

FR-LoRA: Fisher Regularized LoRA for Multilingual Continual Learning

Sayanta Adhikari
sayantaa@amazon.com
Amazon.com Inc.
Bengaluru, India

Sanjay Agrawal
sanjagr@amazon.com
Amazon.com Inc.
Bengaluru, India

Vivek Sembium
viveksem@amazon.com
Amazon.com Inc.
Bengaluru, India

Abstract

Relevance in e-commerce product search is critical to ensuring that results accurately reflect customer intent. While large language models (LLMs) have recently advanced natural language processing capabilities, their high inference latency and significant infrastructure demands make them less suitable for real-time e-commerce applications. Consequently, transformer-based encoder models are widely adopted for relevance classification tasks. These models typically evaluate the relevance of a product to a given query by encoding the query and product title as input features. As e-commerce stores expand into new marketplaces, the need for language- and region-specific relevance models grows, often resulting in the sequential development and maintenance of separate models per marketplace. To address this challenge, we introduce a multilingual continual learning (CL) framework that mitigates catastrophic forgetting. Our proposed method, **FR-LoRA** (Fisher Regularized LoRA), integrates Elastic Weight Consolidation (EWC) with marketplace-specific LoRA modules, where each LoRA is regularized using the Fisher information matrix. FR-LoRA retains the same inference-time footprint as the base model, ensuring zero additional latency while enabling frequent, scalable updates. Empirically, our approach achieves a **~3% ROC-AUC improvement** over single-marketplace baselines and outperforms several recent CL baselines on both proprietary and public datasets.

CCS Concepts

• **Information systems** → **Web searching and information discovery; Query reformulation;** • **Computing methodologies** → **Natural language processing.**

Keywords

Continual Learning, Relevance Classification, LoRA, Fisher Information, Product Search

ACM Reference Format:

Sayanta Adhikari, Sanjay Agrawal, and Vivek Sembium. 2025. FR-LoRA: Fisher Regularized LoRA for Multilingual Continual Learning. In *Proceedings of the 34th ACM International Conference on Information and Knowledge*

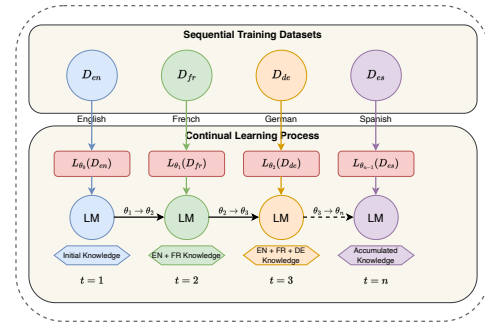


Figure 1: In Continual Learning, the language model (LM) is trained on datasets one at a time—starting with English, followed by French, and so on. The model’s parameters are updated sequentially based on the loss function $L(\cdot)$.

Management (CIKM '25), November 10–14, 2025, Seoul, Republic of Korea. ACM, New York, NY, USA, 8 pages. <https://doi.org/10.1145/3746252.3761531>

1 Introduction

Relevance modeling is a foundational component in e-commerce product search, where the goal is to surface items that align closely with a customer’s intent [2, 16, 28]. With vast product catalogs and highly varied user queries, accurately classifying query-product relevance is critical to delivering satisfying search experiences [4, 20]. While recent advancements in Natural Language Processing (NLP)—particularly through large language models (LLMs)—have demonstrated impressive gains in language understanding, their high inference latency and significant computational costs render them impractical for production-scale e-commerce systems [3, 26]. Consequently, lightweight transformer-based encoders have become the preferred solution for relevance classification tasks due to their efficiency and performance [5, 9].

In a typical deployment, these models ingest a customer query and a product title and produce a relevance score, helping determine which products should appear in search results [3]. However, as global marketplaces grow and diversify, a new challenge emerges: the need for language-specific relevance models. Each new marketplace introduces unique linguistic and cultural nuances, which often necessitate training and maintaining dedicated models for each region [21, 32].

One intuitive solution is to train a single multilingual model on data from all marketplaces to facilitate semantic alignment across languages and improve knowledge transfer. However, this approach is computationally intensive and requires access to all marketplace

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

CIKM '25, Seoul, Republic of Korea.

© 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 979-8-4007-2040-6/2025/11

<https://doi.org/10.1145/3746252.3761531>

data simultaneously. Such requirements hinder scalability, especially when continually expanding into new regions. To address these limitations, we propose a continual learning (CL) framework for multilingual relevance classification [10]. In this setting, marketplace datasets—each associated with a unique language—are introduced sequentially. The model is trained on new language data without accessing previous datasets, making it crucial to preserve performance on earlier tasks while integrating knowledge from new ones. A central challenge in this paradigm is avoiding catastrophic forgetting, where performance on earlier marketplaces deteriorates as new ones are added.

In this work, we introduce **FR-LoRA**, a novel continual learning strategy designed to enable scalable multilingual relevance modeling. FR-LoRA operates in two stages during each training step: first, it updates the base model using Elastic Weight Consolidation (EWC) to preserve shared knowledge across tasks. Then, it trains independent adaptation modules (LoRA or DoRA) to capture marketplace-specific nuances and enhance task-specific performance. To minimize interference between the base model and adaptation modules, the LoRA/DoRA parameters are regularized using the Fisher information matrix computed during the EWC stage. This helps mitigate catastrophic forgetting across sequential tasks. Importantly, FR-LoRA is designed such that only one LoRA or DoRA module is active during inference, keeping the total parameter count identical to the base model and ensuring zero additional inference-time latency. The **key contributions** of our approach are:

1. Develop a novel multilingual continual learning method that enhances relevance classification across multiple marketplaces. This approach takes advantage of cross-lingual transfer to boost the performance on each task while being trained in sequential fashion.
2. Rigorous evaluation on a sequence of 7 multilingual amazon marketplace data, publicly available aicrowd ESCI data (with 3 marketplaces). Our method achieve an average ROC-AUC improvement of 3.02% in amazon data, 0.7% improvement on aicrowd data over vanilla finetuning of base model on each marketplace separately. Our method achieves these boost in performance with zero increase in inference time cost.
3. Conducted several ablations to check the robustness of our approach as well as compare forgetting of our approach against other state-of-the-art CL methods. Our ablations also reveals the importance of our proposed fisher regularization on LoRA/DoRA.

Note that our approach is not limited to query-product relevance prediction and can be easily adapted to support other tasks and evaluation metrics. For example, it can be used for multilingual ad categorization, where new languages or product categories are added over time. These updates can be made without retraining the model from scratch. In addition, our continual learning framework works with any transformer-based model and does not depend on the order in which languages are introduced. This makes it flexible and easy to use for researchers and practitioners who want to try different models and language training setups.

2 Related Work

Continual learning methods are commonly categorized into four main approaches: **Regularization-based methods**: These approaches apply regularization during training on new tasks to preserve knowledge learned from previous ones [17, 23, 34]. For example, methods like Elastic Weight Consolidation (EWC) [17] constrain updates to parameters deemed important for past marketplaces.

Replay-based methods [11, 29, 30]: These methods maintain a buffer of samples from previous tasks. During training on new marketplaces, examples from this buffer are replayed alongside current data, providing implicit regularization and helping to retain prior knowledge. **Distillation-based methods** [6, 13]: In this teacher-student framework, the model trained on previous tasks serves as a teacher. As the student model learns from the current task, it also distills knowledge from the teacher, preserving performance on earlier marketplaces [7, 8]. **Architecture-based methods** [27, 33]: These methods introduce additional parameters to capture task-specific knowledge. In the context of transformer-based language models, such approaches are often implemented through parameter-efficient fine-tuning (PEFT) techniques such as soft prompts [18], prefix-tuning [19], or Low-Rank Adaptation (LoRA) [15].

With the increasing use of large transformer-based language models, it is essential to adapt them to downstream tasks while preserving the general knowledge acquired during pretraining. PEFT methods are widely adopted for this purpose [14]. Among them, **LoRA** [15] has gained significant popularity due to its efficiency. LoRA models the parameter updates as a low-rank matrix $\Delta W = BA$, where B and A are low-dimensional matrices. This formulation enables the model to capture task-specific nuances while significantly reducing the number of trainable parameters compared to full-model fine-tuning.

Several variants of LoRA have been proposed to address its limitations [24]. One notable recent extension is **DoRA** [22], which analyzes the differences in magnitude and direction of parameter changes between full fine-tuning and LoRA. Based on this insight, DoRA decouples the learning of directional and magnitude components, resulting in improved performance. As a generalization of LoRA, DoRA can be applied in any context where LoRA is used.

3 Problem Statement

We define the multilingual continual learning problem [10] in the context of marketplace relevance modeling as follows. Consider a sequence of n distinct marketplaces, each associated with a dataset $\{D_i\}_{i=1}^n$, where each D_i corresponds to a unique combination of language and domain. Our goal is to train a multilingual transformer model sequentially on these datasets, with the constraint that data from previous marketplaces is no longer accessible during future training stages, ensuring computational efficiency and data privacy.

The objective is to learn a model that, at any point in the sequence, performs well across all encountered marketplaces (both past and current), without suffering from *catastrophic forgetting*. Furthermore, the model should leverage *forward transfer*, benefiting from shared patterns across languages and domains, to enhance learning on future marketplaces.

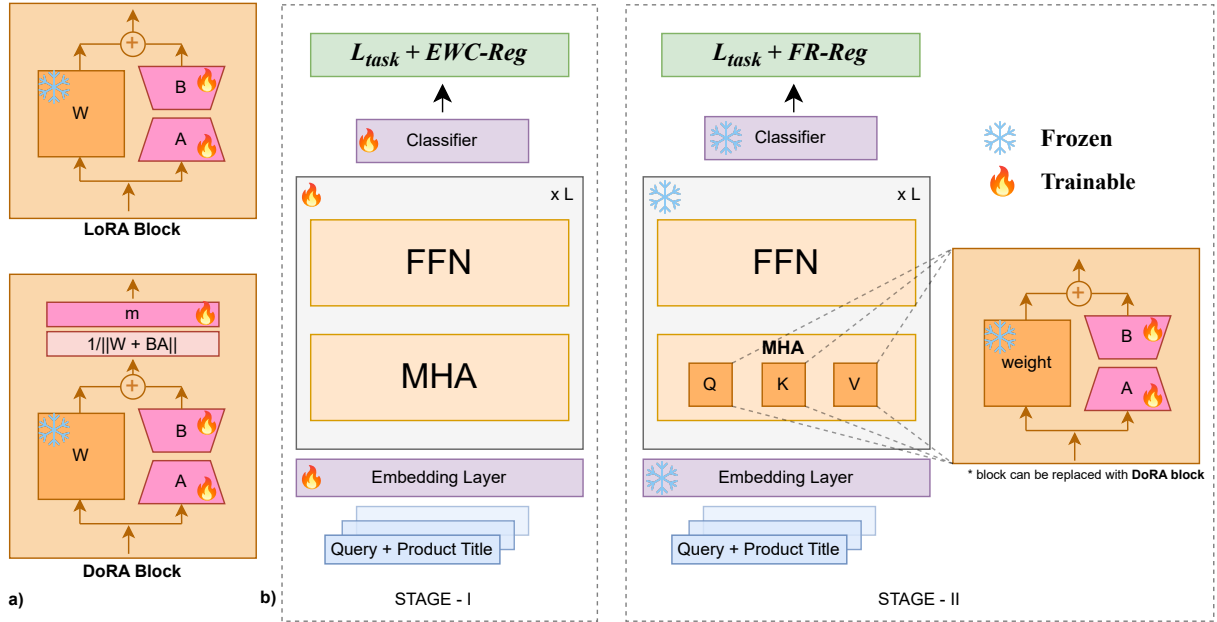


Figure 2: (a) Shows how two different kinds of LoRA variants (Standard LoRA as LoRA-Block and DoRA as DoRA Block) is applied to any weight matrices to adapt them for a downstream task. (b) demonstrate our method FR-LoRA for each marketplace, which is trained in two-stage fashion, where Stage-I loss ($L_{task} + EWC-reg$) as provided in Equation 2 and Stage-II loss ($L_{task} + FR-reg$) as provided in Equation 4 are used. Here, MHA implies Multi-head Attention and FFN implies Feed-Forward Network that is added after MHA in transformer blocks.

4 Proposed Methodology

This paper proposes a query-product relevance classification model designed for sequential multi-marketplace adaptation. Our approach aims to surpass marketplace-specific training performance while eliminating the need for historical data storage and retraining. The model facilitates efficient cross-marketplace knowledge transfer, reducing both computational and storage overhead.

A core challenge in this setting is balancing the retention of prior marketplace knowledge with the adaptation to new, marketplace-specific signals—especially across diverse languages and domains. To address this, we introduce a two-stage training strategy that leverages the strengths of Elastic Weight Consolidation (EWC) [17] and Low-Rank Adaptation (LoRA) [15]. This approach enables the model to preserve shared knowledge across tasks while introducing task-specific flexibility, all without requiring access to historical data.

This section is organized as follows: Section 4.1 discusses the use of LoRA and its recent variant, DoRA, in a continual learning environment. Section 4.2 presents the EWC method, which serves as the foundation of our approach. Section 4.3 introduces our proposed method, FR-LoRA which integrates both LoRA and EWC to mitigate catastrophic forgetting and promote forward transfer.

4.1 LoRA & DoRA in Continual Learning

We use a frozen multilingual transformer backbone θ_b (e.g., multilingual-MiniLM) for all tasks. For each new marketplace $\mathcal{T}_t$, we introduce lightweight, task-specific LoRA [15] modules $\theta_{LoRA,t}$, inserted into

the query, key, and value projection layers of every multi-head attention block, along with a new classification head $\theta_{cls,t}$. During training on dataset D_t , only $\theta_t = \{\theta_{LoRA,t}, \theta_{cls,t}\}$ is optimized, keeping θ_b frozen:

$$\theta_t \leftarrow \arg \min_{\theta_t} \mathcal{L}_t(D_t; \theta_b, \theta_t) \quad (1)$$

After training, only θ_t is stored for each task. Across M marketplaces, the total stored parameters are: $\Theta = \theta_b + \sum_{t=1}^M \theta_t$. LoRA approximates weight updates using low-rank matrices and assumes the pretrained weight magnitudes are optimal, mainly modifying direction. However, this limits expressiveness when magnitude updates are important. To address this, **DoRA** [22] decouples magnitude and direction by reparameterizing a weight matrix W as: $W = \|W\| \cdot \frac{W}{\|W\|}$. Here, direction is learned using low-rank matrices (like LoRA), while the magnitude is modeled using a learnable scalar. This separation improves adaptation with minimal overhead. Figure 2 (a) illustrates the structural differences for LoRA and DoRA in our continual learning framework.

4.2 Elastic Weight Consolidation

EWC reduces catastrophic forgetting by preserving parameters important to previous tasks. When learning task $\mathcal{T}_t$, it adds a regularization term to the current loss, penalizing deviation from the

previous optimum θ_{t-1}^* . The regularized loss is:

$$\mathcal{L}_{\text{EWC}}(\theta_t) = \mathcal{L}_t(D_t; \theta_t) + \mathbf{1}_{\{t>1\}} \left(\frac{\lambda}{2} \sum_i F_i^{(t-1)} (\theta_{t,i} - \theta_{t-1,i}^*)^2 \right) \quad (2)$$

where λ controls regularization strength and $F_i^{(t-1)}$ is the Fisher Information value for parameter θ_i from task $\mathcal{T}_{t-1}$. The Fisher values are estimated as:

$$F_i^{(t-1)} = \mathbb{E}_{(x,y) \sim D_{t-1}} \left[\left(\frac{\partial \log p(y|x; \theta_{t-1}^*)}{\partial \theta_i} \right)^2 \right] \quad (3)$$

This approach helps retain critical knowledge while allowing flexibility for new tasks.

4.3 Proposed Method: FR-LoRA

Overview: In this method we plan to leverage the benefits of EWC and LoRA both, where our base model will get updated based using EWC, capturing the general trends across multiple marketplaces and LoRA which is finetuned on each of the marketplaces independently to capture marketplace specific nuances. In our case then base model that gets updated helps in providing knowledge transfer for the future tasks, where as LoRAs tries to increase the plasticity of the model, improving the performance on the current marketplace. To achieve this we have designed a two-stage training process where first we train the base model using EWC, and then add LoRA to that model and optimize it to capture more specific patters for a particular marketplace. Figure 2(b) demonstrates the two stage setting of our approach.

Stage I (Base Model Training): We finetune the base model using the task-specific objective and the EWC regularization as described in Equation 2. Once this stage is complete, we compute the Fisher Information Matrix F_t for the current task t , which captures the importance of each parameter in the base model with respect to the learned task.

Stage II (LoRA Training): We add LoRA modules to the base model and optimize the LoRA parameters to better capture marketplace-specific information. However, our method also requires the base model θ_b to continue updating for future tasks. This means that if LoRA parameters are trained independently of this changing base, they may become incompatible as the base parameters shift. For example, suppose the base model is at state θ_t during task t , and LoRA parameters $\theta_{\text{LoRA},t}$ are trained in conjunction with it. Later, during task $t+1$, θ_t is updated to θ_{t+1} to fit the new data. Now, $\theta_{\text{LoRA},t}$ may no longer align well with the modified base model θ_{t+1} , leading to performance degradation on earlier tasks.

Fisher-Regularization on LoRA Parameters: To address this issue, we add a regularization term when training the LoRA parameters. This term ensures that LoRA modules focus on modifying only those parts of the base model that were important for the current task. We achieve this by weighting the LoRA parameters with the *inverse* of the Fisher Information Matrix values. This way, parameters with higher importance (i.e., higher Fisher values) receive lower regularization, allowing LoRA to more freely adjust them, while parameters with lower importance are restricted from being changed unnecessarily. The overall loss function for optimizing

LoRA parameters $\theta_{\text{LoRA},t}$ for task t becomes:

$$\begin{aligned} \theta_{\text{LoRA},t} \leftarrow \arg \min_{\theta_{\text{LoRA},t}} \mathcal{L}_t(D_t; \theta_b, \theta_{\text{LoRA},t}) + \\ \frac{\beta}{2} \sum_{i,j} \frac{1}{(F_t)_{i,j} + \epsilon} (\Delta(W_t)_{i,j})^2 \\ \Delta(W_t)_{i,j} = (B_t A_t)_{i,j} \quad \text{where } \theta_{\text{LoRA},t} = \{A_t, B_t\} \end{aligned} \quad (4)$$

Here, $\mathcal{L}_t$ is the task-specific loss, $(F_t)_{i,j}$ is the Fisher Information value for the i, j -th base model parameter, and ϵ is a small constant to prevent division by zero. Note that as the fisher information matrix used over here is associated with weight matrix on which LoRA is applied, which in our case is query, key and value weight matrices. The hyperparameter β controls the strength of the LoRA regularization. With certain modifications, we can apply the same fisher regularization on DoRA parameters as well.

This approach ensures that even when the base model is updated for future tasks, each task's LoRA parameters remain compatible by aligning them with the parts of the base model that were most relevant for that task. As a result, we maintain past performance while still allowing for adaptation to new marketplaces. A detailed algorithm for our approach is presented in the Algorithm 1.

Algorithm 1 FR-LoRA

Input: Base model θ_b , task sequence $\{\mathcal{T}_1, \dots, \mathcal{T}_M\}$, EWC regularization strength λ , FR regularization strength β , lr_decay_rate

- 1: Initialize Fisher matrix $F_t \leftarrow 0$
- 2: Initialize base model with pretrained weights.
- 3: **for** each task $\mathcal{T}_t$ with dataset D_t **do**
- 4: Fine-tune the base model (θ_b) using Equation 2
- 5: Calculate Fisher Information Matrix (Equation 3)
- 6: Freeze the base model
- 7: Add LoRA to the Query, Key and Value weight matrices.
- 8: Optimize LoRA parameters $\theta_{\text{LoRA},t}$ using Equation 4.
- 9: Save trained LoRA parameters $\theta_{\text{LoRA},t}^* \leftarrow \theta_{\text{LoRA},t}$
- 10: Update Fisher Information Matrix (Equation 3) by calculating it for the model with LoRA.
- 11: $\text{lr} \leftarrow \text{lr} \times \text{lr_decay_rate}$
- 12: **end for**
- 13: Save the final base model $\theta_b^* \leftarrow \theta_b$
- 14: **for** $t \in \{1, 2, \dots, M-1\}$ **do** ▶ Performing projection of LoRA parameters
- 15: Project $\theta_{\text{LoRA},t}^*$ to the space of $\theta_{\text{LoRA},M}^*$ using Equation 5
- 16: **end for**

4.4 Practical Modifications

To further improve the stability and efficiency of the proposed approach, we incorporate two additional techniques:

- **Learning-rate Decay:** To minimize excessive drift in the base model across tasks, we used a learning rate decay mechanism. Specifically, after training on each marketplace, we decay the learning rate used for base model updates.
- **LoRA Projection at Inference Time:** We introduce a projection mechanism for LoRA parameters at inference time. Since LoRA modules are trained on top of a specific version

Table 1: Performance comparison across multiple marketplaces. Multilingual training serves as an upper bound, while continual learning methods aim to balance knowledge retention and new task adaptation. Our proposed method, FR-LoRA (Fisher Regularized LoRA), achieves the best trade-off between plasticity and stability, outperforming strong baselines. <value> represents the top performance, and <value> represents second top performance. Mean & std. ($\pm$) for ROC-AUCs are reported based on 5 trials. The reported performance figures are expressed relative to independent fine-tuning, which serves as the baseline and is normalized to 1 $\times$.

Methods	Marketplaces							AVG
	M_a	M_b	M_c	M_d	M_e	M_f	M_g	
Independent Training ER (Buffer = 10000)	1 $\times$ -0.0445	1 $\times$ -0.0279	1 $\times$ <u>+0.0745</u>	1 $\times$ <u>+0.0063</u>	1 $\times$ +0.0029	1 $\times$ <u>+0.0397</u>	1 $\times$ <u>+0.0756</u>	1 $\times$ +0.0181 $\pm$ 0.0.0045
Parameter Isolation								
Prefix FT	+0.0405	-0.0147	+0.0622	-0.0002	+0.0045	+0.0212	+0.0294	+0.0204 $\pm$ 0.0.0005
LoRA FT	<u>+0.0548</u>	<u>+0.0115</u>	+0.0668	+0.0042	+0.0006	+0.0237	+0.0295	+0.0273 $\pm$ 0.0.0007
DoRA FT	<u>+0.0475</u>	<u>+0.0208</u>	+0.0694	+0.0025	+0.0037	+0.0201	+0.0385	+0.0289 $\pm$ 0.0.0008
LoRA Inference-Time Merging Strategy								
linear-combining [12]	-0.1704	-0.2231	-0.1749	-0.2648	-0.2147	-0.1483	-0.2303	-0.2038 $\pm$ 0.0.0015
lora_delta_w_svd [1]	+0.0204	-0.0312	+0.0258	-0.0368	-0.0473	0.0005	-0.0184	-0.0124 $\pm$ 0.0.0024
Continual Full Model Merging Strategy								
MagMax + Independent FT	-0.2057	-0.2658	-0.2237	-0.2756	-0.2283	-0.1745	-0.2674	-0.2344 $\pm$ 0.0.0023
MagMax + Sequential FT	+0.0112	-0.0481	+0.0165	-0.0948	-0.0927	-0.0285	-0.0674	-0.0434 $\pm$ 0.0.0012
OPCM + Independent FT	-0.0115	-0.0544	+0.0013	-0.0918	-0.0973	-0.0288	-0.0708	-0.0505 $\pm$ 0.0.0007
OPCM + Sequential FT	-0.0093	-0.0633	-0.0014	-0.1189	-0.1027	-0.0501	-0.1057	-0.0645 $\pm$ 0.0.0004
Recent SOTA Methods for CL								
LoRA Sequential Stacking	-0.0416	-0.1037	-0.0127	-0.1232	-0.1291	-0.0568	-0.0413	-0.0726 $\pm$ 0.0.0010
O-LoRA	-0.2785	-0.1561	-0.0841	-0.2223	-0.1923	-0.1180	-0.1472	-0.1712 $\pm$ 0.0.0009
Progressive-Prefix	+0.0193	-0.0061	+0.0371	-0.0317	-0.0280	+0.0019	-0.0146	-0.0031 $\pm$ 0.0.0011
OURS + Variants								
FR-LoRA	+0.0242	-0.0092	+0.0689	+0.0050	+0.0057	+0.0266	0.0715	+0.0275 $\pm$ 0.0.0013
FR-LoRA + Proj.	+0.0270	-0.0098	+0.0697	+0.0056	+0.0037	+0.0283	+0.0715	+0.0280 $\pm$ 0.0.0009
FR-LoRA + Lr-Decay	+0.0311	-0.0077	+0.0722	<u>+0.0062</u>	+0.0078	+0.0287	+0.0716	+0.0300 $\pm$ 0.0005
FR-LoRA+Lr-Decay+Proj.	+0.0327	-0.0079	+0.0725	+0.0064	+0.0064	<u>+0.0301</u>	<u>+0.0716</u>	<u>+0.0303</u> $\pm$ 0.0.0006
FR-DoRA	+0.0288	-0.0089	<u>+0.0731</u>	+0.0061	<u>+0.0070</u>	+0.0322	+0.0728	<u>+0.0302</u> $\pm$ 0.0.0010
FR-DoRA + Lr-Decay	+0.0277	-0.0090	+0.0719	<u>+0.0063</u>	<u>+0.0071</u>	+0.0301	+0.0716	+0.0294 $\pm$ 0.0.0007

of the base model, they may become less effective if the base model parameters have significantly changed due to continual updates. Given a matrix $A \in \mathbb{R}^{n \times d}$ and a target matrix $B \in \mathbb{R}^{n \times m}$, the function computes the orthogonal projection of B onto the column space of A using the Moore-Penrose pseudo-inverse. The projection $P_B \in \mathbb{R}^{n \times m}$ is given by:

$$P_B = A(A^\top A)^{-1}A^\top B \quad (5)$$

Here, $A^\top A \in \mathbb{R}^{d \times d}$ is invertible. The matrix $A(A^\top A)^{-1}A^\top$ is the projection matrix onto the column space of A , and P_B is the closest point to B in that subspace, minimizing the squared Euclidean distance. In our case we want to project LoRA from previous task to the LoRA associated to the last task. To perform this we use the above equation for projection for each of the two weight matrices associated with

LoRA. Similar projection can be performed on the weights associated with DoRA as well.

5 Experiments and Results

5.1 Experimental Setup

In this section we provide details regarding our experimental setup and choices of hyperparameter. We use *paraphrase multilingual MiniLM*¹ as our base encoder model.

All baseline models are trained using the *AdamW* optimizer with a learning rate of $5e-5$ for the methods where full model is trained and $3e-4$ for PEFT based methods, batch size of 128, and for 25 epochs per marketplace with an early stopping with patience 3. For

¹the base model weights can be found here: [hugging-face-multilingual-MiniLM](https://huggingface.co/multilingual-MiniLM)

baseline methods specific hyperparameters we considered are the ones that are reported in their paper.

For our method, the optimizer used for both stages are *AdamW* with a learning rate of $5e-5$ (cosine lr-scheduler) for Stage-I training and a learning rate of $1e-5$ (constant lr-scheduler with 1000 step warmup). For Stage-I, hyperparameter λ corresponding to EWC regularization was considered as 1. For Stage-II, hyperparameters of LoRA (same for DoRA) were rank $r = 64$ and $\alpha = 16$ scaling was utilized. The hyperparameter associated with Fisher Regularization on LoRA, β was considered as 0.0001. In both stages the model was trained for 25 epochs with an early stopping (patience = 3).

We report performance using ROC-AUC on all previously seen marketplaces after each task. All experiments are conducted on 8 NVIDIA V100 16GB GPUs.

5.2 Datasets

We benchmarked our approach on two e-commerce dataset, **Amazon proprietary e-commerce data** from seven marketplaces. For query-passage relevance, we use datasets from 7 e-commerce regions across with different languages, which are anonymized, and present as $M_a, M_b, M_c, M_d, M_e, M_f$, and M_g . Each of these marketplaces includes query-product title pairs and their associated ground truth label categorized as either relevant or non-relevant. We also use the publicly available ESCI (E-commerce Search Corpus with Implicit user feedback) dataset from Amazon, which contains search sessions sampled from the Amazon Search Query Logs. The ESCI dataset is used to evaluate the performance across three marketplaces - United States (US, English), Japan (JP, Japanese), and Spain (ES, Spanish) marketplaces.

5.3 Baselines

To evaluate our approach, we compare it against a diverse set of baselines covering both traditional and recent continual learning methods.

1. **Independent Finetuning**: The model is trained individually on each marketplace, with no transfer of information across tasks. This acts as a lower bound that highlights task-specific performance without any continual learning.
2. **Experience Replay** [29]: A subset of data from previous marketplaces is stored in a memory buffer and replayed during training on new marketplaces. This provides implicit regularization and helps reduce forgetting.
3. **PEFT Fine-tuning**: Parameter-efficient fine-tuning (PEFT) methods such as Prefix Tuning [19], LoRA [15], and DoRA [22] are applied independently to each marketplace without any shared parameters or continual learning objective.
4. **MagMax** [25]: A parameter-merging method that aggregates models trained on different tasks by selecting parameter values with maximum magnitude across tasks.
5. **OPCM** [31]: A projection-based merging method where parameters for a new marketplace are projected into the parameter space of the previously learned marketplace to minimize interference.
6. **O-LoRA** [33]: Marketplace-specific LoRA modules are fine-tuned for each task while keeping previously trained LoRA

Table 2: Performance comparison across multiple marketplaces on Public Aicrowd dataset. **<value>** represents the top performance, and **<value>** represents second top performance.

Methods	Marketplaces			
	US	ES	JP	AVG
Independent Training	0.75387	0.83947	0.84432	0.81255
ER (Buffer = 10000)	0.75651	0.84145	0.84959	0.81585
Parameter Isolation				
Prefix FT	0.75293	0.84058	0.83198	0.8085
LoRA FT	0.75271	0.83423	0.83579	0.80758
DoRA FT	0.75478	0.84662	0.8405	0.81397
LoRA Inference-Time Merging Strategy				
linear-combining [12]	0.60432	0.61824	0.65368	0.62541
lora_delta_w_svd [1]	0.7415	0.76616	0.76653	0.75806
Continual Full Model Merging Strategy				
MagMax + Independent FT	0.59107	0.60528	0.64758	0.61464
MagMax + Sequential FT	0.71057	0.69354	0.69642	0.70017
OPCM + Independent FT	0.77668	0.75325	0.74092	0.75695
OPCM + Sequential FT	0.79699	0.81967	0.76454	0.79374
Recent SOTA Methods for CL				
LoRA Sequential Stacking	0.63839	0.63207	0.83028	0.70024
O-LoRA	0.49817	0.63528	0.66559	0.59968
Progressive-Prefix	0.74456	0.79201	0.78641	0.77433
OURS + Variants				
FR-LoRA	0.75551	0.84704	0.85214	0.81823
FR-LoRA + Proj.	0.75638	0.84665	0.85214	0.81839
FR-LoRA + Lr-Decay	0.75601	0.84724	0.85208	0.81845
FR-LoRA + Lr-Decay + Proj.	0.75689	0.84692	0.85208	0.81863
FR-DoRA	0.75591	0.84699	0.85240	0.81843
FR-DoRA + Lr-Decay	0.75692	0.84777	0.85262	0.81911

modules frozen. An orthogonality constraint is applied between LoRA modules of different tasks to reduce representational interference.

7. **Progressive Prefixes** [27]: Independent prefixes are fine-tuned for each marketplace, and during inference, previously trained (and frozen) prefixes are concatenated with the current prefix to promote forward transfer and retain past knowledge.

5.4 Metrics

For classifying relevance and identifying optimal query-product pairs, we use ROC-AUC as our primary metric. Although ranking metrics like precision@k, recall@k and NDCG could be used, however, we opted not to generate results for generating metrics due to the limited number of products per query in our datasets. As for our case we have a sequence of multiple marketplaces, we compare the performances on the basis of average ROC-AUC, i.e., averaged across all the marketplaces.

5.5 Results

Table 1 compares various methods across multiple marketplaces. While independent training performs reasonably well, it lacks cross-market knowledge sharing. Continual learning approaches like FR-LoRA strike a better balance between remembering previous tasks and adapting to new ones. FR-LoRA outperforms strong baselines such as LoRA Finetuning and recent methods like Progressive Prefix. Its variants—especially those with learning rate decay and

Table 3: Table presents ROC-AUC results for the comparison of our methods FR-LoRA and FR-DoRA against EWC. $\langle \text{value} \rangle$ represents the top performance, and $\langle \text{value} \rangle$ represents second top performance. The reported performance figures are expressed relative to independent fine-tuning, which serves as the baseline and is normalized to 1 $\times$.

Method	Marketplaces							AVG
	M_a	M_b	M_c	M_d	M_e	M_f	M_g	
Independent Training	1 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$	1 $\times$
EWC	+0.0260	-0.0115	+0.0596	-0.0037	+0.0042	+0.0299	+0.0705	+0.0250 $\pm$ 0.0059
FR-LoRA	$\langle +0.0364 \rangle$	$\langle -0.0024 \rangle$	$\langle +0.0712 \rangle$	$\langle +0.0030 \rangle$	$\langle +0.0060 \rangle$	$\langle +0.0322 \rangle$	$\langle +0.0753 \rangle$	$\langle +0.0317 \rangle \pm 0.0013$
FR-DoRA	$\langle +0.0288 \rangle$	$\langle -0.0089 \rangle$	$\langle +0.0731 \rangle$	$\langle +0.0061 \rangle$	$\langle +0.0070 \rangle$	$\langle +0.0322 \rangle$	$\langle +0.0728 \rangle$	$\langle +0.0302 \rangle \pm 0.0010$

projection tuning—achieve the highest average ROC-AUC, with FR-LoRA + Lr-Decay + Proj. reaching an improvement of +0.0303.

Table 2 shows similar trends on the Aicrowd dataset. Experience Replay slightly improves over independent training, while parameter isolation methods offer moderate gains. Inference-time and full model merging strategies perform inconsistently, often lacking stability. Among recent methods, Progressive Prefix shows promise but falls short. FR-LoRA and its variants again lead, with FR-DoRA² + Lr-Decay achieving the top score of 0.8191, reinforcing the effectiveness of our approach across diverse marketplaces.

5.6 Ablation Study

5.6.1 Does adding LoRA on top of EWC improve performance? In this section we want to analysis that whether adding LoRA or DoRA on top of EWC, boosts the downstream task performance. To understand that we compare the task-wise performance of EWC alone, with our approach FR-LoRA and FR-DoRA. Table 3 presents the results showing that adding LoRA/ DoRA along with our fisher regularization improves performance on each and every marketplace.

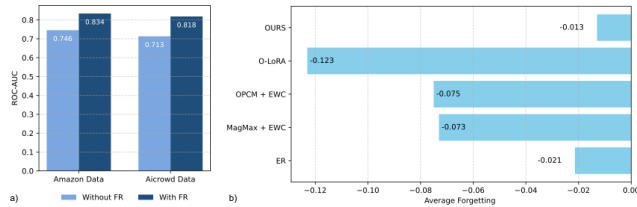


Figure 3: (a) provides a comparison of avg. ROC-AUC for our method with and without FR. (b) provides a comparison of FR-LoRA against other baseline methods for forgetting.

5.6.2 Importance of Fisher Regularization in FR-LoRA. To understand the impact of Fisher regularization in our method, we compared the performance of standard LoRA (without Fisher regularization) to FR-LoRA (with Fisher regularization) across both Amazon and Aicrowd datasets. Note that in both cases the base model is being updated sequentially using EWC, Stage I of our methodology. Figure 3(a) clearly show that incorporating Fisher regularization

significantly improves ROC-AUC scores. This highlights its role in preserving previously learned knowledge while adapting to new marketplaces, making it a crucial component of our continual learning design.

5.6.3 Forgetting Analysis. We measure forgetting by checking how much the ROC-AUC for each task drops after learning new ones. Forgetting is computed as

$$\frac{1}{T} \sum_{t=1}^T \left(\text{ROC-AUC}_t^{(\text{final})} - \max_{k \leq T} \text{ROC-AUC}_t^{(k)} \right)$$

Higher values mean less forgetting. As shown in the Figure 3(b), our method (FR-LoRA) forgets the least, while other methods like ER, MagMax + EWC, OPCM + EWC, and O-LoRA forget more. This shows that Fisher regularization on LoRA along with EWC to update base model helps keep old knowledge better during continual learning.

5.6.4 Latency & Number of Parameter. During training, FR-LoRA updates both the base model and LoRA adapters, resulting in a slightly higher number of trainable parameters compared to methods that train the base model or train PEFTs. However, this additional overhead is only during training. At inference time, LoRA adapters can be merged with the base model weights, ensuring that the final model incurs no additional latency. As a result, inference speed and resource usage remain equivalent to that of a standard base model, making FR-LoRA both efficient and practical for deployment.

6 Conclusion.

We presented FR-LoRA, a simple yet effective continual learning method that combines the benefits of parameter-efficient LoRA fine-tuning with Fisher-based regularization. Our approach consistently outperforms strong baselines and recent state-of-the-art methods across multilingual marketplaces, showing better retention of prior knowledge and strong generalization to new tasks. With minimal inference overhead and strong performance, FR-LoRA offers a practical solution for real-world continual learning scenarios. These gains directly translate to improved query-product relevance, which is critical in e-commerce applications.

²FR-DoRA is same as FR-LoRA with the modification, where in place of LoRA, DoRA is being used.

GenAI Usage Disclosure

In preparing this manuscript, generative AI tools were used solely for language refinement—including improving grammar, style, clarity, and readability of text that was originally authored by the paper’s authors. No sections of the manuscript, including the abstract, related work, methodology, experimental results, or conclusions, were generated by AI systems. All conceptual contributions, technical content, experimental design, data analysis, and interpretations are the authors’ own original work. The authors accept full responsibility for the correctness and integrity of the manuscript.

References

- [1] [n.d.]. Merging Adapters &x2014; AdapterHub documentation — docs.adapterhub.ml. https://docs.adapterhub.ml/merging_adapters.html. [Accessed 13-05-2025].
- [2] Sanjay Agrawal, Faizan Ahemad, and Vivek Varadarajan Sembium. 2025. Rationale-Guided Distillation for E-Commerce Relevance Classification: Bridging Large Language Models and Lightweight Cross-Encoders. In *Proceedings of the 31st International Conference on Computational Linguistics: Industry Track*. 136–148.
- [3] Sanjay Agrawal, Faizan Ahemad, and Vivek Varadarajan Sembium. 2025. Rationale-Guided Distillation for E-Commerce Relevance Classification: Bridging Large Language Models and Lightweight Cross-Encoders. In *Proceedings of the 31st International Conference on Computational Linguistics: Industry Track*, Owen Rambow, Leo Wanner, Marianna Apidianaki, Hend Al-Khalifa, Barbara Di Eugenio, Steven Schockaert, Kareem Darwish, and Apoorv Agarwal (Eds.). Association for Computational Linguistics, Abu Dhabi, UAE, 136–148. <https://aclanthology.org/2025.coling-industry.12/>
- [4] Sanjay Agrawal, Srjana Merugu, and Vivek Sembium. 2023. Enhancing e-commerce product search through reinforcement learning-powered query reformulation. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management*. 4488–4494.
- [5] Sanjay Agrawal, Srjana Merugu, and Vivek Sembium. 2024. Boosting Entity Recognition by leveraging Cross-task Domain Models for Weak Supervision. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*. 4324–4331.
- [6] Sanjay Agrawal, Deep Nayak, and Vivek Varadarajan Sembium. 2025. Multilingual Continual Learning using Attention Distillation. In *Proceedings of the 31st International Conference on Computational Linguistics: Industry Track*. Association for Computational Linguistics, Abu Dhabi, UAE, 91–99. <https://aclanthology.org/2025.coling-industry.8/>
- [7] Sanjay Agrawal and Vivek Sembium. 2025. RTSM: Knowledge distillation with diverse signals for efficient real-time semantic matching in e-commerce. In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 3: Industry Track)*. 9–19.
- [8] Sanjay Agrawal, Vivek Sembium, et al. 2023. Kd-boost: Boosting real-time semantic matching in e-commerce with knowledge distillation. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: Industry Track*. 131–141.
- [9] Qingyao Ai, Ting Bai, Zhao Cao, Yi Chang, Jiawei Chen, Zhumin Chen, Zhiyong Cheng, Shoubin Dong, Zhicheng Dou, Fuli Feng, Shen Gao, Jiafeng Guo, Xiangnan He, Yanyan Lan, Chenliang Li, Yiqun Liu, Ziyu Lyu, Weizhi Ma, Jun Ma, Zhaochun Ren, Pengjie Ren, Zhiqiang Wang, Mingwen Wang, Ji-Rong Wen, Le Wu, Xin Xin, Jun Xu, Dawei Yin, Peng Zhang, Fan Zhang, Weinan Zhang, Min Zhang, and Xiaofei Zhu. 2023. Information Retrieval meets Large Language Models: A strategic report from Chinese IR community. *AI Open* 4 (2023), 80–90. doi:10.1016/j.aiopen.2023.08.001
- [10] Kartikeya Badola, Shachi Dave, and Partha Talukdar. 2023. Parameter-Efficient Finetuning for Robust Continual Multilingual Learning. In *Findings of the Association for Computational Linguistics: ACL 2023*, Anna Rogers, Jordan Boyd-Graber, and Naoaki Okazaki (Eds.). Association for Computational Linguistics, Toronto, Canada, 9763–9780. doi:10.18653/v1/2023.findings-acl.619
- [11] Pietro Buzzega, Matteo Boschini, Angelo Porrello, Davide Abati, and Simone Calderara. 2020. Dark Experience for General Continual Learning: a Strong, Simple Baseline. arXiv:2004.07211 [stat.ML] <https://arxiv.org/abs/2004.07211>
- [12] Alexandra Chronopoulou, Jonas Pfeiffer, Joshua Maynez, Xinyi Wang, Sebastian Ruder, and Priyanka Agrawal. 2024. Language and Task Arithmetic with Parameter-Efficient Layers for Zero-Shot Summarization. arXiv:2311.09344 [cs.CL] <https://arxiv.org/abs/2311.09344>
- [13] Beyza Ermis, Giovanni Zappella, Martin Wistuba, Aditya Rawal, and Cedric Archambeau. 2023. Memory Efficient Continual Learning with Transformers. arXiv:2203.04640 [cs.CL] <https://arxiv.org/abs/2203.04640>
- [14] Zeyu Han, Chao Gao, Jinyang Liu, Jeff Zhang, and Sai Qian Zhang. 2024. Parameter-Efficient Fine-Tuning for Large Models: A Comprehensive Survey. arXiv:2403.14608 [cs.LG] <https://arxiv.org/abs/2403.14608>
- [15] Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, and Weizhu Chen. 2021. LoRA: Low-Rank Adaptation of Large Language Models. *CoRR* abs/2106.09685 (2021). arXiv:2106.09685 <https://arxiv.org/abs/2106.09685>
- [16] Rahul Radhakrishnan Iyer, Rohan Kohli, and Shrimai Prabhumoye. 2020. Modeling Product Search Relevance in e-Commerce. *CoRR* abs/2001.04980 (2020). arXiv:2001.04980 <https://arxiv.org/abs/2001.04980>
- [17] James Kirkpatrick, Razvan Pascanu, Neil C. Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A. Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, Demis Hassabis, Claudia Clopath, Dharshan Kumaran, and Raia Hadsell. 2016. Overcoming catastrophic forgetting in neural networks. *CoRR* abs/1612.00796 (2016). arXiv:1612.00796 <http://arxiv.org/abs/1612.00796>
- [18] Brian Lester, Rami Al-Rfou, and Noah Constant. 2021. The Power of Scale for Parameter-Efficient Prompt Tuning. arXiv:2104.08691 [cs.CL] <https://arxiv.org/abs/2104.08691>
- [19] Xiang Lisa Li and Percy Liang. 2021. Prefix-Tuning: Optimizing Continuous Prompts for Generation. *CoRR* abs/2101.00190 (2021). arXiv:2101.00190 <https://arxiv.org/abs/2101.00190>
- [20] Yiu-Chang Lin, Ankur Datta, and Giuseppe Di Fabbri. 2018. E-commerce Product Query Classification Using Implicit User’s Feedback from Clicks. In *2018 IEEE International Conference on Big Data (Big Data)*. 1955–1959. doi:10.1109/BigData.2018.8622008
- [21] Zhe Lin, Jiwei Tan, Dan Ou, Xi Chen, Shaowei Yao, and Bo Zheng. 2024. Deep Bag-of-Words Model: An Efficient and Interpretable Relevance Architecture for Chinese E-Commerce. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD ’24)*. ACM, 5398–5408. doi:10.1145/3637528.3671559
- [22] Shih-Yang Liu, Chien-Yi Wang, Hongxu Yin, Pavlo Molchanov, Yu-Chiang Frank Wang, Kwang-Ting Cheng, and Min-Hung Chen. 2024. DoRA: Weight-Decomposed Low-Rank Adaptation. arXiv:2402.09353 [cs.CL] <https://arxiv.org/abs/2402.09353>
- [23] David Lopez-Paz and Marc’Aurelio Ranzato. 2022. Gradient Episodic Memory for Continual Learning. arXiv:1706.08840 [cs.LG] <https://arxiv.org/abs/1706.08840>
- [24] Yuren Mao, Yuhang Ge, Yijiang Fan, Wenyi Xu, Yu Mi, Zhonghao Hu, and Yunjun Gao. 2024. A survey on LoRA of large language models. *Frontiers of Computer Science* 19, 7 (2024). doi:10.1007/s11704-024-40663-9
- [25] Daniel Marczak, Bartłomiej Twardowski, Tomasz Trzciński, and Sebastian Cygert. 2024. MagMax: Leveraging Model Merging for Seamless Continual Learning. arXiv:2407.06322 [cs.LG] <https://arxiv.org/abs/2407.06322>
- [26] Duy A. Nguyen, Rishi Kesav Mohan, Van Yang, Pritom Saha Akash, and Kevin Chen-Chuan Chang. 2025. RL-based Query Rewriting with Distilled LLM for online E-Commerce Systems. arXiv:2501.18056 [cs.IR] <https://arxiv.org/abs/2501.18056>
- [27] Anastasia Razdaibiedina, Yuning Mao, Rui Hou, Madian Khabisa, Mike Lewis, and Amjad Almahairi. 2023. Progressive Prompts: Continual Learning for Language Models. In *International Conference on Learning Representations*.
- [28] Zhaochun Ren, Xiangnan He, Dawei Yin, and Maarten de Rijke. 2025. Information Discovery in e-Commerce. arXiv:2410.05763 [cs.IR] <https://arxiv.org/abs/2410.05763>
- [29] David Rolnick, Arun Ahuja, Jonathan Schwarz, Timothy P. Lillicrap, and Greg Wayne. 2018. Experience Replay for Continual Learning. *CoRR* abs/1811.11682 (2018). arXiv:1811.11682 <http://arxiv.org/abs/1811.11682>
- [30] Fan-Keng Sun, Cheng-Hao Ho, and Hung-Yi Lee. 2019. LAMAL: LAnguage Modeling Is All You Need for Lifelong Language Learning. *CoRR* abs/1909.03329 (2019). arXiv:1909.03329 <http://arxiv.org/abs/1909.03329>
- [31] Anke Tang, Enneng Yang, Li Shen, Yong Luo, Han Hu, Bo Du, and Dacheng Tao. 2025. Merging Models on the Fly Without Retraining: A Sequential Approach to Scalable Continual Model Merging. arXiv:2501.09522 [cs.LG] <https://arxiv.org/abs/2501.09522>
- [32] Tian Tang, Zhixing Tian, Zhenyu Zhu, Chenyang Wang, Haiqing Hu, Guoyu Tang, Lin Liu, and Sulong Xu. 2025. LREF: A Novel LLM-based Relevance Framework for E-commerce. *arXiv preprint arXiv:2503.09223* (2025).
- [33] Xiao Wang, Tianze Chen, Qiming Ge, Han Xia, Rong Bao, Rui Zheng, Qi Zhang, Tao Gui, and Xuanjing Huang. 2023. Orthogonal Subspace Learning for Language Model Continual Learning. In *The 2023 Conference on Empirical Methods in Natural Language Processing*. <https://openreview.net/forum?id=L7ZBpZZ8Va>
- [34] Friedemann Zenke, Ben Poole, and Surya Ganguli. 2017. Continual Learning Through Synaptic Intelligence. arXiv:1703.04200 [cs.LG] <https://arxiv.org/abs/1703.04200>