {
        if wire.IsAdbServerErrorMatching(err, readFileNotFoundPredicate) {
            return nil, errors.Errorf(errors.FileNotExistError, "no such file or directory")
        }
        return nil, err
    }

    switch status {
    case wire.StatusSyncData:
        return r.ReadBytes()
    case wire.StatusSyncDone:
        return nil, io.EOF
    default:
        return nil, errors.Errorf(
            errors.AssertionError,
            "expected chunk id '%s' or '%s', but got '%s'",
            wire.StatusSyncData, wire.StatusSyncDone, []byte(status),
        )
    }
}
```

Ruby

```
def read_next_chunk(r)
    status, err = r.read_status("read-chunk")
    if err
        if AdbServerErrorMatching(err, readFileNotFoundPredicate)
            return nil, FileNotExistError.new("no such file or directory")
        end
        return nil, err
    end

    case status
    when wire.StatusSyncData
        return r.read_bytes
    when wire.StatusSyncDone
        return nil, EOFError.new
    else
        return nil, AssertionError.new(
            "expected chunk id '#{wire.StatusSyncData}' or '#{wire.StatusSyncDone}', but got '#{status}'"
        )
    end
end
```

Figure 5: Example of a code snippet translation from *Go* to *Ruby*.

JAVA

```
public void get(URI uri, Credentials credentials) throws IOException {  
    get(uri, null, credentials);  
}
```

C

```
#include <stdio.h>  
#include <stdlib.h>  
#include <string.h>  
  
typedef struct {  
    char *username;  
    char *password;  
} Credentials;  
  
typedef struct {  
    char *scheme;  
    char *host;  
    int port;  
    char *path;  
} URI;  
  
void get(URI *uri, Credentials *credentials) {  
    // Implementation of the get function  
    printf("Fetching resource from %s://%s:%d%s\n", uri->scheme, uri->host, uri->port, uri->path);  
    if (credentials != NULL) {  
        printf("Using credentials: %s:%s\n", credentials->username, credentials->password);  
    }  
}  
  
int main() {  
    URI uri = {"http", "example.com", 80, "/index.html"};  
    Credentials credentials = {"user", "pass"};  
    get(&uri, &credentials);  
    return 0;  
}
```

Figure 6: Example of a code snippet translation from *Java* to *C*.

Python

```
def toString(self):
    result = []
    k, v = self.optimalRepr()
    longest = reduce(lambda x, y: x if x > len(y) else len(y), k, 0)
    for ind in range(len(k)):
        result.append("%s : %s" % (k[ind].ljust(longest), v[ind]))
    return "\n".join(result)
```

PHP

```
public function toString() {

    /**
     * Return a printable view of the dictionary
     */

    $result = [];

    list($k, $v) = $this->optimalRepr();

    $longest = array_reduce($k, function($x, $y) {
        return $x > strlen($y) ? $x : strlen($y);
    }, 0);

    for ($ind = 0; $ind < count($k); $ind++) {
        $result[] = sprintf("%s : %s", ltrim($k[$ind], ' '), str_pad($v[$ind], $longest, ' ', STR_PAD_LEFT));
    }

    return implode("\n", $result);
}
```

Figure 7: Example of a code snippet translation from *Python* to *PHP*.