
Guiding Neural Collapse: Optimising Towards the Nearest Simplex Equiangular Tight Frame

Evan Markou

Australian National University
evan.markou@anu.edu.au

Thalaiyasingam Ajanthan

Australian National University & Amazon
thalaiyasingam.ajanthan@anu.edu.au

Stephen Gould

Australian National University
stephen.gould@anu.edu.au

Abstract

Neural Collapse (NC) is a recently observed phenomenon in neural networks that characterises the solution space of the final classifier layer when trained until zero training loss. Specifically, NC suggests that the final classifier layer converges to a Simplex Equiangular Tight Frame (ETF), which maximally separates the weights corresponding to each class. By duality, the penultimate layer feature means also converge to the same simplex ETF. Since this simple symmetric structure is optimal, our idea is to utilise this property to improve convergence speed. Specifically, we introduce the notion of *nearest simplex ETF geometry* for the penultimate layer features at any given training iteration, by formulating it as a Riemannian optimisation. Then, at each iteration, the classifier weights are implicitly set to the nearest simplex ETF by solving this inner-optimisation, which is encapsulated within a declarative node to allow backpropagation. Our experiments on synthetic and real-world architectures for classification tasks demonstrate that our approach accelerates convergence and enhances training stability¹.

1 Introduction

While modern deep neural networks (DNNs) have demonstrated remarkable success in solving diverse machine learning problems [22, 34, 38], the fundamental mechanisms underlying their training process remain elusive. In recent years, considerable research efforts have focused on delineating the optimisation trajectory and characterising the solution space resulting from the optimisation process in training neural networks [72, 17, 49, 41]. One such finding is that gradient descent algorithms, when combined with certain loss functions, introduce an implicit bias that often favours max-margin solutions, influencing the learned representations and decision boundaries. [44, 57, 33, 27, 21, 54, 70, 31, 49].

In this vein, Neural Collapse (NC) is a recently observed phenomenon in neural networks that characterises the solution space of the final classifier layer in both balanced [50, 76, 74, 28, 48, 63, 43, 32] and imbalanced dataset settings [19, 59, 5]. Specifically, NC suggests that the final classifier layer converges to a Simplex Equiangular Tight Frame (ETF), which maximally separates the weights corresponding to each class, and by duality, the penultimate layer feature means converge to the classifier weights, *i.e.*, to the simplex ETF (formal definitions are provided in Appendix A). This simple, symmetric structure is shown to be the only set of optimal solutions for a variety of loss functions when the features are also assumed to be free parameters, *i.e.*, Unconstrained Feature

¹Code available at <https://github.com/evanmarkou/Guiding-Neural-Collapse.git>.

Models (UFMs) [32, 19, 74, 76, 28, 75]. Nevertheless, even in realistic large-scale deep networks, this phenomenon is observed when trained to convergence, even after attaining zero training error.

Since we can characterise the optimal solution space for the classifier layer, a natural extension is to *leverage the simplex ETF structure of the classifier weights to improve training*. To this end, researchers have tried fixing the classifier weights to a canonical simplex ETF, effectively reducing the number of trainable parameters [76]. However, in practice, this approach does not improve the convergence speed as the backbone network still needs to do the heavy lifting of matching feature means to the chosen fixed simplex ETF.

In this work, we introduce a mechanism for finding the nearest simplex ETF to the features at any given training iteration. Specifically, the nearest simplex ETF is determined by solving a Riemannian optimisation problem. Therefore, our classifier weights are dynamically updated based on the penultimate layer feature means at each iteration, *i.e.*, implicitly defined rather than trained using gradient descent. Additionally, by constructing this inner-optimisation problem as a deep declarative node [23], we allow gradients to propagate through the Riemannian optimisation facilitating end-to-end learning. Our whole framework significantly speeds up convergence to a NC solution compared to the fixed simplex ETF and conventional learnable classifier approaches. We demonstrate the effectiveness of our approach on synthetic UFMs and standard image classification experiments.

Our main contributions are as follows:

1. We introduce the notion of the nearest simplex ETF geometry given the penultimate layer features. Instead of selecting a predetermined simplex ETF (canonical or random), we implicitly fix the classifier as the solution to a Riemannian optimisation problem.
2. To establish end-to-end learning, we encapsulate the Riemannian optimisation problem of determining the nearest simplex ETF geometry within a declarative node. This allows for efficient backpropagation throughout the network.
3. We demonstrate that our method achieves an optimal neural collapse solution more rapidly compared to fixed simplex ETF methods or conventional training approaches, where a learned linear classifier is employed. Additionally, our method ensures training stability by markedly reducing variance in network performance.

2 Related Work

Neural Collapse and Simplex ETFs. Zhu et al. [76] proposed fixing classifier weights to a simplex ETF, reducing parameters while maintaining performance. Simplex ETFs effectively tackle imbalanced learning, as demonstrated by Yang et al. [67], where they fix the target classifier to an arbitrary simplex ETF, relying on the network's over-parameterisation to adapt. Similarly, Yang et al. [68] addressed class incremental learning by fixing the target classifier to a simplex ETF. They advocate adjusting prototype means towards the simplex ETF using a convex combination, smoothly guiding backbone features into the targeted simplex ETF. However, these methods did not yield any benefits regarding convergence speed. The work most relevant to ours is that of Peifeng et al. [51], who argued about the significance of feature directions, particularly in long-tailed learning scenarios. They compared their method against a fixed simplex ETF target, formulating their problem to enable the network to learn feature direction through a rotation matrix. Additionally, they efficiently addressed their optimisation using trivialisation techniques [39, 40]. However, they did not demonstrate any improvements in convergence speed over the fixed simplex ETF, achieving only a minimal increase in test accuracy. Fixing a classifier is not a recent concept, as it has been proposed prior to the emergence of neural collapse [52, 58, 30]. Most notably, Pernici et al. [52] demonstrated improved convergence speed by fixing the classifier to a simplex structure only on ImageNet while maintaining comparable performance on smaller-scale datasets. In contrast, our method shows superior convergence speed compared to both a fixed simplex ETF and a learned classifier across both small and large-scale datasets.

Optimisation on Smooth Manifolds. Our optimisation problem involves orthogonality constraints, characterised by the Stiefel manifold [2, 11]. Due to the nonlinearity of these constraints, efficiently solving such problems requires leveraging Riemannian geometry [18]. A multitude of works are dedicated to solving such problems by either transforming existing classical optimisation techniques into Riemannian equivalent algorithms [1, 73, 20, 64, 55] or by carefully designing penalty functions

to address equivalent unconstrained problems [66, 65]. In our approach, we opt for a retraction-based Riemannian optimisation algorithm [1] to optimally handle orthogonality constraints.

Implicit Differentiable Optimisation. In a neural network setting, end-to-end architectures are commonplace. To backpropagate solutions to optimisation problems, we rely on machinery from implicit differentiation. Pioneering works [4, 3] demonstrated efficient gradient backpropagation when dealing with solutions of convex optimisation problems. This concept was independently introduced as a generalised version by Gould et al. [23, 24] to encompass any twice-differentiable optimisation problem. A key advantage of Deep Declarative Networks (DDNs) lies in their ability to efficiently solve problems at any scale by leveraging the problem’s underlying structure [25]. Our setting involves utilising an equality-constrained declarative node to efficiently backpropagate through the network.

3 Optimising Towards the Nearest Simplex ETF

In this section, we introduce our method to determine the nearest simplex ETF geometry and detail how we can dynamically steer the training algorithm to converge towards this particular solution.

3.1 Problem Setup

Let us first introduce key notation that will be useful when formulating our optimisation problem.

Simplex ETF. Mathematically, a general simplex ETF is a collection of points in $\mathbb{R}^C$ specified by the columns of a matrix

$$\mathbf{M} = \alpha \sqrt{\frac{C}{C-1}} \mathbf{U} \left(\mathbf{I}_C - \frac{1}{C} \mathbf{1}_C \mathbf{1}_C^\top \right). \quad (1)$$

Here, $\alpha \in \mathbb{R}_+$ denotes an arbitrary scale factor, $\mathbf{1}_C$ is the C -dimensional vector of ones, and $\mathbf{U} \in \mathbb{R}^{d \times C}$ (with $d \geq C$) represents a semi-orthogonal matrix ($\mathbf{U}^\top \mathbf{U} = \mathbf{I}_C$). Note that there are many simplex ETFs in $\mathbb{R}^C$ as the rotation $\mathbf{U}$ varies, and $\mathbf{M}$ is rank-deficient. Additionally, the standard simplex ETF with unit Frobenius norm is defined as: $\tilde{\mathbf{M}} = \frac{1}{\sqrt{C-1}} \left(\mathbf{I}_C - \frac{1}{C} \mathbf{1}_C \mathbf{1}_C^\top \right)$.

Mean of Features. Consider a classification dataset $\mathcal{D} = \{(\mathbf{x}_i, y_i) \mid i = 1, \dots, N\}$ where the data $\mathbf{x}_i \in \mathcal{X}$ and labels $y_i \in \mathcal{Y} = \{1, \dots, C\}$. Suppose, n_c is the number of samples correspond to label c , then $\sum_{c=1}^C n_c = N$. Let us consider a scenario where we have a collection of features defined as,

$$\mathbf{H} \triangleq [\mathbf{h}_{c,i} : 1 \leq c \leq C, 1 \leq i \leq n_c] \in \mathbb{R}^{d \times N}. \quad (2)$$

Here, each feature may originate from a nonlinear compound mapping of input data through a neural network, denoted as, $\mathbf{h}_{y_i,i} = \phi_\theta(\mathbf{x}_i)$ for the data sample $(\mathbf{x}_i, y_i)$. Now, for the final layer, our decision variables (weights and biases) are represented as $\mathbf{W} \triangleq [\mathbf{w}_1, \dots, \mathbf{w}_C]^\top \in \mathbb{R}^{C \times d}$, and $\mathbf{b} \in \mathbb{R}^C$, and the logits for the i -th sample is computed as,

$$\psi_\Theta(\mathbf{x}_i) = \mathbf{W} \mathbf{h}_{y_i,i} + \mathbf{b}, \quad \text{where } \mathbf{h}_{y_i,i} = \phi_\theta(\mathbf{x}_i). \quad (3)$$

In UFM, the features are assumed to be free variables, which serves as a rough approximation for neural networks and helps derive theoretical guarantees. Additionally, we define the global mean and per-class mean of the features $\{\mathbf{h}_{c,i}\}$ as:

$$\mathbf{h}_G \triangleq \frac{1}{N} \sum_{c=1}^C \sum_{i=1}^{n_c} \mathbf{h}_{c,i}, \quad \bar{\mathbf{h}}_c \triangleq \frac{1}{n_c} \sum_{i=1}^{n_c} \mathbf{h}_{c,i}, \quad (1 \leq c \leq C), \quad (4)$$

and the globally centred feature mean matrix as,

$$\tilde{\mathbf{H}} \triangleq [\bar{\mathbf{h}}_1 - \mathbf{h}_G, \dots, \bar{\mathbf{h}}_C - \mathbf{h}_G] \in \mathbb{R}^{d \times C}. \quad (5)$$

Finally, we scale the feature mean matrix to have unit Frobenius norm, i.e., $\tilde{\mathbf{H}} = \tilde{\mathbf{H}} / \|\tilde{\mathbf{H}}\|_F$ which will be used in formulation below.

3.2 Nearest Simplex ETF through Riemannian Optimisation

Once we obtain the feature means, our objective is to calculate the nearest simplex ETF based on these means and subsequently adjust the classifier weights $\mathbf{W}$ to align with this particular simplex ETF. The rationale is to identify and establish a simplex ETF that closely corresponds to the feature means at any given iteration. This approach aims to expedite convergence during the training process by providing the algorithm with a starting point that is closer to an optimal solution rather than requiring it to learn a simplex ETF direction or converge towards an arbitrary one.

To find the nearest simplex ETF geometry, we solve the following Riemannian optimisation problem

$$\underset{\mathbf{U} \in St_C^d}{\text{minimize}} \left\| \tilde{\mathbf{H}} - \mathbf{U} \tilde{\mathbf{M}} \right\|_F^2 \quad (6)$$

where $St_C^d = \{\mathbf{X} \in \mathbb{R}^{d \times C} : \mathbf{X}^\top \mathbf{X} = \mathbf{I}_C\}$. Here, $\tilde{\mathbf{M}}$ is the standard simplex ETF with unit Frobenius norm, and the set of the orthogonality constraints St_C^d forms a compact Stiefel manifold [2, 11] embedded in a Euclidean space.

3.3 Proximal Problem

The solution to the Riemannian optimisation problem, denoted as $\mathbf{U}^*$, is not unique since a component of $\mathbf{U}^*$ lies in the null space of $\tilde{\mathbf{M}}$. As simplex ETFs reside in $(C - 1)$ -dimensional space, the matrix $\tilde{\mathbf{M}}$ is rank-one deficient. Consequently, we are faced with a family of solutions, leading to challenges in training stability, as we may oscillate between multiple simplex ETF directions. We address this issue by introducing a proximal term to the problem's objective function. This guarantees the uniqueness of the solution and stabilises the training process, ensuring that our problem converges to a solution closer to the previous one.

So, the original problem in Equation 6 is transformed into:

$$\underset{\mathbf{U} \in St_C^d}{\text{minimize}} \left\| \tilde{\mathbf{H}} - \mathbf{U} \tilde{\mathbf{M}} \right\|_F^2 + \frac{\delta}{2} \left\| \mathbf{U} - \mathbf{U}_{\text{prox}} \right\|_F^2. \quad (7)$$

Here, $\mathbf{U}_{\text{prox}}$ represents the proximal target simplex ETF direction, and $\delta > 0$ serves as the proximal coefficient, handling the trade-off between achieving the optimal solution's proximity to the feature means and its proximity to a given simplex ETF direction. In fact, one can perceive our problem formulation in Equation 7 as a generalisation to a predetermined fixed simplex ETF solution. This is evident when considering that if we significantly increase δ , the optimal direction $\mathbf{U}^*$ would converge towards the fixed proximal direction $\mathbf{U}_{\text{prox}}$.

3.4 General Learning Setting

Our problem formulation, following the general deep neural network architecture in Equation 3, can be seen as a bilevel optimisation problem as follows:

$$\begin{aligned} \underset{\Theta}{\text{minimize}} \quad \mathcal{L}(\mathcal{D}; \Theta, \mathbf{U}^*) &\triangleq -\frac{1}{N} \sum_{i=1}^N \log \left(\frac{\exp(\psi_{\Theta}(\mathbf{x}_i, \mathbf{U}^*)_{y_i})}{\sum_{j=1}^C \exp(\psi_{\Theta}(\mathbf{x}_i, \mathbf{U}^*)_j)} \right), \\ \text{subject to } \mathbf{U}^* &\in \arg \min_{\mathbf{U} \in St_C^d} \left\| \tilde{\mathbf{H}} - \mathbf{U} \tilde{\mathbf{M}} \right\|_F^2 + \frac{\delta}{2} \left\| \mathbf{U} - \mathbf{U}_{\text{prox}} \right\|_F^2, \end{aligned} \quad (8)$$

where $\psi_{\Theta}(\mathbf{x}_i, \mathbf{U}^*) = \tau \mathbf{M} \mathbf{U}^* (\mathbf{h}_i - \mathbf{h}_G)$ with $\mathbf{h}_i = \phi_{\Theta}(\mathbf{x}_i)$. Here, ψ denotes the logits, where the classifier weights are set as $\mathbf{W} = \mathbf{M} \mathbf{U}^*$, and the bias is set to $\mathbf{b} = -\mathbf{M} \mathbf{U}^* \mathbf{h}_G$ to account for feature centring. Furthermore, $\mathbf{M}$ is the standard simplex ETF, $\tilde{\mathbf{M}}$ is its normalised version, and $\tilde{\mathbf{H}}$ is the normalised centred feature matrix. The temperature parameter $\tau > 0$ controls the lower bound of the cross-entropy loss when dealing with normalised features, as defined in [69, Theorem 1].

3.5 Handling Stochastic Updates

In practice, we use stochastic gradient descent updates, and, as such, adjustments to our computations are necessary. With each gradient update now based on a mini-batch, we implement two key

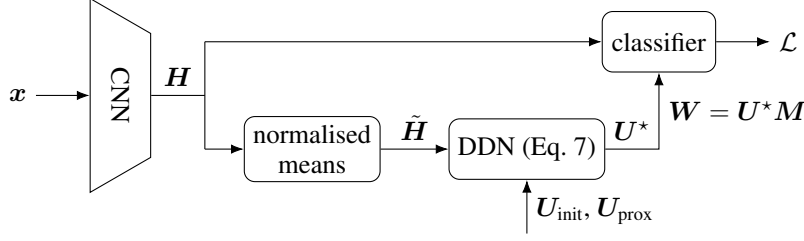


Figure 1: Schematic of our proposed architecture for optimising towards the nearest simplex ETF. The classifier weights $\mathbf{W} = \mathbf{U}^* \mathbf{M}$ are an implicit function of the CNN features $\mathbf{H}$. Note that the parameters of the CNN are updated via two gradient paths from the loss function $\mathcal{L}$, a direct path (top) and an indirect path through $\mathbf{U}^*$ (bottom).

changes. First, rather than directly optimising the problem of finding the nearest simplex ETF geometry concerning the feature means of the mini-batch, we introduce an exponential moving average operation during the computation of the feature means. This operation accumulates statistics and enhances training stability throughout iterations. Formally, at time step t , we have the following equation, where $\alpha \in \mathbb{R}$ represents the smoothing factor:

$$\tilde{\mathbf{H}}_t = \alpha \tilde{\mathbf{H}}_{\text{batch}} + (1 - \alpha) \tilde{\mathbf{H}}_{t-1}. \quad (9)$$

Second, we employ stratified batch sampling to guarantee that all class labels are represented in the mini-batch. This ensures that we avoid degenerate solutions when finding the nearest simplex ETF geometry, as our optimisation problem requires input feature means for all C classes. In cases where the number of classes exceeds the chosen batch size, we compute the per-class feature mean for the class labels present in the given batch. For the remaining class labels, we set their feature mean as the global mean of the batch. We repeat this process for each training iteration until we have sampled examples belonging to the missing class labels. At that point, we update the feature mean of those missing class labels with the new feature statistics. We reserve this method only for cases where the batch size is smaller than the number of labels since it can introduce instability during early iterations.

3.6 Deep Declarative Layer

We can backpropagate through the Riemannian optimisation problem to update the feature means using a declarative node [23]. Then, the features are updated from both the loss and the feature means through auto-differentiation. The motivation for developing the DDN layer lies in recognising that, despite the presence of a proximal term, abrupt and sudden changes to the classifier may occur as the features are updated. These changes can pose challenges for backpropagation, potentially disrupting the stability and convergence of the training process. Incorporating an additional stream of gradients through the feature means to account for such changes, as depicted in Figure 1, assists in stabilising the feature updates during backpropagation.

To efficiently backpropagate through the optimisation problem, we employ techniques described in Gould et al. [23] utilising the implicit function theorem to compute the gradients. In our case, we have a scalar objective function $f : \mathbb{R}^{d \times C} \rightarrow \mathbb{R}$, and a matrix constraint function $J : \mathbb{R}^{d \times C} \rightarrow \mathbb{R}^{C \times C}$. Since we have matrix variables, we use vectorisation techniques [46] to avoid numerically dealing with tensor gradients. More specifically, we have the following:

Proposition 1 (Following directly from Proposition 4.5 in Gould et al. [23]). *Consider the optimisation problem in Equation 7. Assume that the solution exists and that the objective function f and the constraint function J are twice differentiable in the neighbourhood of the solution. If the $\text{rank}(\mathbf{A}) = \frac{C(C+1)}{2}$ and $\mathbf{G}$ is non-singular then:*

$$Dy(\mathbf{U}) = \mathbf{G}^{-1} \mathbf{A}^\top (\mathbf{A} \mathbf{G}^{-1} \mathbf{A}^\top)^{-1} (\mathbf{A} \mathbf{G}^{-1} \mathbf{B}) - \mathbf{G}^{-1} \mathbf{B},$$

where,

$$\mathbf{A} = \text{rvech}(D_{\mathbf{U}} J(\tilde{\mathbf{H}}, \mathbf{U})) \in \mathbb{R}^{\frac{C(C+1)}{2} \times dC},$$

$$\mathbf{B} = \text{rvec}(D_{\tilde{\mathbf{H}}}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) \in \mathbb{R}^{dC \times dC},$$

$$\mathbf{G} = \text{rvec}(D_{\mathbf{U}\mathbf{U}}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) - \mathbf{\Lambda} : \text{rvech}(D_{\mathbf{U}\mathbf{U}}^2 J(\tilde{\mathbf{H}}, \mathbf{U})) \in \mathbb{R}^{dC \times dC}.$$

Here, the double dot product symbol $(:)$ denotes a tensor contraction on appropriate indices between the Lagrange multiplier matrix $\mathbf{\Lambda}$ and a fourth-order tensor Hessian. Also, $\text{rvec}(\cdot)$ and $\text{rvech}(\cdot)$ refer to the row-major vectorisation and half-vectorisation operations, respectively. To find the Lagrange multiplier matrix $\mathbf{\Lambda} \in \mathbb{R}^{C \times \frac{C+1}{2}}$, we solve the following equation where we have vectorised the matrix as $\boldsymbol{\lambda} \in \mathbb{R}^{\frac{C(C+1)}{2}}$,

$$\boldsymbol{\lambda}^\top \mathbf{A} = D_U f(\tilde{\mathbf{H}}, \mathbf{U}).$$

Alternatively, for a more efficient computation of the identity $\mathbf{G}$, we can utilise the embedded gradient field method as defined in Birtea et al. [9, 10]. Therefore, we obtain:

$$\mathbf{G} = \text{rvec}(D_{UU}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) - \mathbf{I}_d \otimes \Sigma(\mathbf{U}) \in \mathbb{R}^{dC \times dC},$$

where $\Sigma(\mathbf{U}) = \frac{1}{2} \left(D_U f(\tilde{\mathbf{H}}, \mathbf{U})^\top \mathbf{U} + \mathbf{U}^\top D_U f(\tilde{\mathbf{H}}, \mathbf{U}) \right)$, and $\otimes$ here denotes Kronecker product.

A detailed derivation of each identity in the proposition can be found in Appendix B.

4 Experiments

In our experiments, we perform feature normalisation onto a hypersphere, a common practice in training neural networks, which improves representation and enhances model performance [71, 26, 53, 42, 61, 62, 12, 16, 35]. We find that combining classifier weight normalisation with feature normalisation accelerates convergence [69]. Given that simplex ETFs are inherently normalised, we include classifier weight normalisation in our standard training procedure to ensure fair method comparisons.

Experimental Setup. In this study, we conduct experiments on three model variants. First, the standard method involves training a model with learnable classifier weights, following conventional practice. Second, in the fixed ETF method, we set the classifier to a predefined simplex ETF. In all experiments, we choose the simplex ETF with canonical direction. In Appendix C, we also include additional experiments for fixed simplex ETFs with random directions generated from a Haar measure [47]. Last, our implicit ETF method, where we set the classifier weights on-the-fly as the simplex ETF closest to the current feature means.

We repeat experiments on each method five times with distinct random seeds and report the median values alongside their respective ranges. For reproducibility and to streamline hyperparameter tuning, we employed Automatic Gradient Descent (AGD) [6]. Following the authors' recommendation, we set the gain/momentum parameter to 10 to expedite convergence, aligning it with other widely used optimisers like Adam [36] and SGD. Our experiments on real datasets run for 200 epochs with batch size 256; for the UFM analysis, we run 2000 iterations.

Our method underwent rigorous evaluation across various UFM sizes and real model architectures trained on actual datasets, including CIFAR10 [37], CIFAR100 [37], STL10 [14], and ImageNet-1000 [15], implemented on ResNet [29] and VGG [56] architectures. More specifically, we trained CIFAR10 on ResNet18 and VGG13, CIFAR100 and STL10 on ResNet50 and VGG13, and ImageNet-1000 on ResNet50. The input images were preprocessed pixel-wise by subtracting the mean and dividing by the standard deviation. Additionally, standard data augmentation techniques were applied, including random horizontal flips, rotations, and crops. All experiments were conducted using Nvidia RTX3090 and A100 GPUs.

Hyperparameter Selection and Riemannian Initialisation Schemes. We solve the Riemannian optimisation problem defined in Equation 7 using a Riemannian Trust-Region method [1] from pyManopt [60]. We maintain a proximal coefficient δ set to 10^{-3} consistently across all experiments. It is worth mentioning that algorithm convergence is robust to the precise value of δ . In our problem, determining values for $\mathbf{U}_{\text{init}}$ and $\mathbf{U}_{\text{prox}}$ is crucial. We explored several methods to initialise these parameters. One approach involved setting both towards the canonical simplex ETF direction. This means initialising them as a partial orthogonal matrix where the first C rows and columns form an identity matrix while the remaining $d - C$ rows are filled with zeros. Another approach is to initialise both of them as random orthogonal matrices from classical compact groups, selected according to a Haar measure [47]. In the end, the approach that yielded the most stable results at initialisation was

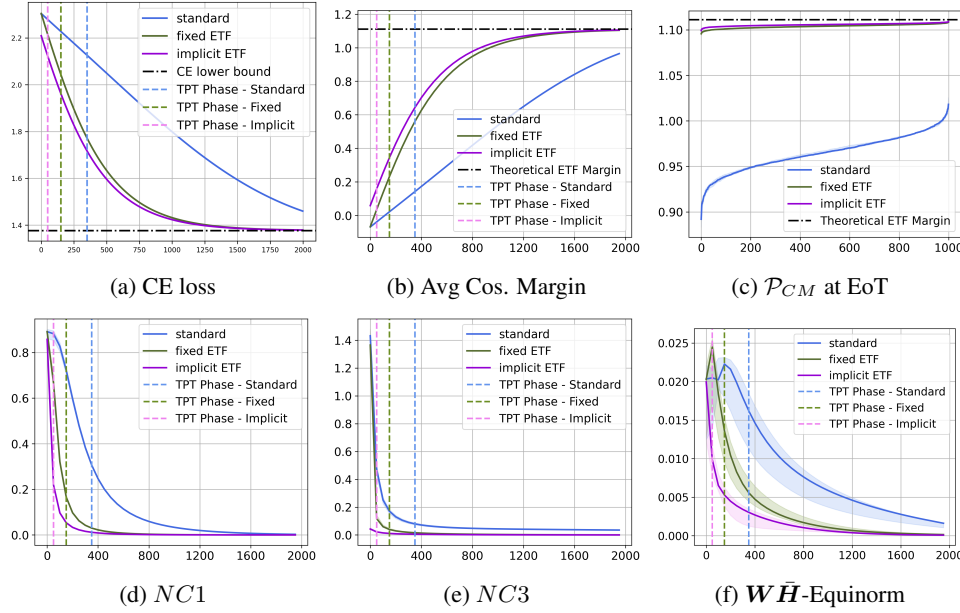


Figure 2: UFM-10 results. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

to employ either of the aforementioned initialisation methods to solve the original problem without the proximal term in Equation 6. We then used the obtained U^* to initialise both U_{init} and U_{prox} for the problem in Equation 7. This process was carried out only for the first gradient update of the first epoch. In subsequent iterations, we update these parameters to the U^* obtained from the previous time step. Importantly, the proximal term is held fixed during each Riemannian optimisation.

Regarding the calculation of the exponential moving average of the feature means, we have found that employing a decay policy on the smoothing factor α yields optimal results. Specifically, we set $\alpha = 2/(T + 1)$, where T represents the number of iterations. Additionally, we include a thresholding value of 10^{-4} , such that if α falls below this threshold, we fix α to be equal to the threshold. This precaution ensures that α does not diminish throughout the iterations, thereby guaranteeing that the newly calculated feature means contribute sufficient statistics to the exponential moving average.

Finally, in our experiments, we set the temperature parameter τ to five. This choice aligns with the findings discussed by Yaras et al. [69], highlighting the influence of the temperature parameter value on the extent of neural collapse statistics with normalised features.

Unconstrained Feature Models (UFMs). Our experiments on UFMs, which provide a controlled setting for evaluating the effectiveness of our method, are done using the following configurations:

- UFM-10: a 10-class UFM containing 1000 features with a dimension of 512.
- UFM-100: a 100-class UFM containing 5000 features with a dimension of 1024.
- UFM-200: a 200-class UFM containing 5000 features, with a dimension of 1024.
- UFM-1000: a 1000-class UFM containing 10000 features, with a dimension of 1024.

Results. We present the results for the synthetic UFM-10 case in Figure 2. The CE loss plot demonstrates that fixing the classifier weights to a simplex ETF achieves the theoretical lower bound of Yaras et al. [69, Thm. 1], indicating the attainment of a globally optimal solution. We also visualise the average cosine margin per epoch and the cosine margin distributions of each example at the end of training, defined in Zhou et al. [74]. The neural collapse metrics, $NC1$ and $NC3$, which measure the features' within-class variability, and the self-duality alignment between the feature means and the classifier weights [76], are also plotted. Last, we depict the absolute difference of the classifier and feature means norms to illustrate their convergence towards equinorms, as described in Pappan et al. [50]. A comprehensive description of the metrics can be found in Appendix A. Collectively, the plots indicate the superior performance of our method in achieving a neural collapse (NC) solution

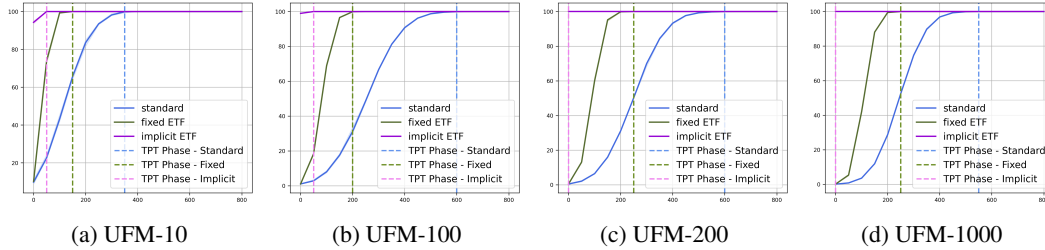


Figure 3: The evolution of convergence measured in top-1 accuracy of the UFM as we increase the number of classes, plotted for the first 800 epochs. We omit the rest of the epochs as all methods have converged and have identical results.

Table 1: Train top-1 accuracy results presented as a median with indices indicating the range of values from five random seeds. Best results are bolded.

Dataset	Network	Train accuracy at epoch 50			Train accuracy at epoch 200		
		Standard	Fixed ETF	Implicit ETF	Standard	Fixed ETF	Implicit ETF
CIFAR10	ResNet18	87.42 ^{89.7} _{86.1}	86.89 ^{88.6} _{84.7}	88.71 ^{89.6} _{88.5}	96.69 ^{98.6} _{96.5}	97.18 ^{97.9} _{95.6}	98.09 ^{98.6} _{97.9}
	VGG13	93.59 ^{97.0} _{90.7}	76.66 ^{85.8} _{53.9}	95.69 ^{96.2} _{95.2}	99.15 ^{99.8} _{97.9}	79.36 ^{99.1} _{59.6}	99.56 ^{99.7} _{99.4}
CIFAR100	ResNet50	58.47 ^{59.6} _{53.9}	63.93 ^{65.2} _{61.1}	72.15 ^{74.1} _{70.1}	95.87 ^{98.6} _{94.7}	91.34 ^{92.1} _{90.4}	96.96 ^{97.3} _{96.2}
	VGG13	82.00 ^{84.0} _{80.5}	81.14 ^{81.9} _{76.0}	88.39 ^{89.4} _{86.9}	99.34 ^{99.6} _{99.2}	94.55 ^{95.3} _{92.5}	98.92 ^{99.0} _{98.8}
STL10	ResNet50	83.86 ^{90.7} _{84.5}	86.76 ^{86.8} _{77.8}	93.54 ^{95.3} _{91.3}	99.42 ^{99.9} _{99.0}	98.38 ^{99.3} _{98.1}	99.72 ^{99.9} _{98.0}
	VGG13	82.66 ^{90.7} _{73.7}	83.60 ^{85.1} _{65.6}	90.14 ^{93.5} _{69.2}	100.0 ¹⁰⁰ ₁₀₀	99.92 ^{99.9} _{99.8}	100.0 ¹⁰⁰ ₁₀₀
ImageNet	ResNet50	58.35 ^{59.1} _{58.0}	70.44 ^{70.7} _{69.5}	74.09 ^{74.5} _{73.8}	77.20 ^{77.3} _{76.5}	83.09 ^{83.6} _{83.1}	88.01 ^{88.5} _{87.5}

faster than other approaches. In Figure 3, we demonstrate under the UFM setting that as we increase the number of classes, our method maintains constant performance and converges at the same rate, while the fixed ETF and the standard approach require more time to reach the interpolation threshold.

Numerical results for the top-1 train and test accuracy are reported in Tables 1 and 2, respectively. The results are provided for snapshots taken at epoch 50 and epoch 200. It is evident that our method achieves a faster convergence speed compared to the competitive methods while ultimately converging to the same performance level. Additionally, it is noteworthy that our method exhibits the smallest degree of variability across different runs, as indicated by the range values provided. Finally, in Figure 4, we present qualitative results that confirm our solution’s ability to converge much faster and reach peak performance earlier than the standard and fixed ETF methods on ImageNet. It’s important to note that the standard method with AGD is reported to converge to the same testing accuracy (65.5%) at epoch 350, as shown in Bernstein et al. [6, Figure 4]. At epoch 200, the authors exhibit a testing accuracy of approximately 51%. Since we have increased the gain parameter on AGD compared to the results reported in the original paper, we report a final 60.67% testing accuracy for the standard method, whereas our method reaches peak convergence at approximately epoch 80. We note that the ImageNet results reported in Tables 1 and 2, as well as Figure 4, are generated solely by solving the Riemannian optimisation problem without considering its gradient stream on the feature updates, due to computational constraints. We discuss the computational requirements of our method in Section 5. We also present qualitative results for all the other datasets and architectures in Appendix C.

5 Discussion: Limitations and Future Directions

Our method involves two gradient streams updating the features, as depicted in Figure 1. Interestingly, empirical observations on small-scale datasets (see Figure 15) indicate that even without the back-propagation through the DDN layer, the performance remains comparable, rendering the gradient calculation of the DDN layer optional. In Figure 15c, we observe a strong impact of the DDN layer gradient on the atomic feature level, with more features reaching the theoretical simplex ETF margin by the end of training. To reach a consensus on the exact effect of the DDN gradient on the learning

Table 2: Test top-1 accuracy results presented as a median with indices indicating the range of values from five random seeds. Best results are bolded.

Dataset	Network	Test accuracy at epoch 50			Test accuracy at epoch 200		
		Standard	Fixed ETF	Implicit ETF	Standard	Fixed ETF	Implicit ETF
CIFAR10	ResNet18	80.47 ^{82.6} _{79.4}	80.63 ^{81.8} _{79.4}	81.76 ^{82.0} _{81.4}	83.97 ^{84.8} _{83.2}	84.53 ^{84.9} _{83.7}	84.78 ^{85.1} _{84.3}
	VGG13	86.70 ^{89.4} _{83.7}	70.99 ^{80.7} _{50.7}	88.30 ^{88.7} _{87.4}	90.34 ^{91.5} _{89.1}	72.48 ^{90.5} _{56.9}	90.98 ^{91.5} _{90.6}
CIFAR100	ResNet50	45.91 ^{46.3} _{42.1}	45.37 ^{45.6} _{43.2}	48.38 ^{49.3} _{48.0}	51.29 ^{51.9} _{50.4}	48.03 ^{49.3} _{47.7}	50.52 ^{51.2} _{50.2}
	VGG13	60.82 ^{61.2} _{59.3}	60.10 ^{61.1} _{57.1}	62.39 ^{63.6} _{61.8}	67.54 ^{68.1} _{66.4}	62.78 ^{63.4} _{61.6}	67.14 ^{67.6} _{66.8}
STL10	ResNet50	55.41 ^{58.3} _{54.8}	58.11 ^{60.0} _{57.1}	57.56 ^{59.2} _{56.4}	63.60 ^{65.2} _{61.8}	64.33 ^{66.5} _{63.9}	62.75 ^{63.0} _{60.7}
	VGG13	66.09 ^{72.3} _{60.0}	66.15 ^{67.7} _{52.2}	68.78 ^{71.3} _{56.2}	81.61 ^{81.9} _{81.3}	79.53 ^{80.3} _{78.6}	79.94 ^{80.9} _{79.5}
ImageNet	ResNet50	52.64 ^{53.3} _{52.3}	58.85 ^{59.1} _{58.3}	63.40 ^{63.8} _{63.2}	60.02 ^{60.7} _{59.8}	61.47 ^{62.0} _{61.4}	65.36 ^{65.7} _{65.0}

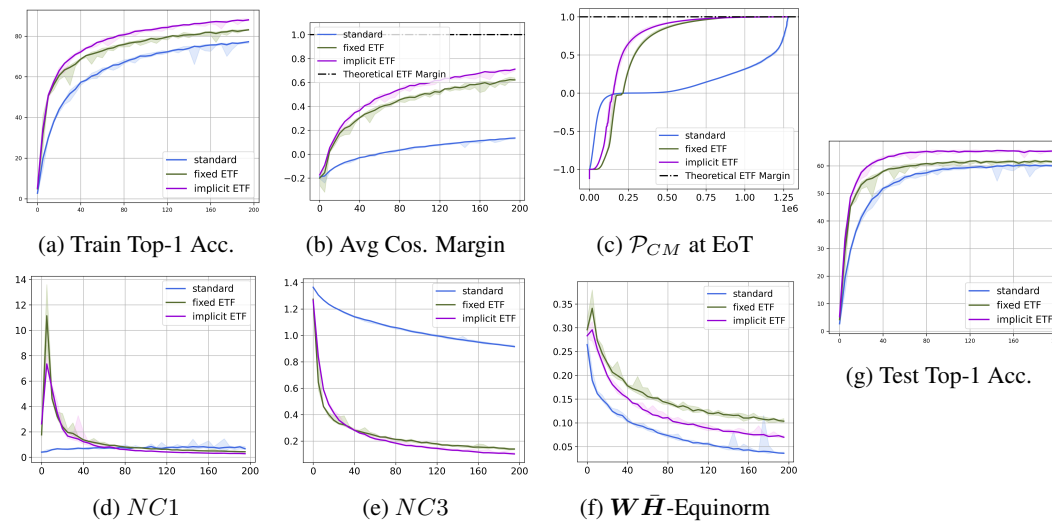


Figure 4: ImageNet results on ResNet-50. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

process, further experiments on large-scale datasets are needed. However, on large-scale datasets with large d and C , such as ImageNet, computing the backward pass of the Riemannian optimisation is challenging due to the memory inefficiency of the current implementation of DDN gradients. This limitation is an area we aim to address in future work. Note that in all other experiments, we use the full gradient computations, including both direct and indirect components, through the DDN layer. We summarise the GPU memory requirements for each method across various datasets in Table 3.

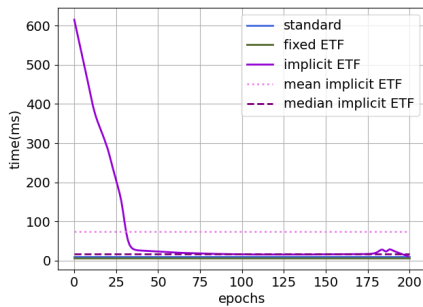
Our discussion so far has focused on convergence speed in terms of the number of epochs required for the network to converge. However, it is also important to consider the time required per epoch. In our case, as training progresses, the time taken by the Riemannian optimization quickly becomes almost negligible compared to the network’s total forward pass time, while it approaches the standard and fixed ETF training forward times, as shown in Figure 5a. However, DDN gradient computation increases considerably when the feature dimension d and the number of classes C increase and starts to dominate the runtime for large datasets such as ImageNet. Nevertheless, for ImageNet, we do not compute the DDN gradients and still outperform other methods. We plan to explore ways to expedite the DDN forward and backward pass in future work.

6 Conclusion

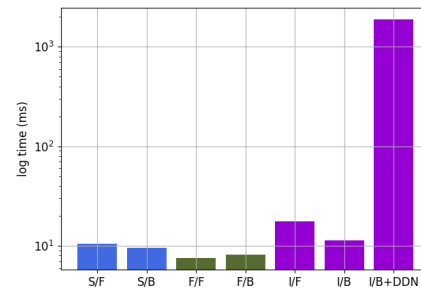
In this paper, we introduced a novel method for determining the nearest simplex ETF to the penultimate features of a neural network and utilising it as our target classifier at each iteration. This contrasts with previous approaches, which either fix to a specific simplex ETF or allow the network

Table 3: GPU memory (in Gigabytes) during training.

Model	Standard	Fixed ETF	Implicit ETF	
			w/o DDN Bwd	w/ DDN Bwd
UFM-10	1.5	1.5	1.5	1.6
UFM-100	1.7	1.7	1.7	10.7
UFM-200	1.7	1.7	1.8	70.3
UFM-1000	1.9	1.9	2.9	N/A
ResNet18 - CIFAR10	2.2	2.2	2.2	2.3
ResNet50 - CIFAR10	2.6	2.6	2.7	2.8
ResNet50 - CIFAR100	2.6	2.6	2.7	18.9
ResNet50 - ImageNet	27.5	27.2	27.8	N/A



(a) Forward pass times in milliseconds.



(b) Forward and backward times in (log) milliseconds.

Figure 5: CIFAR100 computational cost results on ResNet-50. In (a), we plot the forward pass time for each method. For the implicit ETF method, which has dynamic computation times, we also include the mean and median time values. In (b), we plot the computational cost for each forward and backward pass across methods. For the implicit ETF forward pass, we have taken its median time. The notation is as follows: S/F = Standard Forward Pass, S/B = Standard Backward Pass, F/F = Fixed ETF Forward Pass, F/B = Fixed ETF Backward Pass, I/F = Implicit ETF Forward Pass, and I/B = Implicit ETF Backward Pass.

to learn it through gradient descent. Our method involves solving a Riemannian optimisation problem facilitated by a deep declarative node, enabling backpropagation through this process.

We demonstrated that our approach enhances convergence speed across various datasets and architectures while also reducing variability stemming from different random initialisations. By defining the optimal structure of the classifier and efficiently leveraging its rotation invariance property to find the one closest to the backbone features, we anticipate that our method will facilitate the creation of new architectures and the utilisation of new datasets without necessitating specific learning or tuning of the classifier’s structure.

References

- [1] P.-A. Absil, C. G. Baker, and K. A. Gallivan. Trust-region methods on Riemannian manifolds. *Foundations of Computational Mathematics*, 7(3):303–330, 2007. doi: 10.1007/s10208-005-0179-9.
- [2] P.-A. Absil, R. Mahony, and R. Sepulchre. *Optimization Algorithms on Matrix Manifolds*. Princeton University Press, 2008. ISBN 9780691132983. URL <http://www.jstor.org/stable/j.ctt7smmk>.
- [3] Akshay Agrawal, Brandon Amos, Shane Barratt, Stephen Boyd, Steven Diamond, and J Zico Kolter. Differentiable convex optimization layers. *Advances in neural information processing systems*, 32, 2019.

- [4] Brandon Amos and J Zico Kolter. Optnet: Differentiable optimization as a layer in neural networks. In *International Conference on Machine Learning*, pp. 136–145. PMLR, 2017.
- [5] Tina Behnia, Ganesh Ramachandra Kini, Vala Vakilian, and Christos Thrampoulidis. On the implicit geometry of cross-entropy parameterizations for label-imbalanced data. In *International Conference on Artificial Intelligence and Statistics*, pp. 10815–10838. PMLR, 2023.
- [6] Jeremy Bernstein, Chris Mingard, Kevin Huang, Navid Azizan, and Yisong Yue. Automatic Gradient Descent: Deep Learning without Hyperparameters. *arXiv:2304.05187*, 2023.
- [7] Petre Birtea and Dan Comănescu. Geometrical dissipation for dynamical systems. *Communications in Mathematical Physics*, 316:375–394, 2012.
- [8] Petre Birtea and Dan Comănescu. Hessian operators on constraint manifolds. *Journal of Nonlinear Science*, 25:1285–1305, 2015.
- [9] Petre Birtea, Ioan Cașu, and Dan Comănescu. First order optimality conditions and steepest descent algorithm on orthogonal stiefel manifolds. *Optim. Lett.*, 13(8):1773–1791, November 2019.
- [10] Petre Birtea, Ioan Cașu, and Dan Comănescu. Second order optimality on orthogonal stiefel manifolds. *Bulletin des Sciences Mathématiques*, 161:102868, 2020. ISSN 0007-4497. doi: <https://doi.org/10.1016/j.bulsci.2020.102868>. URL <https://www.sciencedirect.com/science/article/pii/S0007449720300385>.
- [11] Nicolas Boumal. *An introduction to optimization on smooth manifolds*. Cambridge University Press, 2023.
- [12] Kwan Ho Ryan Chan, Yaodong Yu, Chong You, Haozhi Qi, John Wright, and Yi Ma. Deep networks from the principle of rate reduction. *arXiv preprint arXiv:2010.14765*, 2020.
- [13] Lingling He Changqing Xu and Zerong Lin. Commutation matrices and commutation tensors. *Linear and Multilinear Algebra*, 68(9):1721–1742, 2020. doi: 10.1080/03081087.2018.1556242. URL <https://doi.org/10.1080/03081087.2018.1556242>.
- [14] Adam Coates, Andrew Ng, and Honglak Lee. An analysis of single-layer networks in unsupervised feature learning. In Geoffrey Gordon, David Dunson, and Miroslav Dudík (eds.), *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, volume 15 of *Proceedings of Machine Learning Research*, pp. 215–223, Fort Lauderdale, FL, USA, 11–13 Apr 2011. PMLR. URL <https://proceedings.mlr.press/v15/coates11a.html>.
- [15] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *Computer Vision and Pattern Recognition, 2009. CVPR 2009. IEEE Conference on*, pp. 248–255. IEEE, 2009. URL <https://ieeexplore.ieee.org/abstract/document/5206848/>.
- [16] Jiankang Deng, Jia Guo, Niannan Xue, and Stefanos Zafeiriou. Arcface: Additive angular margin loss for deep face recognition. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pp. 4690–4699, 2019.
- [17] Simon Du, Jason Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. Gradient descent finds global minima of deep neural networks. In *International conference on machine learning*, pp. 1675–1685. PMLR, 2019.
- [18] Alan Edelman, Tomás A Arias, and Steven T Smith. The geometry of algorithms with orthogonality constraints. *SIAM journal on Matrix Analysis and Applications*, 20(2):303–353, 1998.
- [19] Cong Fang, Hangfeng He, Qi Long, and Weijie J. Su. Exploring deep neural networks via layer-peeled model: Minority collapse in imbalanced training. *Proceedings of the National Academy of Sciences*, 118(43):e2103091118, 2021. doi: 10.1073/pnas.2103091118. URL <https://www.pnas.org/doi/abs/10.1073/pnas.2103091118>.

- [20] Bin Gao, Xin Liu, Xiaojun Chen, and Ya-xiang Yuan. A new first-order algorithmic framework for optimization problems with orthogonality constraints. *SIAM Journal on Optimization*, 28(1):302–332, 2018.
- [21] Gauthier Gidel, Francis Bach, and Simon Lacoste-Julien. Implicit regularization of discrete gradient dynamics in linear neural networks. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL https://proceedings.neurips.cc/paper_files/paper/2019/file/f39ae9ff3a81f499230c4126e01f421b-Paper.pdf.
- [22] Ian J. Goodfellow, Yoshua Bengio, and Aaron Courville. *Deep Learning*. MIT Press, Cambridge, MA, USA, 2016. <http://www.deeplearningbook.org>.
- [23] S. Gould, R. Hartley, and D. Campbell. Deep declarative networks. *IEEE Transactions on Pattern Analysis & Machine Intelligence*, 44(08):3988–4004, aug 2022. ISSN 1939-3539. doi: 10.1109/TPAMI.2021.3059462.
- [24] Stephen Gould, Basura Fernando, Anoop Cherian, Peter Anderson, Rodrigo Santa Cruz, and Edison Guo. On differentiating parameterized argmin and argmax problems with application to bi-level optimization. *arXiv preprint arXiv:1607.05447*, 2016.
- [25] Stephen Gould, Dylan Campbell, Itzik Ben-Shabat, Chamin Hewa Koneputugodage, and Zhiwei Xu. Exploiting problem structure in deep declarative networks: Two case studies. *arXiv preprint arXiv:2202.12404*, 2022.
- [26] Florian Graf, Christoph Hofer, Marc Niethammer, and Roland Kwitt. Dissecting supervised contrastive learning. In *International Conference on Machine Learning*, pp. 3821–3830. PMLR, 2021.
- [27] Suriya Gunasekar, Jason D Lee, Daniel Soudry, and Nati Srebro. Implicit bias of gradient descent on linear convolutional networks. *Advances in neural information processing systems*, 31, 2018.
- [28] X.Y. Han, Vardan Papayan, and David L. Donoho. Neural collapse under MSE loss: Proximity to and dynamics on the central path. In *International Conference on Learning Representations*, 2022. URL https://openreview.net/forum?id=w1UbdvWH_R3.
- [29] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep Residual Learning for Image Recognition. In *Proceedings of 2016 IEEE Conference on Computer Vision and Pattern Recognition*, CVPR '16, pp. 770–778. IEEE, June 2016. doi: 10.1109/CVPR.2016.90. URL <http://ieeexplore.ieee.org/document/7780459>.
- [30] Elad Hoffer, Itay Hubara, and Daniel Soudry. Fix your classifier: the marginal value of training the last weight layer. *arXiv preprint arXiv:1801.04540*, 2018.
- [31] Meena Jagadeesan, Ilya Razenshteyn, and Suriya Gunasekar. Inductive bias of multi-channel linear convolutional networks with bounded weight norm. In *Conference on Learning Theory*, pp. 2276–2325. PMLR, 2022.
- [32] Wenlong Ji, Yiping Lu, Yiliang Zhang, Zhun Deng, and Weijie J Su. An unconstrained layer-peeled perspective on neural collapse. *arXiv preprint arXiv:2110.02796*, 2021.
- [33] Ziwei Ji, Miroslav Dudík, Robert E. Schapire, and Matus Telgarsky. Gradient descent follows the regularization path for general losses. In Jacob Abernethy and Shivani Agarwal (eds.), *Proceedings of Thirty Third Conference on Learning Theory*, volume 125 of *Proceedings of Machine Learning Research*, pp. 2109–2136. PMLR, 09–12 Jul 2020. URL <https://proceedings.mlr.press/v125/ji20a.html>.
- [34] John M. Jumper, Richard Evans, Alexander Pritzel, Tim Green, Michael Figurnov, Olaf Ronneberger, Kathryn Tunyasuvunakool, Russ Bates, Augustin Žídek, Anna Potapenko, Alex Bridgland, Clemens Meyer, Simon A A Kohl, Andy Ballard, Andrew Cowie, Bernardino Romera-Paredes, Stanislav Nikolov, Rishub Jain, Jonas Adler, Trevor Back, Stig Petersen, David Reiman, Ellen Clancy, Michal Zielinski, Martin Steinegger, Michalina Pacholska, Tamas

- Berghammer, Sebastian Bodenstein, David Silver, Oriol Vinyals, Andrew W. Senior, Koray Kavukcuoglu, Pushmeet Kohli, and Demis Hassabis. Highly accurate protein structure prediction with alphafold. *Nature*, 596:583 – 589, 2021. URL <https://api.semanticscholar.org/CorpusID:235959867>.
- [35] Prannay Khosla, Piotr Teterwak, Chen Wang, Aaron Sarna, Yonglong Tian, Phillip Isola, Aaron Maschinot, Ce Liu, and Dilip Krishnan. Supervised contrastive learning. *Advances in neural information processing systems*, 33:18661–18673, 2020.
 - [36] Diederik Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations (ICLR)*, San Diego, CA, USA, 2015.
 - [37] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
 - [38] Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *nature*, 521(7553):436–444, 2015.
 - [39] Mario Lezcano Casado. Trivializations for gradient-based optimization on manifolds. *Advances in Neural Information Processing Systems*, 32, 2019.
 - [40] Mario Lezcano-Casado and David Martinez-Rubio. Cheap orthogonal constraints in neural networks: A simple parametrization of the orthogonal and unitary group. In *International Conference on Machine Learning*, pp. 3794–3803. PMLR, 2019.
 - [41] Hao Li, Zheng Xu, Gavin Taylor, Christoph Studer, and Tom Goldstein. Visualizing the loss landscape of neural nets. *Advances in neural information processing systems*, 31, 2018.
 - [42] Weiyang Liu, Yandong Wen, Zhiding Yu, Ming Li, Bhiksha Raj, and Le Song. Sphereface: Deep hypersphere embedding for face recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 212–220, 2017.
 - [43] Jianfeng Lu and Stefan Steinerberger. Neural collapse under cross-entropy loss. *Applied and Computational Harmonic Analysis*, 59:224–241, 2022.
 - [44] Kaifeng Lyu and Jian Li. Gradient descent maximizes the margin of homogeneous neural networks. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=SJeLIgBKPS>.
 - [45] Jan R. Magnus and H. Neudecker. The elimination matrix: Some lemmas and applications. *SIAM Journal on Algebraic Discrete Methods*, 1(4):422–449, 1980. doi: 10.1137/0601049. URL <https://doi.org/10.1137/0601049>.
 - [46] Jan R. Magnus and Heinz Neudecker. *Matrix differential calculus with applications in statistics and econometrics / Jan Rudolph Magnus and Heinz Neudecker*. Wiley Series in Probability and Statistics. Wiley, Hoboken, NJ, third edition. edition, 2019. ISBN 1-119-54121-2.
 - [47] Francesco Mezzadri. How to generate random matrices from the classical compact groups. *arXiv preprint math-ph/0609050*, 2006.
 - [48] Dustin G. Mixon, Hans Parshall, and Jianzong Pi. Neural collapse with unconstrained features. *CoRR*, abs/2011.11619, 2020. URL <https://arxiv.org/abs/2011.11619>.
 - [49] Behnam Neyshabur, Ryota Tomioka, and Nathan Srebro. In search of the real inductive bias: On the role of implicit regularization in deep learning. *arXiv preprint arXiv:1412.6614*, 2014.
 - [50] Vardan Papyan, X. Y. Han, and David L. Donoho. Prevalence of neural collapse during the terminal phase of deep learning training. *Proceedings of the National Academy of Science*, 117(40):24652–24663, October 2020. doi: 10.1073/pnas.2015509117.
 - [51] Gao Peifeng, Qianqian Xu, Peisong Wen, Zhiyong Yang, Huiyang Shao, and Qingming Huang. Feature directions matter: Long-tailed learning via rotated balanced representation. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett (eds.), *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pp. 27542–27563. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/peifeng23a.html>.

- [52] Federico Pernici, Matteo Bruni, Claudio Baecchi, and Alberto Del Bimbo. Regular polytope networks. *IEEE Transactions on Neural Networks and Learning Systems*, 33(9):4373–4387, 2021.
- [53] Rajeev Ranjan, Carlos D Castillo, and Rama Chellappa. L2-constrained softmax loss for discriminative face verification. *arXiv preprint arXiv:1703.09507*, 2017.
- [54] Ohad Shamir. Gradient methods never overfit on separable data. *Journal of Machine Learning Research*, 22(85):1–20, 2021.
- [55] Jonathan W. Siegel. Accelerated optimization with orthogonality constraints. *Journal of Computational Mathematics*, 39(2):207–226, 2020. ISSN 1991-7139. doi: <https://doi.org/10.4208/jcm.1911-m2018-0242>. URL http://global-sci.org/intro/article_detail/jcm/18372.html.
- [56] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *CoRR*, abs/1409.1556, 2014. URL <http://arxiv.org/abs/1409.1556>.
- [57] Daniel Soudry, Elad Hoffer, and Nathan Srebro. The implicit bias of gradient descent on separable data. In *International Conference on Learning Representations*, 2018. URL <https://openreview.net/forum?id=r1q7n9gAb>.
- [58] Korawat Tanwisuth, Xinjie Fan, Huangjie Zheng, Shujian Zhang, Hao Zhang, Bo Chen, and Mingyuan Zhou. A prototype-oriented framework for unsupervised domain adaptation. *Advances in Neural Information Processing Systems*, 34:17194–17208, 2021.
- [59] Christos Thrampoulidis, Ganesh Ramachandra Kini, Vala Vakilian, and Tina Behnia. Imbalance trouble: Revisiting neural-collapse geometry. *Advances in Neural Information Processing Systems*, 35:27225–27238, 2022.
- [60] J. Townsend, N. Koep, and S. Weichwald. PyManopt: a Python toolbox for optimization on manifolds using automatic differentiation. *Journal of Machine Learning Research*, 17(137):1–5, 2016. URL <https://www.pymanopt.org>.
- [61] Hao Wang, Yitong Wang, Zheng Zhou, Xing Ji, Dihong Gong, Jingchao Zhou, Zhifeng Li, and Wei Liu. Cosface: Large margin cosine loss for deep face recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 5265–5274, 2018.
- [62] Tongzhou Wang and Phillip Isola. Understanding contrastive representation learning through alignment and uniformity on the hypersphere. In *International conference on machine learning*, pp. 9929–9939. PMLR, 2020.
- [63] E Weinan and Stephan Wojtowytsch. On the emergence of simplex symmetry in the final and penultimate layers of neural network classifiers. In *Mathematical and Scientific Machine Learning*, pp. 270–290. PMLR, 2022.
- [64] Zaiwen Wen and Wotao Yin. A feasible method for optimization with orthogonality constraints. *Mathematical Programming*, 142(1):397–434, 2013.
- [65] Nachuan Xiao and Xin Liu. Solving optimization problems over the stiefel manifold by smooth exact penalty function. *arXiv preprint arXiv:2110.08986*, 2021.
- [66] Nachuan Xiao, Xin Liu, and Kim-Chuan Toh. Dissolving constraints for riemannian optimization. *Mathematics of Operations Research*, 49(1):366–397, 2024.
- [67] Yibo Yang, Shixiang Chen, Xiangtai Li, Liang Xie, Zhouchen Lin, and Dacheng Tao. Inducing neural collapse in imbalanced learning: Do we really need a learnable classifier at the end of deep neural network? In Alice H. Oh, Alekh Agarwal, Danielle Belgrave, and Kyunghyun Cho (eds.), *Advances in Neural Information Processing Systems*, 2022. URL https://openreview.net/forum?id=A6EmxI3_Xc.
- [68] Yibo Yang, Haobo Yuan, Xiangtai Li, Jianlong Wu, Lefei Zhang, Zhouchen Lin, Philip H.S. Torr, Bernard Ghanem, and Dacheng Tao. Neural collapse terminus: A unified solution for class incremental learning and its variants. *arXiv pre-print*, 2023.

- [69] Can Yaras, Peng Wang, Zhihui Zhu, Laura Balzano, and Qing Qu. Neural collapse with normalized features: A geometric analysis over the riemannian manifold. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh (eds.), *Advances in Neural Information Processing Systems*, volume 35, pp. 11547–11560. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/4b3cc0d1c897ebcf71aca92a4a26ac83-Paper-Conference.pdf.
- [70] Chong You, Zhihui Zhu, Qing Qu, and Yi Ma. Robust recovery via implicit bias of discrepant learning rates for double over-parameterization. *Advances in Neural Information Processing Systems*, 33:17733–17744, 2020.
- [71] Yaodong Yu, Kwan Ho Ryan Chan, Chong You, Chaobing Song, and Yi Ma. Learning diverse and discriminative representations via the principle of maximal coding rate reduction. *Advances in Neural Information Processing Systems*, 33:9422–9434, 2020.
- [72] Chiyuan Zhang, Samy Bengio, Moritz Hardt, Benjamin Recht, and Oriol Vinyals. Understanding deep learning requires rethinking generalization. In *International Conference on Learning Representations*, 2017. URL <https://openreview.net/forum?id=Sy8gdB9xx>.
- [73] Hongyi Zhang and Suvrit Sra. First-order methods for geodesically convex optimization. In *Conference on learning theory*, pp. 1617–1638. PMLR, 2016.
- [74] Jinxin Zhou, Xiao Li, Tianyu Ding, Chong You, Qing Qu, and Zhihui Zhu. On the optimization landscape of neural collapse under mse loss: Global optimality with unconstrained features. In *International Conference on Machine Learning*, pp. 27179–27202. PMLR, 2022.
- [75] Jinxin Zhou, Chong You, Xiao Li, Kangning Liu, Sheng Liu, Qing Qu, and Zhihui Zhu. Are all losses created equal: A neural collapse perspective. *Advances in Neural Information Processing Systems*, 35:31697–31710, 2022.
- [76] Zhihui Zhu, Tianyu DING, Jinxin Zhou, Xiao Li, Chong You, Jeremias Sulam, and Qing Qu. A geometric analysis of neural collapse with unconstrained features. In A. Beygelzimer, Y. Dauphin, P. Liang, and J. Wortman Vaughan (eds.), *Advances in Neural Information Processing Systems*, 2021. URL <https://openreview.net/forum?id=KRODJAA6pZE>.

Appendices

A Background

Following the definitions of the global mean and class mean of the penultimate-layer features $\{\mathbf{h}_{c,i}\}$ in Equation 4, here we introduce the within-class and between-class covariances,

$$\begin{aligned}\Sigma_W &\triangleq \frac{1}{N} \sum_{c=1}^C \sum_{i=1}^{n_c} (\mathbf{h}_{c,i} - \bar{\mathbf{h}}_c) (\mathbf{h}_{c,i} - \bar{\mathbf{h}}_c)^\top, \\ \Sigma_B &\triangleq \frac{1}{C} \sum_{c=1}^C (\bar{\mathbf{h}}_c - \mathbf{h}_G) (\bar{\mathbf{h}}_c - \mathbf{h}_G)^\top.\end{aligned}\quad (10)$$

We proceed by expanding on the four key properties of the last-layer activations and classifiers, as empirically observed by Papayan et al. [50] at the terminal phase of training (TPT), where we have achieved zero classification error and continue towards zero loss.

NC1 Variability Collapse: Throughout training, feature activation variability diminishes as they converge towards their respective class means.

$$NC1 \triangleq \frac{1}{C} \text{Tr} \left(\Sigma_W \Sigma_B^\dagger \right), \quad (11)$$

where $\dagger$ denotes the Moore–Penrose inverse.

NC2 Convergence to Simplex ETF: The class-mean activation vectors, centred around their global mean, converge to uniform norms while simultaneously maintaining equal-sized and maximally separable angles between them².

$$NC2 \triangleq \left\| \frac{\mathbf{W} \mathbf{W}^\top}{\|\mathbf{W} \mathbf{W}^\top\|_F} - \frac{1}{\sqrt{C-1}} \left(\mathbf{I}_C - \frac{1}{C} \mathbf{1}_C \mathbf{1}_C^\top \right) \right\|_F. \quad (12)$$

NC3 Convergence to Self-duality: The feature class-means and linear classifiers eventually align in a dual vector space up to some scaling.

$$NC3 \triangleq \left\| \frac{\mathbf{W} \bar{\mathbf{H}}}{\|\mathbf{W} \bar{\mathbf{H}}\|_F} - \frac{1}{\sqrt{C-1}} \left(\mathbf{I}_C - \frac{1}{C} \mathbf{1}_C \mathbf{1}_C^\top \right) \right\|_F. \quad (13)$$

NC4 Simplification to Nearest Class-Center (NCC): The network classifier tends to select the class whose mean is closest (in Euclidean distance) to a given deepnet activation.

$$NC4 \triangleq \arg \max_{c'} \langle \mathbf{w}_{c'}, \mathbf{h} \rangle + b_{c'} \rightarrow \arg \min_{c'} \|\mathbf{h} - \bar{\mathbf{h}}_{c'}\|_2. \quad (14)$$

Since the NC2 and NC3 metrics involve normalised matrices, it is not immediately evident whether the linear classifier and the class-mean activations are equinorm. Consequently, we introduce the following definition:

$$\mathbf{W} \bar{\mathbf{H}}\text{-Equinorm} = |W_{\text{equinorm}} - \bar{H}_{\text{equinorm}}|, \quad (15)$$

where $\bar{H}_{\text{equinorm}} = \frac{\text{std}_c(\|\bar{\mathbf{h}}_c - \mathbf{h}_G\|_2)}{\text{avg}_c(\|\bar{\mathbf{h}}_c - \mathbf{h}_G\|_2)}$ and $W_{\text{equinorm}} = \frac{\text{std}_c(\|\mathbf{w}_c\|_2)}{\text{avg}_c(\|\mathbf{w}_c\|_2)}$.

Finally, to measure the extent of the variability collapse of each feature separately, we define the cosine margin metric as follows:

$$CM_{c,i} = \cos \theta_{c,i;j} - \max_{j \neq c} \cos \theta_{c,i;j}, \quad \text{where } \cos \theta_{c,i;j} = \frac{\langle \mathbf{w}_j - \mathbf{w}_G, \mathbf{h}_{c,i} - \mathbf{h}_G \rangle}{\|\mathbf{w}_j - \mathbf{w}_G\|_2 \|\mathbf{h}_{c,i} - \mathbf{h}_G\|_2}. \quad (16)$$

Note that the maximal angle for a simplex ETF vector collection $\{v_c\}_{c=1}^C$ are defined as such:

$$\frac{\langle v_c, v_{c'} \rangle}{\|v_c\|_2 \|v_{c'}\|_2} = \begin{cases} 1, & \text{for } c = c' \\ -\frac{1}{C-1}, & \text{for } c \neq c' \end{cases}. \quad (17)$$

Therefore, we can calculate the theoretical simplex ETF cosine margin as $\frac{C}{C-1}$.

²Mathematically, we have defined NC2 as the collapse of the linear classifiers to a simplex ETF.

B DDN Gradients

The original problem we are trying to solve, as defined in Section 3.2, is the following:

$$\underset{U \in St_C^d}{\text{minimize}} \quad \left\| \tilde{\mathbf{H}} - \mathbf{U} \tilde{\mathbf{M}} \right\|_F^2, \quad (18)$$

or equivalently expanded as such:

$$\begin{aligned} y(\mathbf{U}) \in \arg \min_{\mathbf{U} \in \mathbb{R}^{d \times C}} \quad & \tilde{\mathbf{H}} : \tilde{\mathbf{H}} - \frac{2}{\sqrt{C-1}} \tilde{\mathbf{H}} : \mathbf{U} \tilde{\mathbf{M}} + \frac{1}{C-1} \mathbf{U} : \mathbf{U} \tilde{\mathbf{M}}, \\ \text{subject to} \quad & \mathbf{U}^\top \mathbf{U} = \mathbf{I}_C. \end{aligned} \quad (19)$$

Here, we denote the double-dot operator $:$ as the Frobenius inner product, *i.e.*, $\mathbf{A} : \mathbf{B} = \sum_{i=1}^m \sum_{j=1}^n A_{ij} B_{ij} = \text{Tr}(\mathbf{A}^\top \mathbf{B})$.

To compute the first and second-order gradients of the objective function and constraints, respectively, we need to consider matrices as variables. Utilising matrix differentials and vectorised derivatives, as defined in [46], simplifies the computation of second-order objective derivatives and all constraint derivatives. For the objective function, we have:

$$\begin{aligned} d f(\tilde{\mathbf{H}}, \mathbf{U}) &= -\frac{2}{\sqrt{C-1}} \tilde{\mathbf{H}} : d\mathbf{U} \tilde{\mathbf{M}} + \frac{1}{C-1} \left(d\mathbf{U} : \mathbf{U} \tilde{\mathbf{M}} + \mathbf{U} : d\mathbf{U} \tilde{\mathbf{M}} \right) \\ &= \frac{2}{\sqrt{C-1}} \tilde{\mathbf{H}} \tilde{\mathbf{M}} : d\mathbf{U} + \frac{2}{C-1} \mathbf{U} \tilde{\mathbf{M}} : d\mathbf{U} \\ &= \left(\frac{2}{\sqrt{C-1}} \tilde{\mathbf{H}} \tilde{\mathbf{M}} + \frac{2}{C-1} \mathbf{U} \tilde{\mathbf{M}} \right) : d\mathbf{U} \\ \Rightarrow D_{\mathbf{U}} f(\tilde{\mathbf{H}}, \mathbf{U}) &= \frac{2}{\sqrt{C-1}} \tilde{\mathbf{H}} \tilde{\mathbf{M}} + \frac{2}{C-1} \mathbf{U} \tilde{\mathbf{M}} \in \mathbb{R}^{d \times C}. \end{aligned} \quad (20)$$

$$\begin{aligned} d D_{\mathbf{U}}(\tilde{\mathbf{H}}, \mathbf{U}) &= \frac{2}{C-1} d\mathbf{U} \tilde{\mathbf{M}} \\ \Rightarrow d \text{rvec} D_{\mathbf{U}}(\tilde{\mathbf{H}}, \mathbf{U}) &= \frac{2}{C-1} \text{rvec}(d\mathbf{U} \tilde{\mathbf{M}}) = \frac{2}{C-1} \left(\mathbf{I}_d \otimes \tilde{\mathbf{M}} \right) d \text{rvec} \mathbf{U} \\ \Rightarrow \text{rvec}(D_{\mathbf{U}}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) &= \frac{2}{C-1} \left(\mathbf{I}_d \otimes \tilde{\mathbf{M}} \right) \in \mathbb{R}^{dC \times dC}, \end{aligned} \quad (21)$$

where we have defined here the row-major vectorisation method, *i.e.*, $\text{rvec}(\mathbf{A}) = \text{vec}(\mathbf{A}^\top)$, and we have used the property:

$$\text{rvec}(\mathbf{ABC}) = (\mathbf{A} \otimes \mathbf{C}^\top) \text{rvec}(\mathbf{B}). \quad (22)$$

A similar proof follows for the second-order partial gradient of the objective. We omit the proof here and just declare the result:

$$\mathbf{B} \triangleq \text{rvec}(D_{\mathbf{H}}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) = -\frac{2}{\sqrt{C-1}} \left(\mathbf{I}_d \otimes \tilde{\mathbf{M}} \right) \in \mathbb{R}^{dC \times dC}. \quad (23)$$

Regarding the gradients of the constraint function, we have:

$$\begin{aligned}
dJ(\tilde{\mathbf{H}}, \mathbf{U}) &= d(\mathbf{U}^\top \mathbf{U} - \mathbf{I}_C) = (d\mathbf{U})^\top \mathbf{U} + \mathbf{U}^\top d\mathbf{U} \\
\implies d \operatorname{rvec} J(\tilde{\mathbf{H}}, \mathbf{U}) &= \operatorname{rvec}((d\mathbf{U})^\top \mathbf{U}) + \operatorname{rvec}(\mathbf{U}^\top d\mathbf{U}) \\
&= (\mathbf{I}_C \otimes \mathbf{U}^\top) d \operatorname{rvec} \mathbf{U}^\top + (\mathbf{U}^\top \otimes \mathbf{I}_C) d \operatorname{rvec} \mathbf{U} \\
&= (\mathbf{I}_C \otimes \mathbf{U}^\top) \mathbf{K}_{dC} d \operatorname{rvec} \mathbf{U} + (\mathbf{U}^\top \otimes \mathbf{I}_C) d \operatorname{rvec} \mathbf{U} \\
&= \mathbf{K}_{CC}(\mathbf{U}^\top \otimes \mathbf{I}_C) d \operatorname{rvec} \mathbf{U} + (\mathbf{U}^\top \otimes \mathbf{I}_C) d \operatorname{rvec} \mathbf{U} \\
&= (\mathbf{K}_{CC} + \mathbf{I}_{C^2})(\mathbf{U}^\top \otimes \mathbf{I}_C) d \operatorname{rvec} \mathbf{U} \\
&= (\mathbf{K}_{CC} + \mathbf{I}_{C^2})(\mathbf{U} \otimes \mathbf{I}_C)^\top d \operatorname{rvec} \mathbf{U} \\
\implies \operatorname{rvec}(D_{\mathbf{U}} J(\tilde{\mathbf{H}}, \mathbf{U})) &= (\mathbf{K}_{CC} + \mathbf{I}_{C^2})(\mathbf{U} \otimes \mathbf{I}_C)^\top \in \mathbb{R}^{CC \times dC},
\end{aligned} \tag{24}$$

where we have defined the commutation matrix [46, 13] as $\mathbf{K}_{mn}$, as a matrix which satisfies the two following identities for any given $\mathbf{A} \in \mathbb{R}^{m \times n}$ and $\mathbf{B} \in \mathbb{R}^{r \times q}$:

$$\begin{aligned}
\mathbf{K}_{mn} \operatorname{vec}(\mathbf{A}) &= \operatorname{vec}(\mathbf{A}^\top), \\
\mathbf{K}_{rm}(\mathbf{A} \otimes \mathbf{B})\mathbf{K}_{nq} &= \mathbf{B} \otimes \mathbf{A}.
\end{aligned} \tag{25}$$

The gradient in Equation 24 contains redundant constraints because of the symmetrical nature of the orthogonality constraints. To retain only the non-redundant constraints, we must undertake a half-vectorisation procedure. Given that we already possess the fully vectorised gradients (which are simpler to compute in this scenario), we require an elimination matrix, $\mathbf{L}_C \in \mathbb{R}^{\frac{C(C+1)}{2} \times C^2}$, such that

$$\mathbf{L}_C \operatorname{rvec}(\mathbf{A}) = \operatorname{rvech}(\mathbf{A}). \tag{26}$$

We can define an elimination matrix explicitly as follows [45]:

$$\mathbf{L}_n = \sum_{i \geq j} \mathbf{u}_{ij} \operatorname{vec}(\mathbf{E}_{ij})^\top = \sum_{i \geq j} (\mathbf{u}_{ij} \otimes \mathbf{e}_j^\top \otimes \mathbf{e}_i^\top), \tag{27}$$

where

$$\mathbf{u}_{ij} = \begin{cases} 1 & \text{in position } (j-1)n + i - \frac{1}{2}j(j-1) \\ 0 & \text{elsewhere} \end{cases}. \tag{28}$$

Hence,

$$\mathbf{A} \triangleq \operatorname{rvech}(D_{\mathbf{U}} J(\tilde{\mathbf{H}}, \mathbf{U})) = \mathbf{L}_C(\mathbf{K}_{CC} + \mathbf{I}_{C^2})(\mathbf{U} \otimes \mathbf{I}_C)^\top \in \mathbb{R}^{\frac{C(C+1)}{2} \times dC}. \tag{29}$$

We have the following useful Corollary for the second-order gradient of the constraint function.

Corollary 1 (Magnus & Neudecker [46]). *Let ϕ be a twice differentiable real-valued function of an $n \times q$ matrix $\mathbf{X}$. Then, the following two relationships hold between the second differential and the Hessian matrix of ϕ at $\mathbf{X}$:*

$$d^2 \phi(\mathbf{X}) = \operatorname{Tr}(\mathbf{A}(d\mathbf{X})^\top \mathbf{B} d\mathbf{X}) \iff H\phi(\mathbf{X}) = \frac{1}{2}(\mathbf{A}^\top \otimes \mathbf{B} + \mathbf{A} \otimes \mathbf{B}^\top).$$

With the above corollary established, we can have,

$$\begin{aligned}
J(\tilde{\mathbf{H}}, \mathbf{U}) &= \mathbf{U}^\top \mathbf{U} - \mathbf{I}_C \\
dJ(\tilde{\mathbf{H}}, \mathbf{U}) &= (d\mathbf{U})^\top \mathbf{U} + \mathbf{U}^\top d\mathbf{U} \\
d^2 J(\tilde{\mathbf{H}}, \mathbf{U}) &= (d\mathbf{U})^\top d\mathbf{U} + (d\mathbf{U})^\top d\mathbf{U} = 2(d\mathbf{U})^\top d\mathbf{U} \\
d^2 J(\tilde{\mathbf{H}}, \mathbf{U})_{ij} &= 2\mathbf{e}_i^\top (d\mathbf{U})^\top (d\mathbf{U}) \mathbf{e}_j = 2 \operatorname{Tr}(\mathbf{e}_j \mathbf{e}_i^\top (d\mathbf{U})^\top d\mathbf{U}),
\end{aligned} \tag{30}$$

and hence, we get:

$$\text{rvec}(D_{UU}^2 J(\tilde{\mathbf{H}}, \mathbf{U})_{ij}) = \mathbf{I}_d \otimes (\mathbf{e}_i \mathbf{e}_j^\top) + \mathbf{I}_d \otimes (\mathbf{e}_j \mathbf{e}_i^\top) \in \mathbb{R}^{dC \times dC}. \quad (31)$$

Then, we repeat the process with the elimination matrix to eliminate the redundant constraints. However, while this is one way to compute the derivative, a more efficient approach exists, namely the *embedded gradient vector field* method, which we outline in the following subsection. For a complete overview of the method, we recommend readers to follow through the works of Birtea et al. [9, 10], Birtea & Comănescu [8, 7].

Finally, the gradients for the proximal problem in Equation 7 are straightforward to compute, with the only change in gradients being the added proximal terms in:

$$\begin{aligned} D_U f(\tilde{\mathbf{H}}, \mathbf{U}) &= \frac{2}{K-1} \mathbf{U} \tilde{\mathbf{M}} - \frac{2}{\sqrt{K-1}} \tilde{\mathbf{H}} \tilde{\mathbf{M}} + \delta(\mathbf{U} - \mathbf{U}_{\text{prox}}) \in \mathbb{R}^{d \times C}, \\ \text{rvec}(D_{UU}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) &= \frac{2}{K-1} \left(\mathbf{I}_d \otimes \tilde{\mathbf{M}} \right) + \delta \mathbf{I}_{dC} \in \mathbb{R}^{dC \times dC}. \end{aligned} \quad (32)$$

B.1 Implicit Formulation of Lagrange Multipliers on Differentiable Manifolds

An issue arises in Proposition 1 with the expression for $\mathbf{G}$,

$$\mathbf{G} = \text{rvec}(D_{UU}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) - \mathbf{\Lambda} : \text{rvec}(D_{UU}^2 J(\tilde{\mathbf{H}}, \mathbf{U})) \in \mathbb{R}^{dC \times dC}, \quad (33)$$

where both the calculation of the Lagrange multiplier matrix $\mathbf{\Lambda}$ (solved via a linear system) and the construction of the fourth-order tensor representing the second-order derivatives of the constraint function are complex and challenging. However, by recognising the manifold structure of the problem, we can reformulate Equation 33 in a simpler and more computationally efficient way. The embedded gradient vector field method offers such a solution [8].

For a general Riemannian manifold $(\mathcal{M}, g)$, we can define the Gram matrix for the smooth functions $f_1, \dots, f_s, h_1, \dots, h_r : (\mathcal{M}, g) \rightarrow \mathbb{R}$ as follows,

$$\text{Gram}_{(h_1, \dots, h_r)}^{(f_1, \dots, f_s)} \triangleq \begin{bmatrix} \langle \nabla h_1, \nabla f_1 \rangle & \dots & \langle \nabla h_r, \nabla f_1 \rangle \\ \vdots & \ddots & \vdots \\ \langle \nabla h_1, \nabla f_s \rangle & \dots & \langle \nabla h_r, \nabla f_s \rangle \end{bmatrix} \in \mathbb{R}^{s \times r}. \quad (34)$$

In our problem, we are working with a compact Stiefel manifold, which is an embedded sub-manifold of $\mathbb{R}^{d \times C}$, and we identify the isomorphism (via vec) between $\mathbb{R}^{d \times C}$ and $\mathbb{R}^{dC}$. A Stiefel manifold $St_C^d = \{\mathbf{U} \in \mathbb{R}^{d \times C} \mid \mathbf{U}^\top \mathbf{U} = \mathbf{I}_C\}$ can be characterised by a set of constraint functions, $j_s, j_{pq} : \mathbb{R}^{dC} \rightarrow \mathbb{R}$, as follows:

$$\begin{aligned} j_s(\mathbf{u}) &= \frac{1}{2} \|\mathbf{u}_s\|^2, & 1 \leq s \leq C, \\ j_{pq}(\mathbf{u}) &= \langle \mathbf{u}_p, \mathbf{u}_q \rangle, & 1 \leq p < q \leq C. \end{aligned} \quad (35)$$

We consider a smooth cost function $\tilde{f} : St_C^d \rightarrow \mathbb{R}$ and define $f : \mathbb{R}^{d \times C} \rightarrow \mathbb{R}$ as a smooth extension (or prolongation) of $\tilde{f}$. The embedded gradient vector field (after applying the isomorphism) is defined on the open set $\mathbb{R}_{reg}^{dC} \subset \mathbb{R}^{dC}$ formed with the regular leaves of the constrained function $\mathbf{j} : \mathbb{R}^{dC} \rightarrow \mathbb{R}^{\frac{C(C+1)}{2}}$ and is tangent to the foliation of this function. The vector field takes the form [9]:

$$\partial f(\mathbf{u}) = \nabla f(\mathbf{u}) - \sum_{1 \leq s \leq C} \sigma_s(\mathbf{u}) \nabla j_s(\mathbf{u}) - \sum_{1 \leq p < q \leq C} \sigma_{pq}(\mathbf{u}) \nabla j_{pq}(\mathbf{u}), \quad (36)$$

where σ_s, σ_{pq} are the Lagrange multiplier functions defined as such [7]:

$$\sigma_s(\mathbf{u}) = \frac{\det \left(\text{Gram}_{(j_1, \dots, j_{s-1}, f, j_{s+1}, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}^{(j_1, \dots, j_{s-1}, j_s, j_{s+1}, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}(\mathbf{u}) \right)}{\det \left(\text{Gram}_{(j_1, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}^{(j_1, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}(\mathbf{u}) \right)} \triangleq \frac{\det(\text{Gram}_s(\mathbf{u}))}{\det(\text{Gram}(\mathbf{u}))}, \quad (37)$$

$$\sigma_{pq}(\mathbf{u}) = \frac{\det \left(\text{Gram}_{(j_1, \dots, j_C, j_{12}, \dots, j_{pq-1}, f, j_{pq+1}, \dots, j_{C-1}, C)}^{(j_1, \dots, j_C, j_{12}, \dots, j_{pq-1}, j_{pq}, j_{pq+1}, \dots, j_{C-1}, C)}(\mathbf{u}) \right)}{\det \left(\text{Gram}_{(j_1, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}^{(j_1, \dots, j_C, j_{12}, \dots, j_{C-1}, C)}(\mathbf{u}) \right)} \triangleq \frac{\det(\text{Gram}_{pq}(\mathbf{u}))}{\det(\text{Gram}(\mathbf{u}))}.$$

The Gram matrix in the denominator of the Lagrange multiplier functions can be defined as a block matrix as follows,

$$\text{Gram}(\mathbf{u}) = \begin{bmatrix} \mathbf{A} & \mathbf{C}^\top \\ \mathbf{C} & \mathbf{B} \end{bmatrix}, \quad (38)$$

where

$$\mathbf{A} = \begin{bmatrix} \langle \nabla j_1, \nabla j_1 \rangle & \dots & \langle \nabla j_C, \nabla j_1 \rangle \\ \vdots & \ddots & \vdots \\ \langle \nabla j_1, \nabla j_C \rangle & \dots & \langle \nabla j_C, \nabla j_C \rangle \end{bmatrix},$$

$$\mathbf{B} = \begin{bmatrix} \langle \nabla j_{12}, \nabla j_{12} \rangle & \dots & \langle \nabla j_{C-1, C}, \nabla j_{12} \rangle \\ \vdots & \ddots & \vdots \\ \langle \nabla j_{12}, \nabla j_{C-1, C} \rangle & \dots & \langle \nabla j_{C-1, C}, \nabla j_{C-1, C} \rangle \end{bmatrix},$$

$$\mathbf{C} = \begin{bmatrix} \langle \nabla j_1, \nabla j_{12} \rangle & \dots & \langle \nabla j_C, \nabla j_{12} \rangle \\ \vdots & \ddots & \vdots \\ \langle \nabla j_1, \nabla j_{C-1, C} \rangle & \dots & \langle \nabla j_C, \nabla j_{C-1, C} \rangle \end{bmatrix}.$$

For $\text{Gram}_s(\mathbf{u})$ and $\text{Gram}_{pq}(\mathbf{u})$, we can similarly define block decompositions by replacing the appropriate column of the $\text{Gram}(\mathbf{u})$ based on the index. For instance, to define $\text{Gram}_s(\mathbf{u})$, we replace the column s with $[\langle \nabla f(\mathbf{u}), \nabla j_i(\mathbf{u}) \rangle]_{i=1}^{C-1, C}$.

It has been proved [9] that if $\mathbf{U} \in St_C^d$ is a critical point of the function $\tilde{f}$, i.e., $\partial f(\mathbf{u}) = 0$, then $\sigma_s(\mathbf{u}), \sigma_{pq}(\mathbf{u})$ become the classical Lagrange multipliers.

Proposition 2 (Lagrange multiplier functions for Stiefel Manifolds). *The Lagrange multiplier functions are described as a function of the constraint functions in Equation 35 in the following way:*

$$\sigma_s(\mathbf{u}) = \langle \nabla f(\mathbf{u}), \nabla j_s(\mathbf{u}) \rangle = \left\langle \frac{\partial f}{\partial \mathbf{u}_s}(\mathbf{u}), \mathbf{u}_s \right\rangle, \quad (39)$$

$$\sigma_{pq}(\mathbf{u}) = \frac{1}{2} \langle \nabla f(\mathbf{u}), \nabla j_{pq}(\mathbf{u}) \rangle = \frac{1}{2} \left(\left\langle \frac{\partial f}{\partial \mathbf{u}_q}(\mathbf{u}), \mathbf{u}_p \right\rangle + \left\langle \frac{\partial f}{\partial \mathbf{u}_p}(\mathbf{u}), \mathbf{u}_q \right\rangle \right).$$

Proof. We take the definition of the Lagrange multiplier functions in Equation 37. Starting with the denominator, we observe that the Gram matrix is of the form:

$$\text{Gram}(\mathbf{u}) = \begin{bmatrix} \mathbf{I}_C & \mathbf{0} \\ \mathbf{0} & 2\mathbf{I}_{\frac{C(C-1)}{2}} \end{bmatrix}, \quad (40)$$

since for each block matrix $(\mathbf{A}, \mathbf{B}, \mathbf{C})$ of the Gramian, we get for each of their elements:

- $\forall A_{s,r}: \langle \nabla j_s(\mathbf{u}), \nabla j_r(\mathbf{u}) \rangle = \delta_{sr}$
- $\forall B_{\gamma\tau, \alpha\beta}: \langle \nabla j_{\gamma\tau}(\mathbf{u}), \nabla j_{\alpha\beta}(\mathbf{u}) \rangle = 2\delta_{\gamma\alpha}\delta_{\tau\beta} + 2\delta_{\tau\alpha}\delta_{\gamma\beta}$

$$\bullet \forall C_{s,\alpha\beta}: \langle \nabla j_s(\mathbf{u}), \nabla j_{\alpha\beta}(\mathbf{u}) \rangle = 2\delta_{s\alpha}\delta_{s\beta} = 0$$

where δ_{ij} is defined as the Kronecker delta symbol, i.e., $\delta_{ij} = [i = j]$.

Thus, the determinant of this Gram matrix is:

$$\det(\text{Gram}(\mathbf{u})) = \det \begin{bmatrix} \mathbf{I}_C & \mathbf{0} \\ \mathbf{0} & 2\mathbf{I}_{\frac{C(C-1)}{2}} \end{bmatrix} = \det(\mathbf{I}_C) \det(2\mathbf{I}_{\frac{C(C-1)}{2}}) = 2^{\frac{C(C-1)}{2}}. \quad (41)$$

Now, let us turn our focus to the numerator of the Lagrange multiplier function $\sigma_s(\mathbf{u})$.

$$\begin{aligned} \det(\text{Gram}_s(\mathbf{u})) &= \det \left[\begin{array}{ccc|ccc} 1 & \dots & \langle \nabla f(\mathbf{u}), \nabla j_1(\mathbf{u}) \rangle & \dots & 0 & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & \langle \nabla f(\mathbf{u}), \nabla j_s(\mathbf{u}) \rangle & \dots & 0 & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & \langle \nabla f(\mathbf{u}), \nabla j_C(\mathbf{u}) \rangle & \dots & 1 & 0 & \dots & 0 \\ \hline 0 & \dots & \langle \nabla f(\mathbf{u}), \nabla j_{12}(\mathbf{u}) \rangle & \dots & 0 & 2 & \dots & 0 \\ \vdots & \ddots & \vdots & & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & \langle \nabla f(\mathbf{u}), \nabla j_{C-1,C}(\mathbf{u}) \rangle & \dots & 0 & 0 & \dots & 2 \end{array} \right] \\ &= \langle \nabla f(\mathbf{u}), \nabla j_s(\mathbf{u}) \rangle \cdot \det \begin{bmatrix} \mathbf{I}_{C-1} & \mathbf{0} \\ \mathbf{0} & 2\mathbf{I}_{\frac{C(C-1)}{2}} \end{bmatrix} \\ &= 2^{\frac{C(C-1)}{2}} \langle \nabla f(\mathbf{u}), \nabla j_s(\mathbf{u}) \rangle. \end{aligned} \quad (42)$$

Similarly, for $\sigma_{pq}(\mathbf{u})$, we get:

$$\begin{aligned} \det(\text{Gram}_{pq}(\mathbf{u})) &= \langle \nabla f(\mathbf{u}), \nabla j_{pq}(\mathbf{u}) \rangle \cdot \det \begin{bmatrix} \mathbf{I}_C & \mathbf{0} \\ \mathbf{0} & 2\mathbf{I}_{\frac{C(C-1)}{2}-1} \end{bmatrix} \\ &= 2^{\frac{C(C-1)}{2}-1} \langle \nabla f(\mathbf{u}), \nabla j_{pq}(\mathbf{u}) \rangle. \end{aligned} \quad (43)$$

Finally, we get the result by combining the determinants' results for each Lagrange multiplier function and computing the gradients of the constraint functions.

□

Similarly to the gradient vector field, we can show that the Hessian for a constraint manifold (e.g., Stiefel manifold, $\text{Hess}\tilde{f}(\mathbf{u}) : T_{\mathbf{u}}\text{St}_C^d \times T_{\mathbf{u}}\text{St}_C^d \rightarrow \mathbb{R}$) can be given as such [10]:

$$\begin{aligned} \text{Hess}\tilde{f}(\mathbf{u}) &= \left(\text{Hess}f(\mathbf{u}) - \sum_{1 \leq s \leq C} \sigma_s(\mathbf{u}) \text{Hess}j_s(\mathbf{u}) \right. \\ &\quad \left. - \sum_{1 \leq p < q \leq C} \sigma_{pq}(\mathbf{u}) \text{Hess}j_{pq}(\mathbf{u}) \right)_{|T_{\mathbf{u}}\text{St}_C^d \times T_{\mathbf{u}}\text{St}_C^d}. \end{aligned} \quad (44)$$

We can transform the results into a matrix form in the following way. First, we denote,

$$\nabla f(\mathbf{U}) \triangleq \text{vec}^{-1}(\nabla f(\mathbf{u})) \in \mathbb{R}^{d \times C}; \quad \partial f(\mathbf{U}) \triangleq \text{vec}^{-1}(\partial f(\mathbf{u})) \in \mathbb{R}^{d \times C}. \quad (45)$$

We then introduce a symmetric matrix $\Sigma(\mathbf{U}) \triangleq [\sigma_{pq}(\mathbf{u})] \in \mathbb{R}^{C \times C}$, where we also include the Lagrange multiplier function's symmetrical component. For the Stiefel manifold, the matrix form of the embedded gradient vector field (after some simple computation) is given by,

$$\partial f(\mathbf{U}) = \nabla f(\mathbf{U}) - \mathbf{U}\Sigma(\mathbf{U}), \quad (46)$$

where $\Sigma(\mathbf{U}) = \frac{1}{2} (\nabla f(\mathbf{U})^\top \mathbf{U} + \mathbf{U}^\top \nabla f(\mathbf{U}))$.

Regarding the Hessian in Equation 44, to write it in matrix form, we first need to compute the Hessian matrices of the constraint functions as follows³:

$$\begin{aligned}
 [\text{Hess} j_s(\mathbf{U})] &= \begin{bmatrix} \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \mathbf{0}_d & \cdots & \mathbf{I}_d & \cdots & \mathbf{0}_d \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d \end{bmatrix} = \mathbf{I}_d \otimes (\mathbf{e}_s \otimes \mathbf{e}_s^\top); \\
 [\text{Hess} j_{pq}(\mathbf{U})] &= \begin{bmatrix} \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{I}_d & \cdots & \mathbf{0}_d \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ \mathbf{0}_d & \cdots & \mathbf{I}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d \\ \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d & \cdots & \mathbf{0}_d \end{bmatrix} = \mathbf{I}_d \otimes (\mathbf{e}_p \otimes \mathbf{e}_q^\top + \mathbf{e}_q \otimes \mathbf{e}_p^\top).
 \end{aligned} \tag{47}$$

Finally, the matrix form of the Hessian of the cost function $\tilde{f} : St_C^d \rightarrow \mathbb{R}$ is given by,

$$\text{Hess} \tilde{f}(\mathbf{U}) = (\text{Hess} f(\mathbf{U}) - \mathbf{I}_d \otimes \Sigma(\mathbf{U}))|_{T_U St_C^d \times T_U St_C^d}, \tag{48}$$

where from there, we can re-define the expression $\mathbf{G}$ in Equation 33 as follows,

$$\mathbf{G} = \text{rvec}(D_{UU}^2 f(\tilde{\mathbf{H}}, \mathbf{U})) - \mathbf{I}_d \otimes \Sigma(\mathbf{U}) \in \mathbb{R}^{dC \times dC}. \tag{49}$$

C Additional Experimental Results

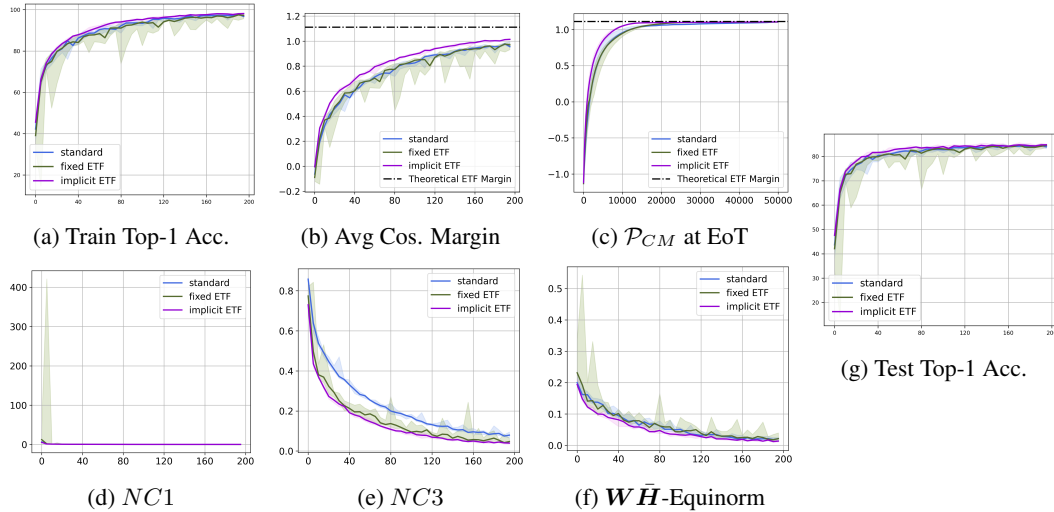


Figure 6: CIFAR10 results on ResNet-18. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

³Here the resulting Kronecker product is expressed in a row-major way.

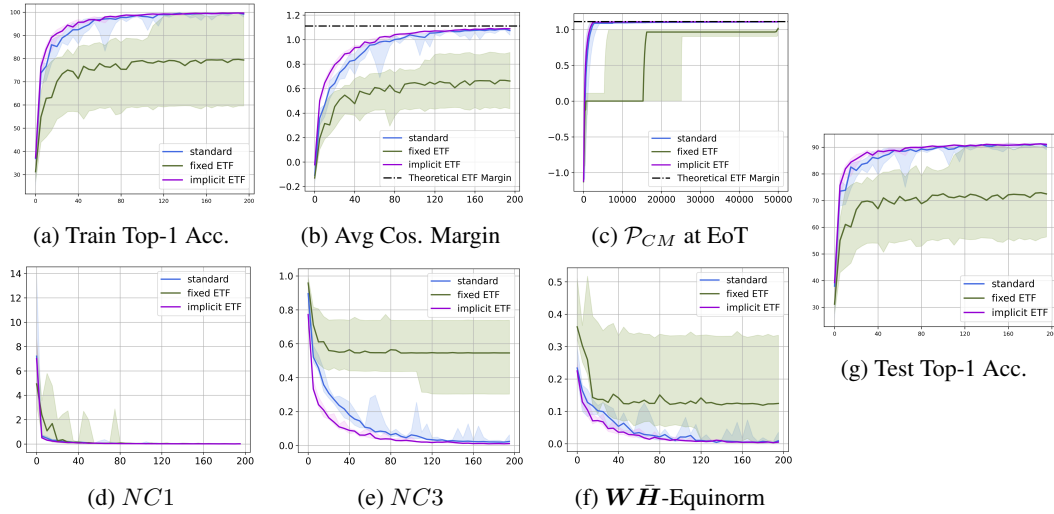


Figure 7: CIFAR10 results on VGG-13. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

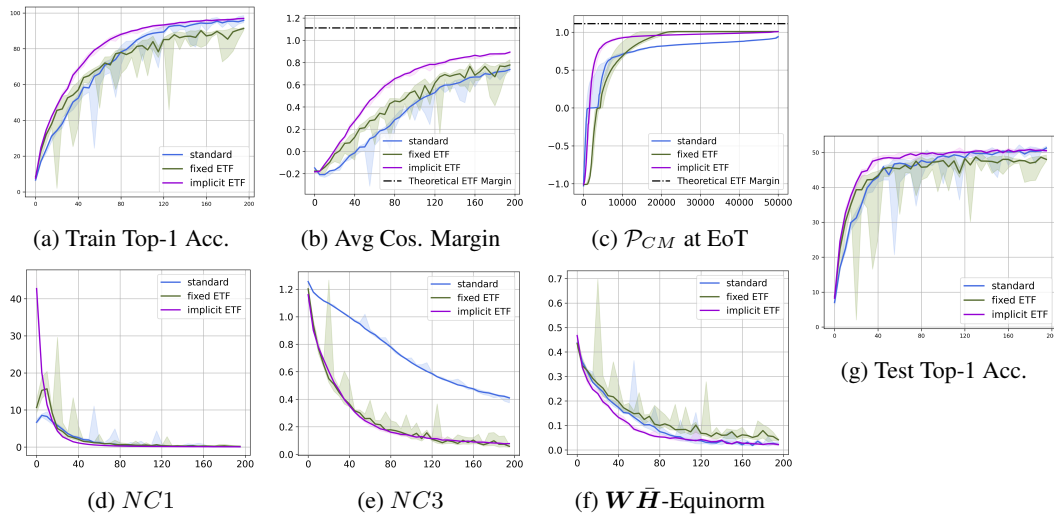


Figure 8: CIFAR100 results on ResNet-50. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

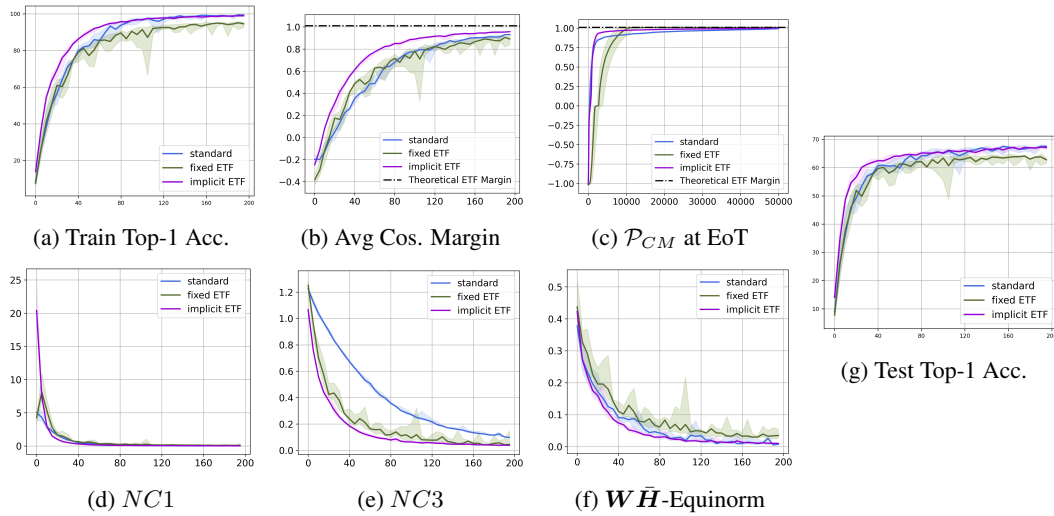


Figure 9: CIFAR100 results on VGG-13. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

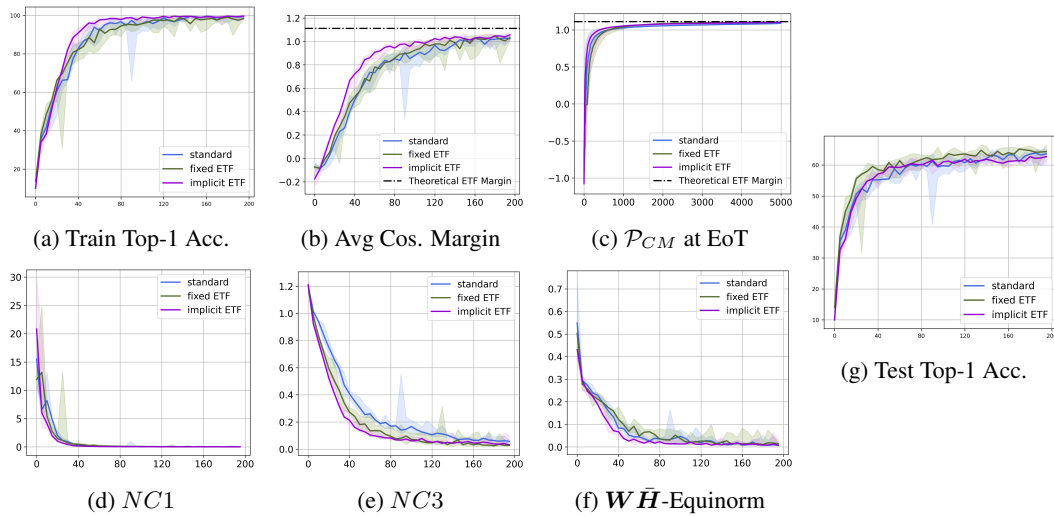


Figure 10: STL10 results on ResNet-50. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

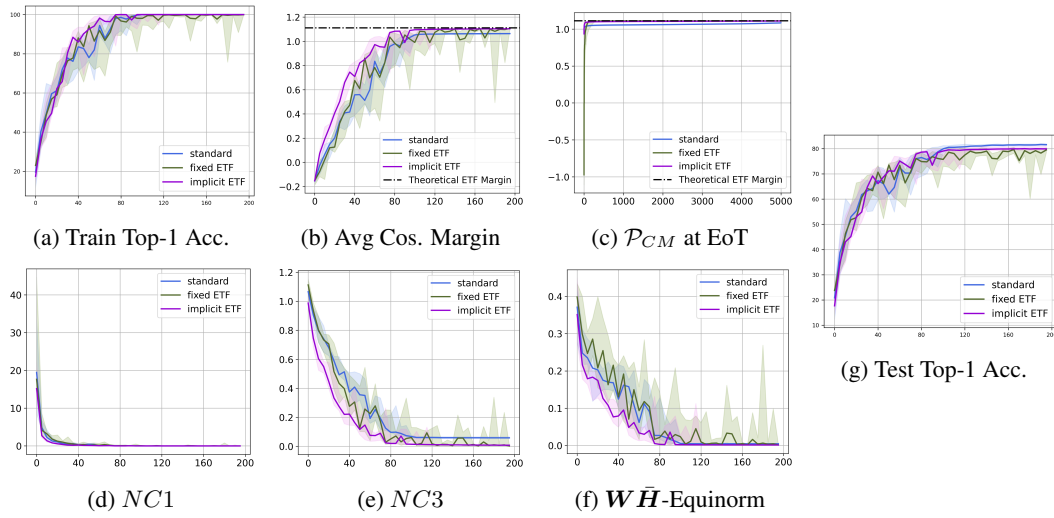


Figure 11: STL10 results on VGG-13. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

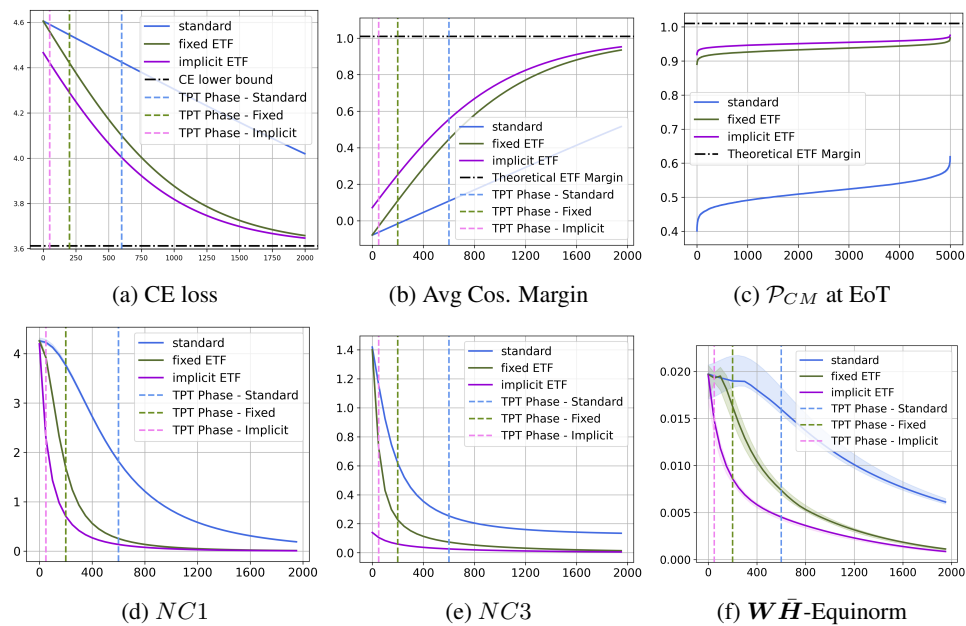


Figure 12: UFM-100 results. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

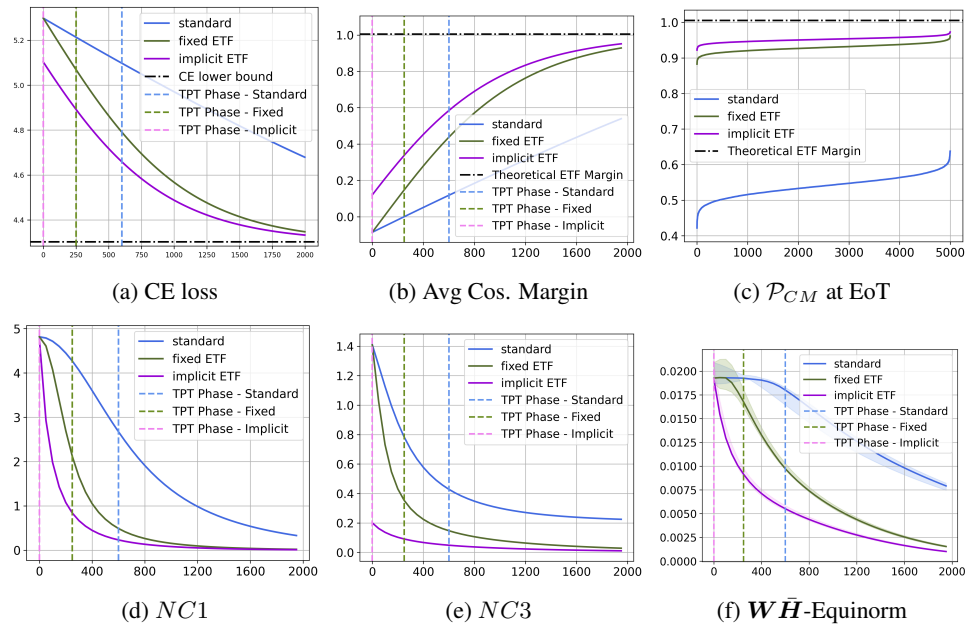


Figure 13: UFM-200 results. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

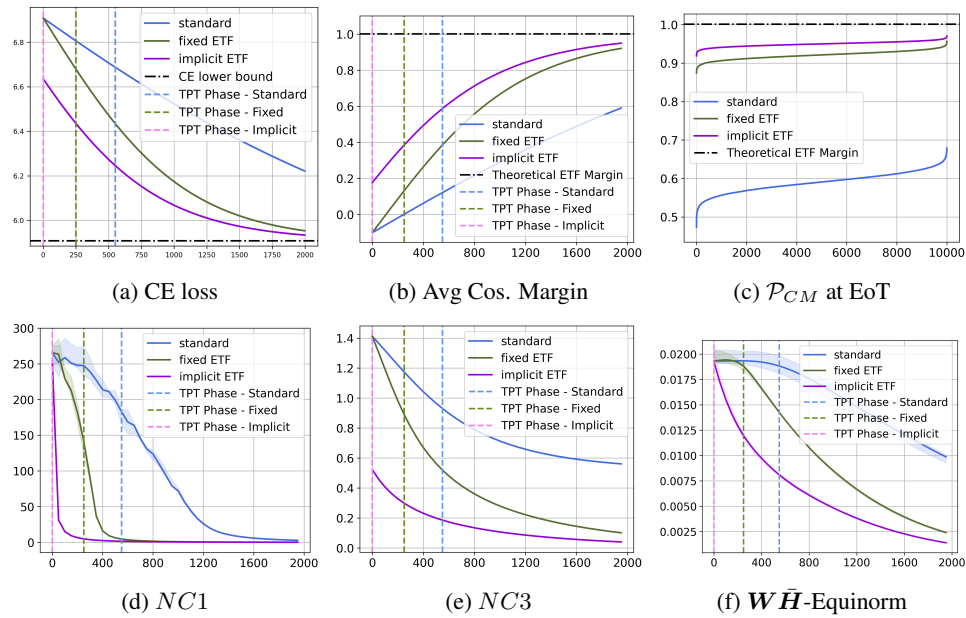


Figure 14: UFM-1000 results. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

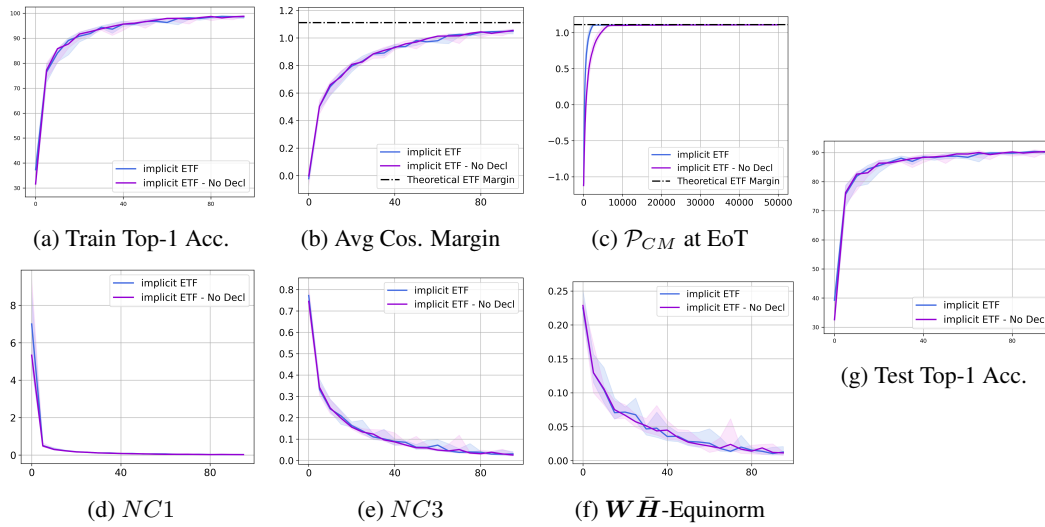
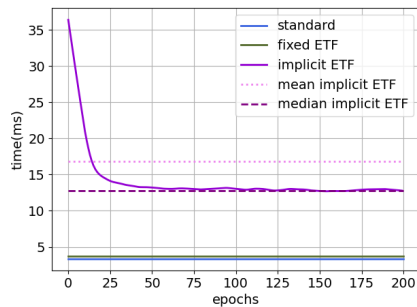
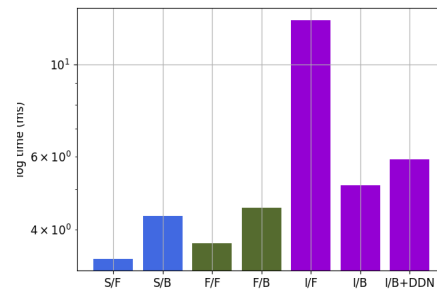


Figure 15: CIFAR10 results on VGG-13, comparing the implicit ETF method in two scenarios: one where the DDN gradient is computed and included in the SGD update, and another where the DDN gradient computation is omitted from the update. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.



(a) Forward pass times in milliseconds.



(b) Forward and backward times in (log) milliseconds.

Figure 16: CIFAR10 computational cost results on ResNet-18. In (a), we plot the forward pass time for each method. For the implicit ETF method, which has dynamic computation times, we also include the mean and median time values. In (b), we plot the computational cost for each forward and backward pass across methods. For the implicit ETF forward pass, we have taken its median time. The notation is as follows: S/F = Standard Forward Pass, S/B = Standard Backward Pass, F/F = Fixed ETF Forward Pass, F/B = Fixed ETF Backward Pass, I/F = Implicit ETF Forward Pass, and I/B = Implicit ETF Backward Pass.

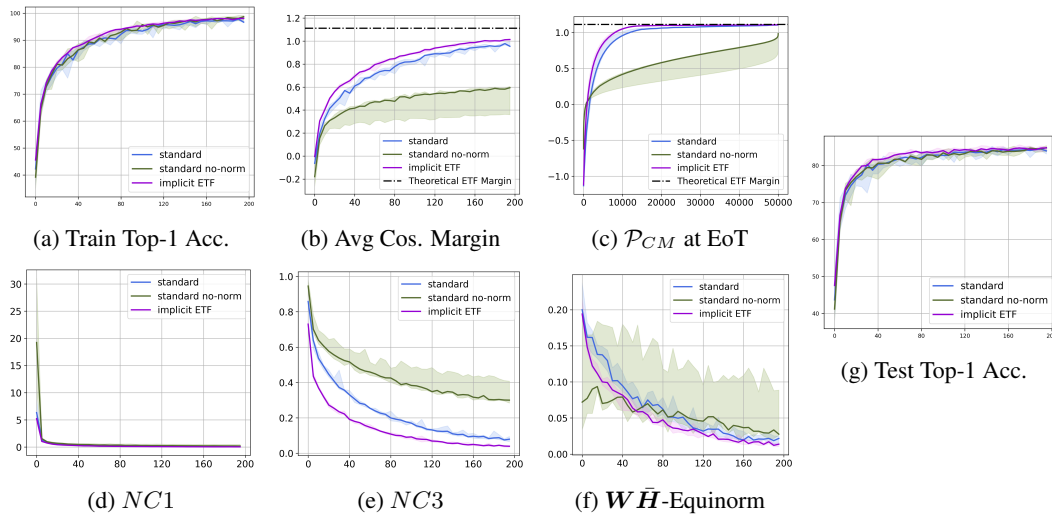


Figure 17: CIFAR10 results on ResNet-18, where in standard no-norm we do not perform feature and weight normalisation. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

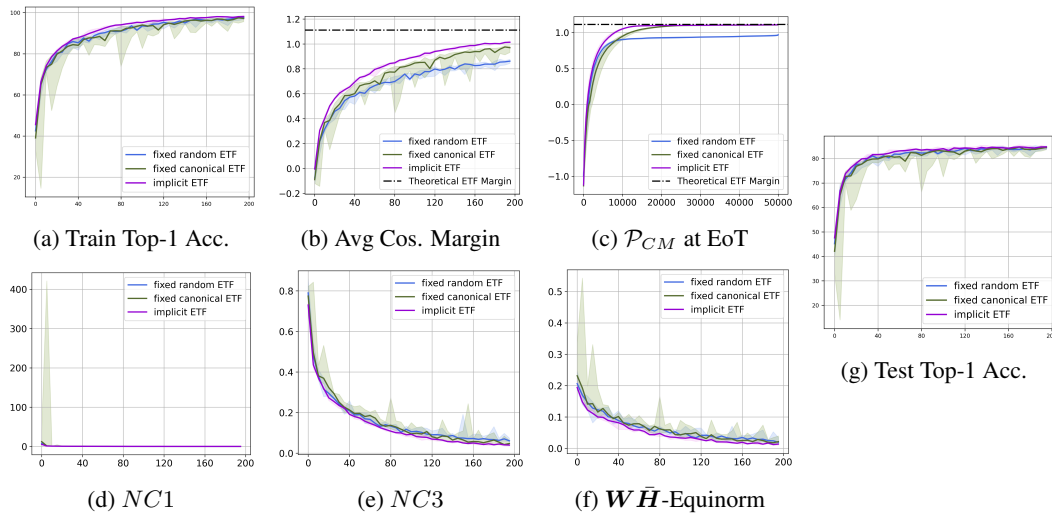


Figure 18: CIFAR10 results on ResNet-18, where in fixed random ETF we choose a random orthogonal direction instead of the canonical one. In all plots, the x-axis represents the number of epochs, except for plot (c), where the x-axis denotes the number of training examples.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: We introduced a novel method leveraging the simplex ETF structure and the neural collapse phenomenon to enhance training convergence and stability in neural networks, and we provided experimental results supporting our claims.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We discussed the limitations of our approach in Section 5 of the paper.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: The paper does not propose new theoretical results. We proposed a new method and its mathematical formulation as well as the appropriate mathematical background.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: All the necessary details required for reproducing the results are presented in the paper. Also, the code will be made publicly available upon acceptance.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: Due to anonymity reasons, the code is not included in the submission. However, it will be made available in a GitHub repository upon acceptance.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: All the specific implementation details can be found in Section 4 of the paper and the appendix.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We have tested our results using different random seeds, and we present the median values along with the range.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have specified the GPUs that were used for our experiments.

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines?>

Answer: [Yes]

Justification: The paper conforms with the NeurIPS code of ethics.

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There is no societal impact of the work performed in this paper.

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper does not pose such risks.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: Yes, every model architecture and dataset used were cited.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.