

Continual Learning with Scaled Gradient Projection

Gobinda Saha, Kaushik Roy

Elmore Family School of Electrical and Computer Engineering
Purdue University, West Lafayette, Indiana, USA
gsaha@purdue.edu, kaushik@purdue.edu

Abstract

In neural networks, continual learning results in gradient interference among sequential tasks, leading to catastrophic forgetting of old tasks while learning new ones. This issue is addressed in recent methods by storing the important gradient spaces for old tasks and updating the model orthogonally during new tasks. However, such restrictive orthogonal gradient updates hamper the learning capability of the new tasks resulting in sub-optimal performance. To improve new learning while minimizing forgetting, in this paper we propose a Scaled Gradient Projection (SGP) method, where we combine the orthogonal gradient projections with scaled gradient steps along the important gradient spaces for the past tasks. The degree of gradient scaling along these spaces depends on the importance of the bases spanning them. We propose an efficient method for computing and accumulating importance of these bases using the singular value decomposition of the input representations for each task. We conduct extensive experiments ranging from continual image classification to reinforcement learning tasks and report better performance with less training overhead than the state-of-the-art approaches. Code available at <https://github.com/sahagobinda/sgp>.

Introduction

Continual learning (CL) (Ring 1998; Thrun and Mitchell 1995) aims to endue artificial intelligent (AI) systems with human-like adaptation capability in dynamically changing environments. In this learning paradigm, an AI model, commonly a deep neural network (DNN), learns from a sequence of tasks over time with the aim of accumulating and maintaining past knowledge and transferring it to future tasks. This objective is hard to achieve since standard optimization methods for training DNNs overwrite the parametric representations of past tasks with new input representations during model update (Hadsell et al. 2020). This leads to ‘Catastrophic Forgetting’ (McCloskey and Cohen 1989; Ratcliff 1990) where performance of the past tasks degrades drastically, making continual learning a challenging problem.

Many continual learning approaches for fixed capacity DNNs have been proposed that aim to balance two competing objectives: maintaining stability of past knowledge while providing sufficient plasticity for new learning. One line of works (Kirkpatrick et al. 2017; Schwarz et al. 2018;

Ebrahimi et al. 2020; Saha et al. 2021a; Kao et al. 2021) achieves this goal by penalizing or preventing changes to the most important weights of the model for the past tasks while learning new tasks. Other works minimize forgetting by either storing samples from old tasks in the memory (Robins 1995; Lopez-Paz and Ranzato 2017; Chaudhry et al. 2019b, 2021; Saha and Roy 2023) or synthesizing old data by generative models (Shin et al. 2017) for rehearsal. Despite varying degrees of success, the stability-plasticity balance in such methods breaks down under long sequence of learning. Recently, gradient projection methods (Zeng et al. 2019; Farajtabar et al. 2020; Saha et al. 2021b; Wang et al. 2021; Guo et al. 2022) have shown superior performance to the other approaches. These methods minimize forgetting by minimizing gradient interference among tasks where new tasks are learned in the orthogonal gradient directions or spaces of the old tasks. Despite their remarkable stability (nearly zero forgetting) on the old tasks, performance of new tasks may suffer due to such restrictive gradient updates. To facilitate new learning, orthogonal gradient update is complemented with task similarity-based weight scaling (Lin et al. 2022) and dynamic gradient scaling (Deng et al. 2021). However, these approaches either add significant training overheads or require old data storage, raising privacy concerns.

In this paper, we propose a new method - Scaled Gradient Projection (SGP) for continual learning that improves performance of the new task, while minimizing forgetting of the past tasks. Utilizing the framework introduced by Saha et al. (2021b), we find the bases of the core (important) gradient spaces (CGS) for the past tasks by Singular Value Decomposition (SVD) of the input activations and store them in the memory. From the singular value distributions, we propose an efficient method for finding the (relative) importance of each basis of CGS. For each new task, we decompose the gradients into two components (Figure 1(a)) - one is along the past CGS and another is in the orthogonal direction to the CGS. In our method, the orthogonal gradient component is kept without modification and combined with a scaled version of the gradient component along the CGS. The degree of gradient scaling along each CGS basis is determined by importance of each basis stored in the memory. With this resultant gradient, the model is updated to achieve the CL objective. We evaluate our method in directly comparable settings on image classification datasets (CIFAR-100 and mini-

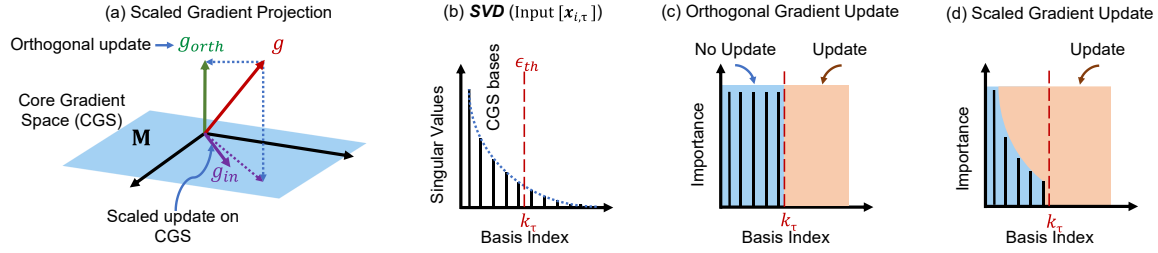


Figure 1: Scaled Gradient Projection (SGP) framework, (a) where new task is learned both in the orthogonal direction (g_{orth}) and along (g_{in}) the core gradient spaces (CGS) of the past tasks. (b) Bases of the CGS and importance of each basis are obtained from SVD on input representations. Unlike (c) orthogonal gradient projection where no update is allowed along the CGS, (d) in SGP, based on the importance of CGS bases, scaled gradient updates are made along those directions enabling better learning.

ImageNet) and reinforcement learning tasks (consisting of six Atari games) and report better performance than the relevant baselines. We summarize our contributions as follows:

- We present Scaled Gradient Projection (SGP) for continual learning which, in addition to orthogonal gradient updates, allows scaled updates along the important gradient spaces of the past tasks, thus improving new learning.
- We propose a method for finding and accumulating importance of each basis of those gradient spaces using singular values from the SVD on input representations. Our method does not store or use data from old tasks and has very little computing overhead.
- We evaluate SGP in diverse application domains: continual image classification and reinforcement learning tasks and obtain better performance than state-of-the-arts. We show, better task-level generalization and minimum forgetting enable superior performance in SGP.

Related Works

Continual learning approaches can be broadly divided into three categories (Lange et al. 2021).

Expansion-based Methods: These methods prevent forgetting by expanding the network capacity. Notably, Progressive Neural Network (PNN) (Rusu et al. 2016) adds new sub-network for each task and freezes them for next tasks. RCL (Xu and Zhu 2018) and BNS (Qin et al. 2021) perform network expansion with reinforcement learning, whereas Dynamic Expansion Network (DEN) (Yoon et al. 2018) uses both model compression and expansion. APD (Yoon et al. 2020) minimizes such network growth by additive decomposition of parameters into task-specific and task-shared parts. In contrast, we learn continually in a fixed capacity model.

Regularization-based Methods: These methods penalize changes in network parameters when learning new tasks based on their importance for prior tasks. Notably, Elastic Weight Consolidation (EWC) (Kirkpatrick et al. 2017) computes this importance from the Fisher diagonal matrix after training, whereas Zenke et al. (2017) find them during training from the loss sensitivity of the parameters. Other methods completely prevent updates to these important parameters by masking. For instance, PackNet (Mallya and Lazebnik 2018) uses weight-level binary masks and HAT (Serrà et al. 2018) learns attention masks for neurons to prevent gradient updates. BLIP (Shi et al. 2021) uses weight quantization and protects past information by bit-level freezing.

Memory-based Methods: Memory-based methods can be divided into two sub-categories - (a) Experience Replay and (b) Gradient Projection based methods. In experience replay (ER) (Rebuffi et al. 2017; Chaudhry et al. 2019b) data samples from the past tasks are stored in the memory and during new tasks, the model is jointly optimized with new data and memory data to mitigate forgetting. Meta Experience Replay (Riemer et al. 2019) and Look-ahead MAML (Gupta et al. 2020) combine experience replay with meta learning to achieve maximum knowledge transfer with minimum forgetting, whereas Dark Experience Replay (Buzzega et al. 2020) augments memory replay with logit distillation loss to improve performance. Alternatively, Gradient Episodic Memory (GEM) (Lopez-Paz and Ranzato 2017) and Averaged-GEM (A-GEM) (Chaudhry et al. 2019a) use memory data to compute gradient constraints for the new task so that loss on the past task does not increase.

Gradient projection methods, in contrast, keep memory of gradient information from the past tasks and update model in the orthogonal directions of the past tasks to minimize forgetting. For instance, Orthogonal Weight Modulation (OWM) (Zeng et al. 2019) updates model in the orthogonal input directions of the old tasks, whereas Wang et al. (2021) performs model update in the null-space of the previous tasks. Orthogonal Gradient Descent (Farajtabar et al. 2020) stores gradient directions of past tasks in the memory and project new gradient orthogonally to these memory gradients. Saha et al. (2021b) reduced training complexity and memory overhead of such methods in Gradient Projection Memory (GPM), where important gradient spaces for the past tasks are stored and new learning progresses in the orthogonal directions to these spaces. To relax orthogonal gradient constraint and encourage knowledge transfer to the new tasks, Trust Region Gradient Projection (TRGP) (Lin et al. 2022) augmented GPM with task-correlation based scaling of the weights of the old tasks. To achieve similar objective, Deng et al. (2021) in FS-DGPM combined experience replay with gradient projection to dynamically scale the gradients. In a similar vein, we propose to improve new learning by scaled gradient projection. In contrast to previous methods, we do not learn these scales during training. Rather, we find and accumulate these scales in a one-shot manner after learning each task without using any old data. Thus we achieve better performance with significantly less training overhead while preserving data privacy.

Method

In this section, we describe the steps for continual learning with our Scaled Gradient Projection (SGP) method. We consider supervised learning setup (Saha et al. 2021b) where T tasks are sequentially learned. Each task has a task descriptor, $\tau \in \{1, 2, \dots, T\}$ with a dataset, $\mathbb{D}_\tau = \{(\mathbf{x}_{i,\tau}, \mathbf{y}_{i,\tau})_{i=1}^{n_\tau}\}$ containing n_τ example pairs. We train a neural network having parameter set, $\mathbb{W}_\tau = \{(\mathbf{W}_\tau^l)_{l=1}^L\}$, where L is the number of layers. For each task we follow a three step process - (1) **training** the model on current task, (2) **finding bases** of the important gradient spaces for that task, and (3) **computing importance** of individual basis. These bases and corresponding importance are computed for each layer of the network and stored in the memory. The method described below is for a typical layer of a network and is generally applicable to all the layers. Hence we drop the layer notation. We use $\mathbf{x}_{i,\tau}$ to denote input or corresponding representation (activation) of input inside DNN for task τ interchangeably.

Learning Task, $\tau = 1$

Train: Starting from random initialization ($\mathbb{W}_0$), we train the model on dataset, $\mathbb{D}_1$ and obtain learned parameter set $\mathbb{W}_1$. During training on the first task no gradient scaling or constraint is imposed.

Find Important Gradient Space: Next, we find important gradient spaces for this task using SVD on the input representations as described in GPM. For that, we randomly select n_s training samples from $\mathbb{D}_1$ and pass it through the network. Then, for each layer we construct a representation matrix, $\mathbf{R}_1 = [\mathbf{x}_{1,1}, \mathbf{x}_{2,1}, \dots, \mathbf{x}_{n_s,1}]$ by concatenating the corresponding input representations. To find the top- k_1 important basis of this input representation, SVD is performed on $\mathbf{R}_1 = \mathbf{U}_1 \mathbf{\Sigma}_1 \mathbf{V}_1^T$, where $\mathbf{U}_1$ and $\mathbf{V}_1$ are orthonormal matrices, and $\mathbf{\Sigma}_1$ has sorted singular values ($\sigma_{i,1}$) along its diagonal. This step is followed by k_1 -rank approximation $(\mathbf{R}_1)_{k_1}$ of $\mathbf{R}_1$ using the norm-based criteria :

$$\|(\mathbf{R}_1)_{k_1}\|_F^2 \geq \epsilon_{th} \|\mathbf{R}_1\|_F^2. \quad (1)$$

Here, $\|\cdot\|_F^2$ is the Frobenius norm of the matrix. A brief description of SVD and k -rank approximation is given in Appendix A. Threshold hyperparameter, $\epsilon_{th} \in (0, 1)$ controls the value of k_1 selected. The first k_1 (column) vectors in $\mathbf{U}_1$ spans the most important space of inputs (representations) for task 1. Saha et al. (2021b) showed that these bases equivalently span the most important gradient space for task 1. We store these bases in $\mathbf{M} = [\mathbf{u}_{1,1}, \mathbf{u}_{2,1}, \dots, \mathbf{u}_{k_1,1}]$, which is denoted as Gradient Projection Memory (GPM).

Compute Importance of Gradient Bases: Typical gradient projection methods (e.g. Saha et al. 2021b) assume all the bases in GPM are equally important, hence no gradient step is taken along these directions for the new tasks to prevent forgetting. We find this assumption too restrictive for new learning and propose to assign variable importance to different bases. Our motivation comes from the observation that singular values ($\sigma_{i,1}$) corresponding to these bases ($\mathbf{u}_{i,1}$) have non-uniform distribution (Figure 1(b)) that implies their non-uniform importance for input or gradient preservation. Leveraging this distribution, we assign importance to individual basis. For that, we first get the singular

value vector, $\boldsymbol{\sigma}_1 = \text{diag}(\mathbf{\Sigma}_1, k_1)$, where $\boldsymbol{\sigma}_1 \in \mathbb{R}^{k_1 \times 1}$, corresponding to the top- k_1 basis in $\mathbf{U}_1$. We then compute importance for i^{th} basis, where $i \in [1, k_1]$, by following:

$$\lambda_{i,1} = \frac{(\alpha + 1)\sigma_{i,1}}{\alpha\sigma_{i,1} + \max(\boldsymbol{\sigma}_1)}, \quad (2)$$

where α is a non-negative scale coefficient hyperparameter. The value of $\lambda_{i,1}$ will range from 0 to 1 as we are concerned with the non-negative singular values. We construct importance vector for this task by $\boldsymbol{\lambda}_1 = [\lambda_{1,1}, \dots, \lambda_{k_1,1}]^T$. Equation 2 ensures that maximum importance (of 1) is assigned to the basis with highest singular value and other bases are given importance (< 1) relative to this maximum. In our formulation, $\lambda_{i,1} = 1$ means no gradient step is allowed along the corresponding basis direction for the new tasks whereas along other bases gradients are scaled by the factor of $(1 - \lambda_{i,1})$. In GPM (Saha et al. (2021b)), all the bases in the memory are assumed to have maximum importance (of 1). This prevents any update along those directions (Figure 1(c)). In contrast, we allow a scaled gradient update along those bases (Figure 1(d)) enabling higher plasticity for new tasks, while importance based scaling ensures desirable stability of past tasks. Moreover, setting α to very high value (ideally ∞) in Equation 2 would push all the importance values to 1. Thus for a large value of α our method converges to GPM. We transfer the basis matrix, $\mathbf{M}$ and accumulated importance vector $\boldsymbol{\lambda}$ to the next task. Here, $\boldsymbol{\lambda} = \boldsymbol{\lambda}_1$, $\boldsymbol{\lambda} \in \mathbb{R}^{k \times 1}$ with $k = k_1$ is the number of bases after task 1.

Learning Task, $\tau \in [2, T]$

Train: We learn the τ^{th} task sequentially using dataset, $\mathbb{D}_\tau$ only. Let L_τ denote loss for τ^{th} task. To prevent catastrophic forgetting and facilitate new learning we perform scaled gradient projection of new gradient, $\nabla_{\mathbf{W}_\tau} L_\tau$ as follows:

$$\nabla_{\mathbf{W}_\tau} L_\tau = \nabla_{\mathbf{W}_\tau} L_\tau - (\mathbf{M} \boldsymbol{\Lambda} \mathbf{M}^T)(\nabla_{\mathbf{W}_\tau} L_\tau), \quad (3)$$

where importance matrix, $\boldsymbol{\Lambda}$ is a diagonal matrix of size $\mathbb{R}^{k \times k}$ containing $\boldsymbol{\lambda}$ in its diagonal. Here k is the number of bases in $\mathbf{M}$ till task $\tau - 1$. Such gradient projection ensures (a) the gradient components along orthogonal directions to $\mathbf{M}$ (GPM) will not be changed (as in Saha et al. (2021b)), but (b) along the GPM bases the gradient components will be scaled by $(1 - \lambda_i)$ based on the stored importance in $\boldsymbol{\lambda}$.

Update Important Gradient Space: At the end of task τ , we update $\mathbf{M}$ by adding the important gradient space for this task. For that, we construct $\mathbf{R}_\tau = [\mathbf{x}_{1,\tau}, \mathbf{x}_{2,\tau}, \dots, \mathbf{x}_{n_s,\tau}]$ using samples from $\mathbb{D}_\tau$ only. Now, SVD on $\mathbf{R}_\tau$ may produce some basis directions that are already present in $\mathbf{M}$ or are linear combinations of the bases in $\mathbf{M}$. These redundant bases are removed with the following step:

$$\hat{\mathbf{R}}_\tau = \mathbf{R}_\tau - (\mathbf{M} \mathbf{M}^T) \mathbf{R}_\tau = \mathbf{R}_\tau - \mathbf{R}_{\tau, \mathbf{M}}. \quad (4)$$

Next, SVD is performed on $\hat{\mathbf{R}}_\tau = \hat{\mathbf{U}}_\tau \hat{\mathbf{\Sigma}}_\tau \hat{\mathbf{V}}_\tau^T$ and new k_τ bases are chosen for minimum k_τ satisfying the criteria:

$$\|\mathbf{R}_{\tau, \mathbf{M}}\|_F^2 + \|(\hat{\mathbf{R}}_\tau)_{k_\tau}\|_F^2 \geq \epsilon_{th} \|\mathbf{R}_\tau\|_F^2. \quad (5)$$

Gradient space in GPM is updated (after $\boldsymbol{\lambda}$ update) by adding these new bases to $\mathbf{M}$ as $\mathbf{M} = [\mathbf{M}, \hat{\mathbf{u}}_{1,\tau}, \dots, \hat{\mathbf{u}}_{k_\tau,\tau}]$.

Update Importance of Gradient Bases: This step involves assigning importance to the newly added k_τ bases and updating the importance of the old k bases. Such importance update comes with several challenges. Firstly, SVD is performed on a part of the representation $\hat{\mathbf{R}}_\tau$, thus $\hat{\Sigma}_\tau$ does not capture the full singular values distributions in $\mathbf{R}_\tau$. Hence, relative importance computation for the new k_τ bases using only the singular values of $\hat{\Sigma}_\tau$ based on Equation 2 would not be appropriate. Secondly, from Equation 4 and Inequality 5 we only know that old k bases in $\mathbf{M}$ are responsible for capturing a fraction ($\|\mathbf{R}_{\tau,M}\|_F^2 / \|\mathbf{R}_\tau\|_F^2$) of total norm of $\mathbf{R}_\tau$. However, computing their individual contributions in $\|\mathbf{R}_{\tau,M}\|_F^2$, hence their importance for task τ is non-trivial.

Here we propose a method for finding and updating the importance without using data from old tasks. First, we perform SVD on $\mathbf{R}_{\tau,M} = \mathbf{U}_{\tau,M} \Sigma_{\tau,M} \mathbf{V}_{\tau,M}^T$. By construction (Equation 4) the first k (column) vectors in $\mathbf{U}_{\tau,M}$ will be linear combinations of the bases in $\mathbf{M}$. Interestingly, SVD provides singular value vector, $\sigma_{\tau,M} = \text{diag}(\Sigma_{\tau,M}, k)$ where each element signifies the importance of corresponding bases in $\mathbf{U}_{\tau,M}$ for the current task τ . Since these are redundant bases and will not be stored in $\mathbf{M}$, we transfer their singular values (hence importance) to the bases of $\mathbf{M}$. For that we create following projection coefficient matrix:

$$\mathbf{C} = \mathbf{M}^T \mathbf{U}_{\tau,M}, \quad (6)$$

where each element $c_{i,j}$ is the dot product between i^{th} columns of $\mathbf{M}$ and j^{th} columns of $\mathbf{U}_{\tau,M}$. Then we compute ‘surrogate singular values’ for the bases in $\mathbf{M}$ by:

$$\sigma'_{\tau,M} = \sqrt{(\mathbf{C} \odot \mathbf{C})(\sigma_{\tau,M})^2}. \quad (7)$$

Here $\odot$ denotes element-wise multiplication, $(\cdot)^2$ and $\sqrt{(\cdot)}$ denote element-wise square and square root operations respectively. Next, we obtain the singular value vector, $\hat{\sigma}_\tau = \text{diag}(\hat{\Sigma}_\tau, k_\tau) \in \mathbb{R}^{k_\tau \times 1}$, for the newly added k_τ bases for the current task from SVD of $\hat{\mathbf{R}}_\tau$. Finally, we create the full (k_τ -rank approximated) singular value vector, σ_τ as:

$$\sigma_\tau = \begin{bmatrix} \sigma'_{\tau,M} \\ \hat{\sigma}_\tau \end{bmatrix} \in \mathbb{R}^{(k+k_\tau) \times 1}. \quad (8)$$

The above construction of σ_τ ensures that the Inequality 5 is satisfied. That can be re-written as:

$$\sum_{i=1}^k (\sigma'_{i,\tau,M})^2 + \sum_{i=1}^{k_\tau} (\hat{\sigma}_{i,\tau})^2 \geq \epsilon_{th} \|\mathbf{R}_\tau\|_F^2. \quad (9)$$

Therefore, we can use σ_τ to obtain the basis importance vector, $\lambda_\tau = f(\sigma_\tau, \alpha)$ from Equation 2 for the given τ^{th} task. Finally, we compute the accumulated importance over tasks and update the accumulated importance vector λ . For $i \in [1, k]$ we update importance of old k bases by:

$$\lambda_i = \begin{cases} 1, & \text{if } (\lambda_i + \lambda_{i,\tau}) \geq 1 \\ \lambda_i + \lambda_{i,\tau}, & \text{otherwise} \end{cases} \quad (10)$$

We then add the importance of new k_τ bases in λ as $\lambda = [\lambda^T, \lambda_{k+1,\tau}, \dots, \lambda_{k+k_\tau,\tau}]^T$. After task τ , the updated number of bases is $k = k + k_\tau$. We transfer $\mathbf{M}$ and λ to the next task and repeat the same procedure. The pseudocode of the algorithm is given in Algorithm 1 in the Appendix B.

Continual Image Classification Tasks

Experimental Setup

Datasets and Models: First, we evaluate our method on standard image classification benchmarks (Saha et al. 2021b; Lin et al. 2022) including **Split CIFAR-100**, **CIFAR-100 Superclass** and **Split miniImageNet** datasets. The Split CIFAR-100 is constructed by dividing 100 classes of CIFAR-100 (Krizhevsky 2009) into 10 tasks with 10 distinct classes per task. CIFAR-100 Superclass (Yoon et al. 2020) is divided into 20 tasks where each task contains 5 semantically related classes from CIFAR-100 dataset, whereas Split miniImageNet is constructed by dividing 100 classes of miniImageNet (Vinyals et al. 2016) into 20 tasks with 5 distinct classes per task. We do not use any data augmentation in our experiments. Similar to Saha et al. (2021b), we use 5-layer AlexNet for Split CIFAR-100, LeNet for CIFAR-100 Superclass, and reduced ResNet18 for Split miniImageNet experiments. We train and test all the methods in task-incremental learning (Hsu et al. 2018) setup where each task has a separate classifier head. Details of dataset statistics and network architectures are given in Appendix C.

Baselines: As baselines for comparison, we use EWC (Kirkpatrick et al. 2017) and HAT (Serrà et al. 2018) from the regularization-based methods, OWM (Zeng et al. 2019), GPM (Saha et al. 2021b), TRGP (Lin et al. 2022) and FS-DGPM (Deng et al. 2021) from gradient projection methods, and Experience Replay with reservoir sampling (ER_Res) (Chaudhry et al. 2019b) and A-GEM (Chaudhry et al. 2019a) from replay methods. To provide an upper limit for performance, we include an oracle baseline - Multitask where all the tasks are jointly learned in a single network using the entire dataset. Additionally, for CIFAR-100 Superclass we provide comparisons with PNN, RCL, DEN, and APD (Yoon et al. 2020) from the expansion-based methods.

Training Details: All the baselines and our method are trained with plain stochastic gradient descent (SGD). Following GPM setup, each task in Split CIFAR-100 and CIFAR-100 Superclass is trained for 200 and 50 epochs respectively with early stopping criteria with batch size of 64, whereas each Split miniImageNet task is trained for 10 epochs with batch size of 10. We use the same threshold values (ϵ_{th}) as GPM and use scale coefficient (α) of 1, 3 and 10 for Split miniImageNet, CIFAR-100 Superclass and Split CIFAR-100 datasets respectively. More details of training setup, hyperparameters considered in baselines and our methods and implementation are given in Appendix C.

Evaluation Metrics: For evaluations, we use two performance metrics (Lopez-Paz and Ranzato 2017): (1) **ACC** - measures average test classification accuracy of all tasks and (2) **BWT** (backward transfer) - measures influence of new learning on the past knowledge where negative BWT implies forgetting. ACC and BWT are defined as:

$$\text{ACC} = \frac{1}{T} \sum_{i=1}^T R_{T,i}; \text{BWT} = \frac{1}{T-1} \sum_{i=1}^{T-1} R_{T,i} - R_{i,i}. \quad (11)$$

Here, T is the total number of tasks. $R_{T,i}$ is the accuracy of the model on i^{th} task after learning the T^{th} task sequentially.

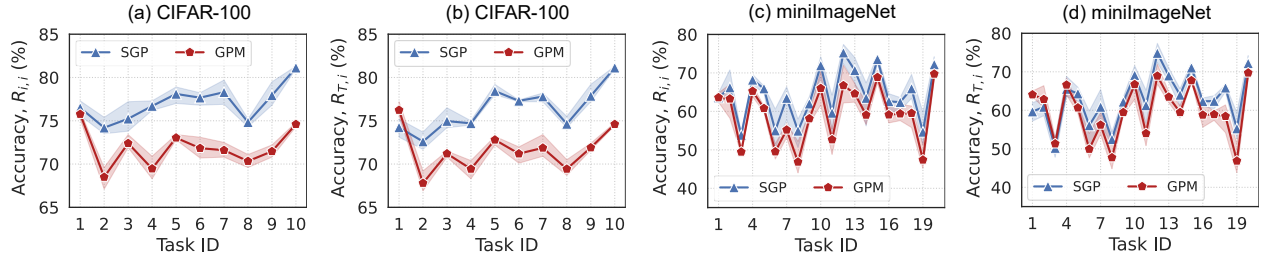


Figure 2: Test accuracy of each Split CIFAR-100 task (a) after learning that task, (b) after learning all the tasks. Test accuracy of each Split miniImageNet task (c) after learning that task, (d) after learning all the tasks.

Methods	CIFAR-100		miniImageNet	
	ACC (%)	BWT	ACC (%)	BWT
Multitask [†]	79.58 ± 0.54	-	69.46 ± 0.62	-
OWM [†]	50.94 ± 0.60	-0.30 ± 0.01	-	-
A-GEM [†]	63.98 ± 1.22	-0.15 ± 0.02	57.24 ± 0.72	-0.12 ± 0.01
ER_Res [†]	71.73 ± 0.63	-0.06 ± 0.01	58.94 ± 0.85	-0.07 ± 0.01
EWC [†]	68.80 ± 0.88	-0.02 ± 0.01	52.01 ± 2.53	-0.12 ± 0.03
HAT [†]	72.06 ± 0.50	-0.00 ± 0.00	59.78 ± 0.57	-0.03 ± 0.00
FS-DGPM [‡]	74.33 ± 0.31	-0.03 ± 0.00	-	-
TRGP	74.46 ± 0.32 [*]	-0.01 ± 0.00 [*]	60.93 ± 0.94	-0.00 ± 0.00
GPM [†]	72.48 ± 0.40	-0.00 ± 0.00	60.41 ± 0.61	-0.00 ± 0.00
SGP (ours)	76.05 ± 0.43	-0.01 ± 0.00	62.83 ± 0.33	-0.01 ± 0.01

Table 1: Results (mean ± std in 5 runs) on Split CIFAR-100 and Split miniImageNet. [†], [‡] and ^{*} denote the results from GPM, FS-DGPM and TRGP respectively.

Results and Discussions

Performance Comparison: First, we compare the ACC and BWT of our method with the baselines. Table 1 shows the comparisons for Split CIFAR-100 and Split miniImageNet dataset. In Split CIFAR-100 tasks, SGP obtains highest ACC (76.05 ± 0.43%) with accuracy gain of ~ 3.5% over GPM and ~ 1.6% over TRGP and FS-DGPM. Compared to the best performing regularization method HAT and replay method ER_Res, SGP obtains ~ 4% improvement in accuracy. In terms of BWT, SGP shows ~ 1% forgetting which is lower than FS-DGPM, comparable to TRGP and slightly higher than GPM and HAT. Similar trend is observed for longer (20) tasks sequence in miniImageNet dataset, where SGP obtains accuracy gain of ~ 2.4% over GPM and ~ 2% over TRGP with BWT of -1.5%. Compared to GPM, this slight increase in forgetting (negative BWT) in SGP is intuitive, as accuracy gain in SGP comes from selective model updates along the important gradient spaces of the past tasks which interfere (minimally) with old learning. Finally, we compare the performance of gradient projection methods with the expansion-based methods in Table 2 for 20-task CIFAR-100 Superclass dataset. Here, SGP achieves highest ACC (59.05 ± 0.50%) with -1.5% BWT obtaining accuracy improvement of 2.2% over APD with 30% fewer network parameters. It has 1.3% ACC gain over GPM and per-

forms marginally better than TRGP and FS-DGPM. Overall, improved performance in SGP indicates that scaled gradient projections offer better stability-plasticity balance in CL compared to strict orthogonal projections.

Task-Level Generalization and Forward Transfer: To quantify task-level generalization and forward transfer, in Figure 2(a) and (c) we plot $R_{i,i}$, which is the test accuracy of i^{th} task immediately after learning i^{th} task during CL process, for all Split CIFAR-100 and Split miniImageNet tasks. Analysis for CIFAR-100 Superclass is given in the Appendix D (Figure 6). Compared to GPM, in SGP each new task ($\tau > 1$) has better accuracy which indicates relaxing constraints with gradient scaling (along CGS) facilitates learning and improves generalization capability of the new tasks. Forward transfer (FWT) (Veniat et al. 2021; Lin et al. 2022) is quantified by the difference in accuracy when a task is learned sequentially and when that task is learned standalone in a single network. Positive FWT thus means sequential learning has improved the performance, hence a desired attribute of CL. Two CL methods (e.g. A & B) can be compared in terms of relative FWT (Lin et al. 2022) as:

$$\Delta_{A-B} = \frac{1}{T} \sum_{i=1}^T R_{i,i}^A - R_{i,i}^B. \quad (12)$$

From Figure 2(a) and (c), we find that relative forward transfer between SGP and GPM, $\Delta_{SGP-GPM}$ is approximately +5% and +4% for Split CIFAR-100 and Split miniImageNet respectively. This results again show better task-level generalization is achieved in CL with SGP.

In Figure 2(b) and (d) we plot $R_{T,i}$ - final test accuracy of each task after learning all T tasks, for Split CIFAR-100 and miniImageNet. Comparing them to Figure 2(a) and (c), we find that there is a slight difference in final and initial task accuracies ($R_{T,i} - R_{i,i}$) in SGP which translates to backward transfer according to Equation 11. Moreover, earlier tasks have a slightly higher drop, whereas later tasks show negligible to no drop in accuracy. Overall, almost all the tasks (except few earlier ones) have higher final accuracy in SGP than GPM which translates to better average accuracy (ACC).

Training Time and Memory Overhead: In Figure 3(a) we compare the training complexity of gradient projection methods in terms of relative (wall-clock) time needed for sequentially training all the tasks. Here the normalization is performed with respect to GPM. Compared to GPM, SGP tasks only up to 17% more time which is used for basis importance calculations. TRGP takes up to 78% more

Metric	Methods									
	STL [†]	PNN [†]	DEN [†]	RCL [†]	APD [†]	EWC [‡]	FS-DGPM [‡]	TRGP [†]	GPM [†]	SGP (ours)
ACC (%)	61.00	50.76	51.10	51.99	56.81	50.26	58.81	58.25	57.72	59.05
Capacity (%)	2000	271	191	184	130	100	100	100	100	100

Table 2: Results (averaged over 5 task-order) for CIFAR-100 Superclass. [†] and [‡] denote the results from TRGP and FS-DGPM respectively. Single task learning (STL) is a non-continual learning baseline where each task is learned in a separate network.

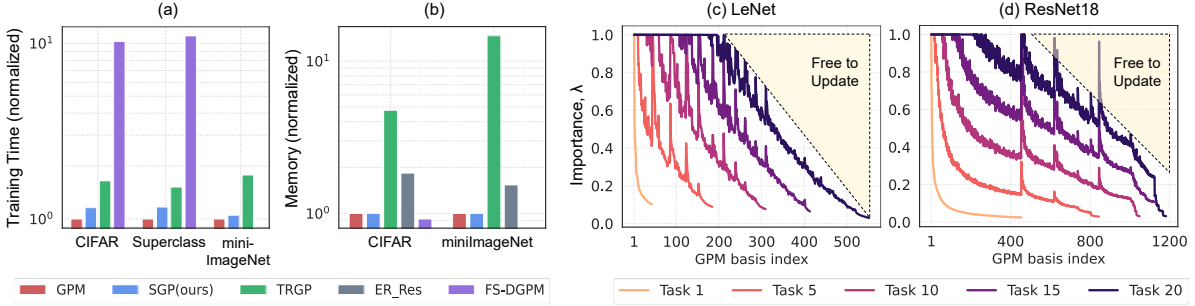


Figure 3: (a) Training time and (b) Memory comparisons (both normalized with respect to GPM). Smaller means better. Basis importance distribution at the last layer before classifier of the model in (c) CIFAR-100 Superclass and (d) miniImageNet tasks.

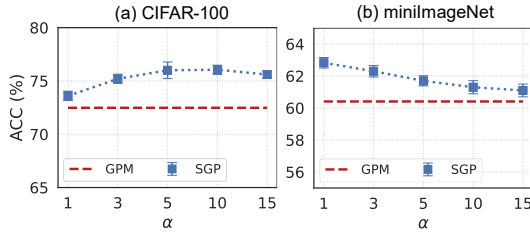


Figure 4: Impact of varying α on SGP performance.

time which includes trust regions selection and weight scale learning. Whereas, FS-DGPM takes an order of magnitude more time than GPM and SGP primarily due to iterative sharpness evaluation steps used for dynamic gradient projection, potentially limiting its application in deeper networks. More details of time measurements and comparison with other methods are provided in Appendix D. Next, in Figure 3(b), we show the extra memory used by different memory-based methods either for gradient or data storage. Compared to GPM, SGP takes only 0.1% extra memory for storing basis importance. TRGP uses up to an order of magnitude more memory for storing weight scales for each layer for each task. Thus, SGP provides performance gain in CL with significantly less training time and memory overhead.

Importance Accumulation Dynamics: Next, we analyze the basis importance accumulation dynamics across tasks in SGP. Figure 3(c) and (d) show importance distributions across different tasks at the last layer before the classifier of LeNet and ResNet18 in CIFAR-100 Superclass and miniImageNet experiments. These distributions reflect non-uniformity of the singular value distributions from which they are generated and help us understand the degree of gradient constraint in SGP optimization. For task 1, only one basis per layer has the importance of $\lambda_1 = 1$, whereas large number of bases (especially in Figure. 3(d)) have very low

importance. Unlike GPM, where gradients along all these directions are blocked, in SGP large steps are encouraged along the less important basis directions relaxing the gradient constraints significantly. As new tasks are learned, new bases are added and importance of the old bases is updated (Equation 10). Due to these accumulative updates, importance of the bases gradually moves towards unity making the optimization gradually restrictive along those directions to encounter forgetting. This process still leaves significant degree of freedom for future learning. As in the last layer of LeNet and ResNet18, even after learning 20 tasks, SGP fully restricts updates along only $\sim 37\%$ and $\sim 18\%$ of the stored basis directions, whereas along the remaining directions gradient steps with variable scaling are allowed. In Appendix D (Figure 7 & 8), importance (λ) distributions in all the layers (except classifier) for Split CIFAR-100 and CIFAR-100 Superclass are shown where a similar trend is observed.

Impact of Scale Coefficient (α): In Figure 4 we show the impact of varying α on SGP performance. In Figure 3(c), (d) we observe that the importance distributions vary with the network architecture and dataset. Thus depending on these factors optimum performance will occur at different α . Ultimately, at large enough α the SGP performance will be similar to GPM. For Split CIFAR-100 we find that for α in the range of 5 to 10, SGP has maximum ACC. In miniImageNet at $\alpha = 1$ best performance is obtained and increasing α pushes SGP performance towards GPM. CIFAR-100 Superclass results are given in Appendix D (Figure 6(c)).

Continual Reinforcement Learning Tasks

Experimental Setup

Unlike majority of CL works, we test the scalability of our algorithm in challenging yet highly relevant continual reinforcement learning benchmarks. For that, we fol-

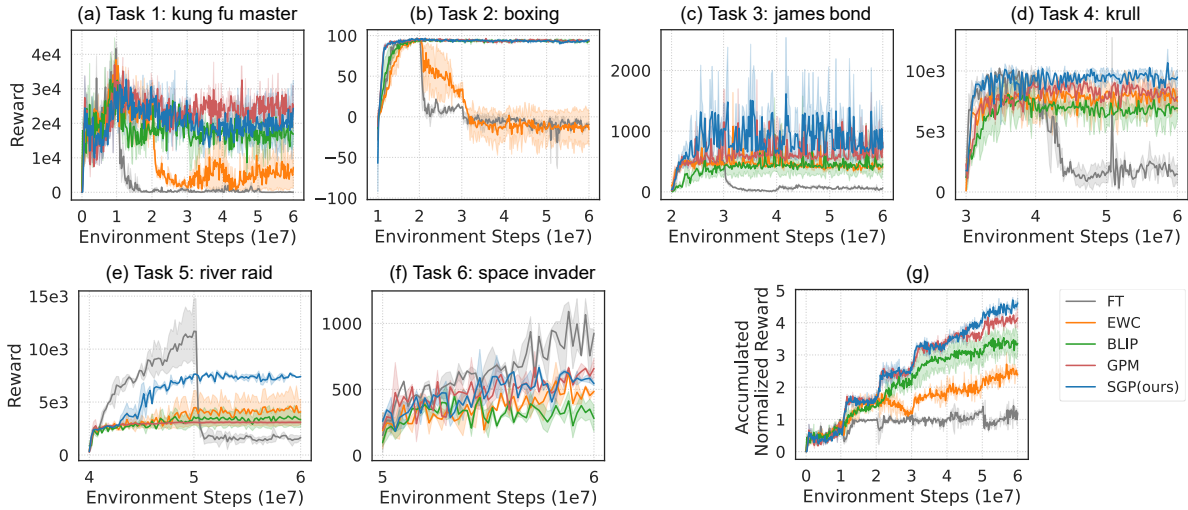


Figure 5: (a)-(f) Rewards on six sequential Atari games (tasks). Each task is trained for 10 million environment steps, and 60 million environment steps are used in total. Results are averaged over 3 random seeds. (g) Accumulated Normalized Rewards.

low the state-of-the-art setup in BLIP (Shi et al. 2021), where the agent sequentially learns to play the following six Atari (Mnih et al. 2013) games: **kung fu master, boxing, james bond, krull, river raid, space invaders**. We use a PPO agent (Schulman et al. 2017) with 3 convolution layers (32-64-128 filters) followed by a fully connected layer (1024 neurons) that is trained with initial learning rate 2.5×10^{-4} and entropy regularization coefficient 0.01. We sample 10 million steps in each (task) environment to train our agents. We use EWC, BLIP, and naive baseline of sequential fine-tuning (FT) for comparison. Like baselines, we could not use standard Adam optimizer (Kingma and Ba 2014) in GPM and SGP as it shows catastrophic forgetting when projected gradients are passed into Adam (see Appendix Figure 9). Hence we take the gradient output from Adam and then perform orthogonal or scaled projection on that (see modified Adam algorithm in Appendix E). For GPM and SGP, $\epsilon_{th} = 0.995$ and for SGP $\alpha = 25$ is used.

Results and Discussions

To evaluate and compare performance, following BLIP, we score each game (task) by average reward obtained after playing that game for 10 episodes. In Figure 5(a)-(f), we plot rewards of the PPO agents on each task from the point they start to train on the task to the end of the whole continual learning process. In each of these plots, for the first 10 million steps, the agent is trained on that task and at regular intervals within that period the agent’s scores (reward) on that task are plotted. In the remaining period, the agent is not trained on that task, only its performance on that task is tested and plotted. From these plots, we observe that FT trained agent suffers from catastrophic forgetting as its reward on one task rapidly decreases after learning next task. EWC trained agent forgets first two tasks and maintains the performance of last four tasks. BLIP trained agent, on the other hand, maintains performance for each task, however from task 3 to 6, compared to SGP, it obtains significantly less rewards primarily due to capacity constraints. Similar

trend is observed in GPM, especially in task 5 GPM underperforms SGP trained agent by a large margin due to strict gradient constraints. Overall, in all these tasks SGP has comparable or better (especially later tasks) performance than other methods taking advantage of scaled gradient updates.

Finally, to compare the performance of each method in an overall manner, in Figure 5(g) we plot the accumulated reward, which is the sum of rewards on all learned tasks at each environment step. We normalize rewards for each task on a scale of 0 to 1. This means an agent’s accumulated normalized reward should increase steadily from 0 to 6 if it learns and remembers all six tasks. In contrast, when an agent continually forgets past tasks, which is the case for FT in Figure 5, then the accumulated normalized reward should oscillate around 1. In this metric, SGP outperforms all the other methods in the entire continual learning sequence, particularly accumulating 12% and 36% more rewards compared to GPM and BLIP respectively. These experiments show that SGP can provide stability to past learning, while improving new learning by better utilizing gradient capacity in DNNs, even under long sequence of gradient updates.

Conclusions

In summary, to improve continual learning performance, we propose a scaled gradient projection method that combines orthogonal gradient projections with a scaled gradient updates along the past important gradient spaces. We introduce a principled method for finding the importance of each basis spanning these spaces and use them to guide the gradient scaling. With quantitative analyses, we show that our method enables better task-level generalization with minimum forgetting by effectively utilizing the capacity of gradient spaces in the DNNs. Moreover, we show our method has very little training overhead and is scalable to CL applications in multiple domains including image classification and reinforcement learning. Future works should take complementary benefits of scaled gradient projection by combining it with other methods in various CL setups.

Acknowledgements

This work was supported in part by the National Science Foundation, Vannevar Bush Faculty Fellowship, Army Research Office, MURI, DARPA AI Exploration (ShELL), and by Center for Brain-Inspired Computing (C-BRIC), one of six centers in JUMP, a SRC program sponsored by DARPA.

References

- Buzzega, P.; Boschini, M.; Porrello, A.; Abati, D.; and Calderara, S. 2020. Dark Experience for General Continual Learning: a Strong, Simple Baseline. In *Advances in Neural Information Processing Systems*, volume 33, 15920–15930. Curran Associates, Inc.
- Chaudhry, A.; Gordo, A.; Dokania, P. K.; Torr, P. H.; and Lopez-Paz, D. 2021. Using Hindsight to Anchor Past Knowledge in Continual Learning. In *AAAI*.
- Chaudhry, A.; Ranzato, M.; Rohrbach, M.; and Elhoseiny, M. 2019a. Efficient Lifelong Learning with A-GEM. In *International Conference on Learning Representations*.
- Chaudhry, A.; Rohrbach, M.; Elhoseiny, M.; Ajanthan, T.; Dokania, P. K.; Torr, P. H. S.; and Ranzato, M. 2019b. Continual Learning with Tiny Episodic Memories. *ArXiv*, abs/1902.10486.
- Deng, D.; Chen, G.; Hao, J.; Wang, Q.; and Heng, P.-A. 2021. Flattening Sharpness for Dynamic Gradient Projection Memory Benefits Continual Learning. *Advances in Neural Information Processing Systems*, 34: 18710–18721.
- Ebrahimi, S.; Elhoseiny, M.; Darrell, T.; and Rohrbach, M. 2020. Uncertainty-guided Continual Learning with Bayesian Neural Networks. In *International Conference on Learning Representations*.
- Farajtabar, M.; Azizan, N.; Mott, A.; and Li, A. 2020. Orthogonal gradient descent for continual learning. In *International Conference on Artificial Intelligence and Statistics*, 3762–3773. PMLR.
- Guo, Y.; Hu, W.; Zhao, D.; and Liu, B. 2022. Adaptive Orthogonal Projection for Batch and Online Continual Learning. *Proceedings of AAAI-2022*, 2.
- Gupta, G.; Yadav, K.; and Paull, L. 2020. Look-ahead meta learning for continual learning. *Advances in Neural Information Processing Systems*, 33: 11588–11598.
- Hadsell, R.; Rao, D.; Rusu, A. A.; and Pascanu, R. 2020. Embracing change: Continual learning in deep neural networks. *Trends in cognitive sciences*, 24(12): 1028–1040.
- Hsu, Y.-C.; Liu, Y.-C.; Ramasamy, A.; and Kira, Z. 2018. Re-evaluating Continual Learning Scenarios: A Categorization and Case for Strong Baselines. In *NeurIPS Continual learning Workshop*.
- Kao, T.-C.; Jensen, K.; van de Ven, G.; Bernacchia, A.; and Hennequin, G. 2021. Natural continual learning: success is a journey, not (just) a destination. *Advances in Neural Information Processing Systems*, 34: 28067–28079.
- Kingma, D. P.; and Ba, J. 2014. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*.
- Kirkpatrick, J.; Pascanu, R.; Rabinowitz, N. C.; Veness, J.; Desjardins, G.; Rusu, A. A.; Milan, K.; Quan, J.; Ramalho, T.; Grabska-Barwinska, A.; Hassabis, D.; Clopath, C.; Kumaran, D.; and Hadsell, R. 2017. Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114: 3521 – 3526.
- Krizhevsky, A. 2009. Learning multiple layers of features from tiny images. Technical report, University of Toronto, Toronto.
- Lange, M. D.; Aljundi, R.; Masana, M.; Parisot, S.; Jia, X.; Leonardis, A.; Slabaugh, G.; and Tuytelaars, T. 2021. A continual learning survey: Defying forgetting in classification tasks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 1–1.
- Lin, S.; Yang, L.; Fan, D.; and Zhang, J. 2022. TRGP: Trust Region Gradient Projection for Continual Learning. In *International Conference on Learning Representations*.
- Lopez-Paz, D.; and Ranzato, M. A. 2017. Gradient Episodic Memory for Continual Learning. In *Advances in Neural Information Processing Systems*, volume 30.
- Mallya, A.; and Lazebnik, S. 2018. PackNet: Adding Multiple Tasks to a Single Network by Iterative Pruning. *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 7765–7773.
- Mccloskey, M.; and Cohen, N. J. 1989. Catastrophic Interference in Connectionist Networks: The Sequential Learning Problem. *The Psychology of Learning and Motivation*, 24: 104–169.
- Mnih, V.; Kavukcuoglu, K.; Silver, D.; Graves, A.; Antonoglou, I.; Wierstra, D.; and Riedmiller, M. 2013. Playing atari with deep reinforcement learning. *arXiv preprint arXiv:1312.5602*.
- Qin, Q.; Hu, W.; Peng, H.; Zhao, D.; and Liu, B. 2021. BNS: Building Network Structures Dynamically for Continual Learning. *Advances in Neural Information Processing Systems*, 34: 20608–20620.
- Ratcliff, R. 1990. Connectionist models of recognition memory: constraints imposed by learning and forgetting functions. *Psychological review*, 97 2: 285–308.
- Rebuffi, S.-A.; Kolesnikov, A.; Sperl, G.; and Lampert, C. H. 2017. iCaRL: Incremental Classifier and Representation Learning. *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 5533–5542.
- Riemer, M.; Cases, I.; Ajemian, R.; Liu, M.; Rish, I.; Tu, Y.; and Tesauro, G. 2019. Learning to Learn without Forgetting By Maximizing Transfer and Minimizing Interference. In *International Conference on Learning Representations*.
- Ring, M. B. 1998. Child: A First Step Towards Continual Learning. In *Learning to Learn*.
- Robins, A. V. 1995. Catastrophic Forgetting, Rehearsal and Pseudorehearsal. *Connect. Sci.*, 7: 123–146.
- Rusu, A. A.; Rabinowitz, N. C.; Desjardins, G.; Soyer, H.; Kirkpatrick, J.; Kavukcuoglu, K.; Pascanu, R.; and Hadsell, R. 2016. Progressive Neural Networks. *ArXiv*, abs/1606.04671.

- Saha, G.; Garg, I.; Ankit, A.; and Roy, K. 2021a. SPACE: Structured Compression and Sharing of Representational Space for Continual Learning. *IEEE Access*, 9: 150480–150494.
- Saha, G.; Garg, I.; and Roy, K. 2021b. Gradient Projection Memory for Continual Learning. In *International Conference on Learning Representations*.
- Saha, G.; and Roy, K. 2023. Saliency Guided Experience Packing for Replay in Continual Learning. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, 5273–5283.
- Schulman, J.; Wolski, F.; Dhariwal, P.; Radford, A.; and Klimov, O. 2017. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*.
- Schwarz, J.; Czarnecki, W.; Luketina, J.; Grabska-Barwinska, A.; Teh, Y. W.; Pascanu, R.; and Hadsell, R. 2018. Progress & Compress: A scalable framework for continual learning. In *ICML*.
- Serrà, J.; Surís, D.; Miron, M.; and Karatzoglou, A. 2018. Overcoming Catastrophic Forgetting with Hard Attention to the Task. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, 4548–4557. PMLR.
- Shi, Y.; Yuan, L.; Chen, Y.; and Feng, J. 2021. Continual learning via bit-level information preserving. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 16674–16683.
- Shin, H.; Lee, J. K.; Kim, J.; and Kim, J. 2017. Continual Learning with Deep Generative Replay. In *Advances in Neural Information Processing Systems*, volume 30.
- Thrun, S.; and Mitchell, T. M. 1995. Lifelong robot learning. *Robotics Auton. Syst.*, 15: 25–46.
- Veniat, T.; Denoyer, L.; and Ranzato, M. 2021. Efficient Continual Learning with Modular Networks and Task-Driven Priors. In *International Conference on Learning Representations*.
- Vinyals, O.; Blundell, C.; Lillicrap, T.; Kavukcuoglu, K.; and Wierstra, D. 2016. Matching Networks for One Shot Learning. In *Advances in Neural Information Processing Systems*, volume 29.
- Wang, S.; Li, X.; Sun, J.; and Xu, Z. 2021. Training networks in null space of feature covariance for continual learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 184–193.
- Xu, J.; and Zhu, Z. 2018. Reinforced Continual Learning. In *Advances in Neural Information Processing Systems*, volume 31, 899–908.
- Yoon, J.; Kim, S.; Yang, E.; and Hwang, S. J. 2020. Scalable and Order-robust Continual Learning with Additive Parameter Decomposition. In *International Conference on Learning Representations*.
- Yoon, J.; Yang, E.; Lee, J.; and Hwang, S. J. 2018. Lifelong Learning with Dynamically Expandable Networks. In *6th International Conference on Learning Representations*.
- Zeng, G.; Chen, Y.; Cui, B.; and Yu, S. 2019. Continual learning of context-dependent processing in neural networks. *Nature Machine Intelligence*, 1(8): 364–372.
- Zenke, F.; Poole, B.; and Ganguli, S. 2017. Continual Learning Through Synaptic Intelligence. In *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, 3987–3995. PMLR.