

Understanding and Training Deep Diagonal Circulant Neural Networks

Alexandre Araujo^{1,2} Benjamin Negrevergne² Yann Chevaleyre² Jamal Atif²

Abstract. In this paper, we study deep diagonal circulant neural networks, which are deep neural networks in which weight matrices are the product of diagonal and circulant ones. Besides making a theoretical analysis of their expressivity, we introduce principled techniques for training these models: we devise an initialization scheme and propose a smart use of non-linearity functions in order to train deep diagonal circulant networks. Furthermore, we show that these networks outperform recently introduced deep networks with other types of structured layers. We conduct a thorough experimental study to compare the performance of deep diagonal circulant networks with state-of-the-art models based on structured matrices and with dense models. We show that our models achieve better accuracy than other structured approaches while requiring 2x fewer weights than the next best approach. Finally, we train compact and accurate deep diagonal circulant networks on a real world video classification dataset with over 3.8 million training examples.

1 Introduction

The deep learning revolution has yielded models of increasingly large size. In recent years, designing compact and accurate neural networks with a small number of trainable parameters has been an active research topic. It is motivated by practical applications in embedded systems (to reduce memory footprint [23]), federated and distributed learning (to reduce communication [15]), derivative-free optimization in reinforcement learning (to simplify the computation of the approximated gradient [8]), etc. Besides a number of practical applications, it is also an important research question whether or not models really need to be this large or if smaller networks can achieve similar accuracy [5].

Structured matrices are at the very core of most of the work on compact networks. In these models, dense weight matrices are replaced by matrices with a prescribed structure (e.g. low rank matrices, Toeplitz matrices, circulant matrices, LDR, etc.). Despite substantial efforts (e.g. [7, 20]), the performance of compact models is still far from achieving an acceptable accuracy motivating their use in real-world scenarios. This raises several questions about the effectiveness of such models and about our ability to train them. In particular two main questions call for investigation:

Q1 How to efficiently train deep neural networks with a large number of structured layers?

Q2 What is the expressive power of structured layers compared to dense layers?

In this paper, we provide principled answers to these questions for the particular case of deep neural networks based on diagonal and circulant matrices (a.k.a. Diagonal-circulant neural networks or DCNNs).

The idea of using diagonal and circulant matrices together comes from a series of results in linear algebra by Muller et al. [21] and Huhtanen et al. [13]. The most recent result from Huhtanen et al. [13] demonstrates that any matrix A in $\mathbb{C}^{n \times n}$ can be decomposed into the product of $2n - 1$ alternating diagonal and circulant matrices. The diagonal-circulant decomposition inspired Moczulski et al. [20] to design the *Structured Efficient Linear Layers* (SELL), which is the building block of DCNNs. However, they were not able to train deep neural networks based on these layers.

To answer **Q1**, we first describe a theoretically sound initialization procedure for DCNN which allows the signal to propagate through the network without vanishing or exploding. Furthermore, we provide a number of empirical insights to explain the behaviour of DCNNs and show the impact of the number of the non-linearities in the network on the convergence rate and the accuracy of the network. By combining all these insights, we are able (for the first time) to train large and deep DCNNs and demonstrate the good performance of these networks on a large scale application (the *YouTube-8M* video classification problem) and obtain very competitive accuracy.

To answer **Q2**, we propose an analysis of the expressivity of DCNNs by extending the results by Huhtanen et al. [13]. We introduce a new bound on the number of diagonal-circulant products required to approximate a matrix that depends on its rank. Building on this result, we demonstrate that a DCNN with bounded width and small depth can approximate any dense networks with ReLU activations.

Outline of the paper: We present in Section 2 the related work on structured neural networks and several compression techniques. Section 3 introduces circulant matrices, our new result extending the one from Huhtanen et al. [13]. Section 4 proposes a theoretical analysis on the expressivity on DCNNs. Section 5 describes two efficient techniques for training deep diagonal circulant neural networks. Finally, Section 6 presents extensive experiments to compare the performance of deep diagonal circulant neural networks in different settings with respect to other state of the art approaches. Section 7 provides a discussion and concluding remarks.

2 Related Work

Structured matrices exhibit a number of good properties which have been exploited by deep learning practitioners, mainly to compress large neural networks architectures into smaller ones. For example, Hinrichs et al. [11] have demonstrated that a single circulant matrix

¹ Wavestone, Paris, France, email: alexandre.araujo@dauphine.psl.eu

² Université Paris-Dauphine, PSL Research University, LAMSADE, CNRS, UMR 7243, Paris, France

can be used to approximate the Johnson-Lindenstrauss transform, often used in machine learning to perform dimensionality reduction. Building upon this result, Cheng et al. [7] proposed to replace the weight matrix of a fully connected layer by a circulant matrix effectively replacing the complex transform modeled by the fully connected layer by a simple dimensionality reduction. Despite the reduction of expressivity, the resulting network demonstrated good accuracy using only a fraction of its original size (90% reduction).

Comparison with ACDC. Moczulski et al. [20] have introduced two *Structured Efficient Linear Layers* (SELL) called AFDF and ACDC, where A and D are diagonal matrices and F and C are the Fourier and cosine transform respectively. The AFDF structured layer benefits from the theoretical results introduced by Huhtanen et al. [13] and can be seen as the building block of DCNNs. However, Moczulski et al. [20] only experiment using ACDC, a different type of layer that does not involve circulant matrices. As far as we can tell, the theoretical guarantees available for the AFDF layer do not apply on the ACDC layer since the cosine transform does not diagonalize circulant matrices [24]. Another possible limit of the ACDC paper is that they only train large neural networks involving ACDC layers combined with many other expressive layers. Although the resulting network demonstrates good accuracy, it is difficult to characterize the true contribution of the ACDC layers in this setting.

Comparison with Low displacement rank structures. More recently, Thomas et al. [27] have generalized these works by proposing neural networks with low-displacement rank matrices (LDR), that are structured matrices encompassing a large family of structured matrices, including Toeplitz-like, Vandermonde-like, Cauchy-like and more notably DCNNs. To obtain this result, LDR represents a structured matrix using two displacement operators and a low-rank residual. Despite being elegant and general, we found that the LDR framework suffers from several limits which are inherent to its generality and makes it difficult to use in the context of large and deep neural networks. First, the training procedure for learning LDR matrices is highly involved and implies many complex mathematical objects such as Krylov matrices. Then, as acknowledged by the authors, the number of parameters required to represent a given structured matrix (e.g. a Toeplitz matrix) in practice is unnecessarily high (higher than required in theory).

Other compression techniques. Besides structured matrices, a variety of techniques have been proposed to build more compact deep learning models. These include *model distillation* [12], Tensor Train [22], Low-rank decomposition [9], to mention a few. However, Circulant networks show good performances in several contexts (the interested reader can refer to the results reported by Moczulski et al. [20] and Thomas et al. [27]).

3 A primer on circulant matrices and a new result

An n -by- n circulant matrix C is a matrix where each row is a cyclic right shift of the previous one as illustrated below.

$$C = \text{circ}(c) = \begin{bmatrix} c_0 & c_{n-1} & c_{n-2} & \dots & c_1 \\ c_1 & c_0 & c_{n-1} & & c_2 \\ c_2 & c_1 & c_0 & & c_3 \\ \vdots & & & \ddots & \vdots \\ c_{n-1} & c_{n-2} & c_{n-3} & & c_0 \end{bmatrix}$$

Circulant matrices exhibit several interesting properties from the perspective of numerical computations. Most importantly, any n -by- n circulant matrix C can be represented using only n coefficients instead

of the n^2 coefficients required to represent classical unstructured matrices. In addition, the matrix-vector product is simplified from $O(n^2)$ to $O(n \log n)$ using the convolution theorem.

As we will show in this paper, circulant matrices also have a strong expressive power. So far, we know that a single circulant matrix can be used to represent a variety of important linear transforms such as random projections [11]. When they are combined with diagonal matrices, they can also be used as building blocks to represent any linear transform [25, 13] with an arbitrary precision. Huhtanen et al. [13] were able to bound the number of factors that is required to approximate any matrix A with arbitrary precision.

Relation between diagonal circulant matrices and low rank matrices We recall this result in Theorem 1 as it is the starting point of our theoretical analysis (note that in the rest of the paper, $\|\cdot\|$ denotes the ℓ_2 norm when applied to vectors, and the operator norm when applied to matrices).

Theorem 1. (Reformulation from Huhtanen et al. [13]) For every matrix $A \in \mathbb{C}^{n \times n}$, for any $\epsilon > 0$, there exists a sequence of matrices $B_1 \dots B_{2n-1}$ where B_i is a circulant matrix if i is odd, and a diagonal matrix otherwise, such that $\|B_1 B_2 \dots B_{2n-1} - A\| < \epsilon$.

Unfortunately, this theorem is of little use to understand the expressive power of diagonal-circulant matrices when they are used in deep neural networks. This is because: 1) the bound only depends on the dimension of the matrix A , not on the matrix itself, 2) the theorem does not provide any insights regarding the expressive power of m diagonal-circulant factors when m is much lower than $2n - 1$ as it is the case in most practical scenarios we consider in this paper.

In the following theorem, we enhance the result by Huhtanen et al. [13] by expressing the number of factors required to approximate A , as a function of the rank of A . This is useful when one deals with low-rank matrices, which is common in machine learning problems.

Theorem 2. ³ (Rank-based circulant decomposition) Let $A \in \mathbb{C}^{n \times n}$ be a matrix of rank at most k . Assume that n can be divided by k . For any $\epsilon > 0$, there exists a sequence of $4k + 1$ matrices $B_1, \dots, B_{4k+1}$, where B_i is a circulant matrix if i is odd, and a diagonal matrix otherwise, such that $\|B_1 B_2 \dots B_{4k+1} - A\| < \epsilon$.

A direct consequence of Theorem 2, is that if the number of diagonal-circulant factors is set to a value K , we can represent all linear transform A whose rank is $\frac{K-1}{4}$.

Compared to [13], this result shows that structured matrices with fewer than $2n$ diagonal-circulant matrices (as it is the case in practice) can still represent a large class of matrices. As we will show in the following section, this result will be useful to analyze the expressivity of neural networks based on diagonal and circulant matrices.

4 Analysis of Diagonal Circulant Neural Networks (DCNNs)

Zhao et al. [31] have shown that circulant networks with 2 layers and unbounded width are universal approximators. However, results on unbounded networks offer weak guarantees and two important questions have remained open until now: 1) Can we approximate any function with a bounded-width circulant networks? 2) What function can we approximate with a circulant network that has a bounded

³ All proofs are in the extended version of the paper.
<https://arxiv.org/abs/1901.10255>

width and a small depth? We answer these two questions in this section.

First, we introduce some necessary definitions regarding neural networks and we provide a theoretical analysis of their approximation capabilities.

Definition 1 (Deep ReLU network). *Given L weight matrices $W = (W_1, \dots, W_L)$ with $W_i \in \mathbb{C}^{n \times n}$ and L bias vectors $b = (b_1, \dots, b_L)$ with $b_i \in \mathbb{C}^n$, a deep ReLU network is a function $f_{W_L, b_L} : \mathbb{C}^n \rightarrow \mathbb{C}^n$ such that $f_{W_L, b_L}(x) = (f_{W_L, b_L} \circ \dots \circ f_{W_1, b_1})(x)$ where $f_{W_i, b_i}(x) = \phi(W_i x + b_i)$ and $\phi(\cdot)$ is a ReLU non-linearity⁴. In the rest of this paper, we call L and n respectively the depth and the width of the network. Moreover, we call total rank k , the sum of the ranks of the matrices $W_1 \dots W_L$. i.e. $k = \sum_{i=1}^L \text{rank}(W_i)$.*

We also need to introduce DCNNs, similarly to Moczulski et al. [20].

Definition 2 (Diagonal Circulant Neural Networks). *Given L diagonal matrices $D = (D_1, \dots, D_L)$ with $D_i \in \mathbb{C}^{n \times n}$, L circulant matrices $C = (C_1, \dots, C_L)$ with $C_i \in \mathbb{C}^{n \times n}$ and L bias vectors $b = (b_1, \dots, b_L)$ with $b_i \in \mathbb{C}^n$, a Diagonal Circulant Neural Networks (DCNN) is a function $f_{W_L, b_L} : \mathbb{C}^n \rightarrow \mathbb{C}^n$ such that $f_{D, C, b}(x) = (f_{D_L, C_L, b_L} \circ \dots \circ f_{D_1, C_1, b_1})(x)$ where $f_{D_i, C_i, b_i}(x) = \phi_i(D_i C_i x + b_i)$ and where $\phi_i(\cdot)$ is a ReLU non-linearity or the identity function.*

We can now show that bounded-width DCNNs can approximate any Deep ReLU Network, and as a corollary, that they are universal approximators.

Lemma 1. *Let $\mathcal{N}$ be a deep ReLU network of width n and depth L , and let $\mathcal{X} \subset \mathbb{C}^n$ be a bounded set. For any $\epsilon > 0$, there exists a DCNN $\mathcal{N}'$ of width n and of depth $(2n-1)L$ such that $\|\mathcal{N}(x) - \mathcal{N}'(x)\| < \epsilon$ for all $x \in \mathcal{X}$.*

We can now state the universal approximation corollary:

Corollary 1. *Bounded width DCNNs are universal approximators in the following sense: for any continuous function $f : [0, 1]^n \rightarrow \mathbb{R}_+$ of bounded supremum norm, for any $\epsilon > 0$, there exists a DCNN $\mathcal{N}_\epsilon$ of width $n+3$ such that $\forall x \in [0, 1]^{n+3}$, $|f(x_1 \dots x_n) - (\mathcal{N}_\epsilon(x))_1| < \epsilon$, where $(\cdot)_i$ represents the i^{th} component of a vector.*

This is a first result, however $(2n+5)L$ is not a small depth (in our experiments, n can be over 300 000), and a number of work provided empirical evidences that DCNN with small depth can offer good performances (e.g. [3, 7]). To improve our result, we introduce our main theorem which studies the approximation properties of these small depth networks.

Theorem 3 (Rank-based expressive power of DCNNs). *Let $\mathcal{N}$ be a deep ReLU network of width n , depth L and a total rank k and assume n is a power of 2. Let $\mathcal{X} \subset \mathbb{C}^n$ be a bounded set. Then, for any $\epsilon > 0$, there exists a DCNN with ReLU activation $\mathcal{N}'$ of width n such that $\|\mathcal{N}(x) - \mathcal{N}'(x)\| < \epsilon$ for all $x \in \mathcal{X}$ and the depth of $\mathcal{N}'$ is bounded by $9k$.*

Remark that in the theorem, we require that n is a power of 2. We conjecture that the result still holds even without this condition. This result refines Lemma 1, and answer our second question: a DCNN

⁴ Because our networks deal with complex numbers, we use an extension of the ReLU function to the complex domain. The most straightforward extension defined in [28] is as follows: $\text{ReLU}(z) = \text{ReLU}(\Re(z)) + i\text{ReLU}(\Im(z))$, where $\Re$ and $\Im$ refer to the real and imaginary parts of z .

of bounded width and small depth can approximate a Deep ReLU network of low total rank. Note that the converse is not true: because n -by- n circulant matrix can be of rank n , approximating a DCNN of depth 1 can require a deep ReLU network of total rank equals to n .

Expressivity of DCNNs For the sake of clarity, we highlight the significance of these results with the two following properties.

Properties. Given an arbitrary fixed integer n , let $\mathcal{R}_k$ be the set of all functions $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ representable by a deep ReLU network of total rank at most k and let $\mathcal{C}_l$ the set of all functions $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ representable by deep diagonal-circulant networks of depth at most l , then:

$$\forall k, \exists l \quad \mathcal{R}_k \subsetneq \mathcal{C}_l \quad (1)$$

$$\forall l, \nexists k \quad \mathcal{C}_l \subseteq \mathcal{R}_k \quad (2)$$

We illustrate the meaning of this properties using Figure 1. As we can see, the set $\mathcal{R}_k$ of all the functions representable by a deep ReLU network of total rank k is strictly included in the set $\mathcal{C}_{9k}$ of all DCNN of depth $9k$ (as by Theorem 3).

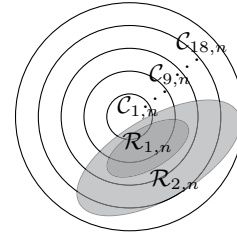


Figure 1. Illustration of Properties 1 and 2.

These properties are interesting for many reasons. First, Property 2 shows that diagonal-circulant networks are *strictly more expressive* than networks with low total rank. Second and most importantly, in standard deep neural networks, it is known that the most of the singular values are close to zero (see e.g. [26, 4]). Property 1 shows that these networks can efficiently be approximated by diagonal-circulant networks. Finally, several publications have shown that neural networks can be trained explicitly to have low-rank weight matrices [17, 10]. This opens the possibility of learning compact and accurate diagonal-circulant networks.

5 How to train very deep DCNNs

Training DCNNs has revealed to be a challenging problem. We devise two techniques to facilitate the training of deep DCNNs. First, we propose an initialization procedure which guarantee the signal is propagated across the network without vanishing nor exploding. Secondly, we study the behavior of DCNNs with different non-linearity functions and determine the best parameters for different settings.

Initialization scheme The following initialization procedure which is a variant of Xavier initialization. First, for each circulant matrix $C = \text{circ}(c_1 \dots c_n)$, each c_i is randomly drawn from $\mathcal{N}(0, \sigma^2)$, with $\sigma = \sqrt{\frac{2}{n}}$. Next, for each diagonal matrix $D = \text{diag}(d_1 \dots d_n)$, each d_i is drawn randomly and uniformly from $\{-1, 1\}$ for all i . Finally, all biases in the network are randomly drawn from $\mathcal{N}(0, \sigma'^2)$, for some small value of σ' . The following proposition states that the covariance matrix at the output of any layer in a DCNN, independent of the depth, is constant.

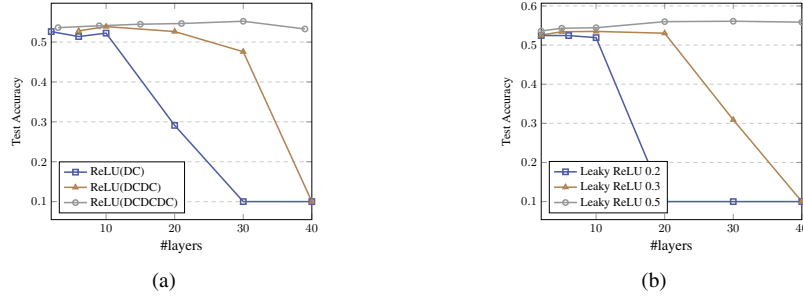


Figure 2. Experiments on training DCNNs and other structured neural networks on CIFAR-10. Figure 2(a): impact of increasing the number of ReLU activations in a DCNN. Deep DCNNs with fewer ReLUs are easier to train. Figure 2(b): impact of increasing the slope of a Leaky-ReLU in DCNNs. Deep DCNNs with a larger slope are easier to train.

Proposition 4. Let $\mathcal{N}$ be a DCNN of depth L initialized according to our procedure, with $\sigma' = 0$. Assume that all layers 1 to $L - 1$ have ReLU activation functions, and that the last layer has the identity activation function. Then, for any $x \in \mathbb{R}^n$, the covariance matrix of $\mathcal{N}(x)$ is $\frac{2 \cdot Id}{n} \|x\|_2^2$. Moreover, note that this covariance does not depend on the depth of the network.

Proof. Let $\mathcal{N} = f_{D_L, C_L} \circ \dots \circ f_{D_1, C_1}$ be a L layer DCNN. All matrices are initialized as described in the statement of the proposition. Let $y = D_1 C_1 x$. Lemma 2 shows that $cov(y_i, y_{i'}) = 0$ for $i \neq i'$ and $var(y_i) = \frac{2}{n} \|x\|_2^2$. For any $j \leq L$, define $z^j = f_{D_j, C_j} \circ \dots \circ f_{D_1, C_1}(x)$. By a recursive application of lemma 2, we get that then $cov(z_i^j, z_{i'}^j) = 0$ and $var(z_i^j) = \frac{2}{n} \|x\|_2^2$. $\square$

Lemma 2. Let $c_1 \dots c_n, d_1 \dots d_n, b_1 \dots b_n$ be random variables in $\mathbb{R}$ such that $c_i \sim \mathcal{N}(0, \sigma^2)$, $b_i \sim \mathcal{N}(0, \sigma'^2)$ and $d_i \sim \{-1, 1\}$ uniformly. Define $C = \text{circ}(c_1 \dots c_n)$ and $D = \text{diag}(d_1 \dots d_n)$. Define $y = DCu$ and $z = CDu$ for some vector u in $\mathbb{R}^n$. Also define $\bar{y} = y + b$ and $\bar{z} = z + b$. Then, for all i , the p.d.f. of $y_i, \bar{y}_i, z_i$ and $\bar{z}_i$ are symmetric. Also:

- Assume $u_1 \dots u_n$ is fixed. Then, we have for $i \neq i'$:

$$\begin{aligned} cov(y_i, y_{i'}) &= cov(z_i, z_{i'}) = cov(\bar{y}_i, \bar{y}_{i'}) = cov(\bar{z}_i, \bar{z}_{i'}) = 0 \\ var(y_i) &= var(z_i) = \sum_j u_j^2 \sigma^2 \\ var(\bar{y}_i) &= var(\bar{z}_i) = \sigma'^2 + \sum_j u_j^2 \sigma^2 \end{aligned}$$

- Let $x_1 \dots x_n$ be random variables in $\mathbb{R}$ such that the p.d.f. of x_i is symmetric for all i , and let $u_i = \text{ReLU}(x_i)$. We have for $i \neq i'$:

$$\begin{aligned} cov(y_i, y_{i'}) &= cov(z_i, z_{i'}) = cov(\bar{y}_i, \bar{y}_{i'}) = cov(\bar{z}_i, \bar{z}_{i'}) = 0 \\ var(y_i) &= var(z_i) = \frac{1}{2} \sum_j var(x_i) \cdot \sigma^2 \\ var(\bar{y}_i) &= var(\bar{z}_i) = \sigma'^2 + \frac{1}{2} \sum_j var(x_i) \cdot \sigma^2 \end{aligned}$$

Proof. By an abuse of notation, we write $c_0 = c_n, c_{-1} = c_{n-1}$ and so on. First, note that: $y_i = \sum_{j=1}^n c_{j-i} u_j d_j$ and $z_i = \sum_{j=1}^n c_{j-i} u_j d_i$. Observe that each term $c_{j-i} u_j d_j$ and $c_{j-i} u_j d_i$ have symmetric p.d.f. because of d_i and d_j . Thus, y_i and z_i have symmetric p.d.f. Now let us compute the covariance.

$$\begin{aligned} cov(y_i, y_{i'}) &= \sum_{j, j'=1}^n cov(c_{j-i} u_j d_j, c_{j'-i'} u_{j'} d_{j'}) \\ &= \sum_{j, j'=1}^n \mathbb{E}[c_{j-i} u_j d_j c_{j'-i'} u_{j'} d_{j'}] \\ &\quad - \mathbb{E}[c_{j-i} u_j d_j] \mathbb{E}[c_{j'-i'} u_{j'} d_{j'}] \end{aligned}$$

Observe that $\mathbb{E}[c_{j-i} u_j d_j] = \mathbb{E}[c_{j-i} u_j] \mathbb{E}[d_j] = 0$ because d_j is independent from $c_{j-i} u_j$. Also, observe that if $j \neq j'$ then $\mathbb{E}[d_j d_{j'}] = 0$ and thus $\mathbb{E}[c_{j-i} u_j d_j c_{j'-i'} u_{j'} d_{j'}] = \mathbb{E}[d_j d_{j'}] \mathbb{E}[c_{j-i} u_j c_{j'-i'} u_{j'}] = 0$. Thus, the only non null terms are those for which $j = j'$. We get:

$$\begin{aligned} cov(y_i, y_{i'}) &= \sum_{j=1}^n \mathbb{E}[c_{j-i} u_j d_j c_{j-i'} u_j d_{j'}] \\ &= \sum_{j=1}^n \mathbb{E}[c_{j-i} c_{j-i'} u_j^2] \end{aligned}$$

Assume u is a fixed vector. Then, $var(y_i) = \sum_{j=1}^n u_j^2 \sigma^2$ and $cov(y_i, y_{i'}) = 0$ for $i \neq i'$ because c_{j-i} is independent from $c_{j-i'}$. Now assume that $u_j = \text{ReLU}(x_j)$ where x_j is a r.v. Clearly, u_j^2 is independent from c_{j-i} and $c_{j-i'}$. Thus:

$$cov(y_i, y_{i'}) = \sum_{j=1}^n \mathbb{E}[c_{j-i} c_{j-i'}] \mathbb{E}[u_j^2]$$

For $i \neq i'$, then c_{j-i} and $c_{j-i'}$ are independent, and thus $\mathbb{E}[c_{j-i} c_{j-i'}] = \mathbb{E}[c_{j-i}] \mathbb{E}[c_{j-i'}] = 0$. Therefore, $cov(y_i, y_{i'}) = 0$ if $i \neq i'$. Let us compute the variance. We get $var(y_i) = \sum_{j=1}^n var(c_{j-i}) \cdot \mathbb{E}[u_j^2]$. Because the p.d.f. of x_j is symmetric, $\mathbb{E}[x_j^2] = 2\mathbb{E}[u_j^2]$ and $\mathbb{E}[x_j] = 0$. Thus, $var(y_i) = \frac{1}{2} \sum_{j=1}^n var(c_{j-i}) \cdot \mathbb{E}[x_j^2] = \frac{1}{2} \sum_{j=1}^n var(c_{j-i}) \cdot var(x_j)$.

Finally, note that $cov(\bar{y}_i, \bar{y}_{i'}) = cov(y_i, y_{i'}) + cov(b_i, b_{i'})$. This yields the covariances of $\bar{y}$.

To derive $cov(z_i, z_{i'})$ and $cov(\bar{z}_i, \bar{z}_{i'})$, the required calculus is nearly identical. We let the reader check by himself/herself. $\square$

Non-linearity function We empirically found that reducing the number of non-linearities in the networks simplifies the training of deep neural networks. To support this claim, we conduct a series of experiments on various DCNNs with a varying number of ReLU activations (to reduce the number of non-linearities, we replace some ReLU activations with the identity function). In a second experiment,

we replace the ReLU activations with Leaky-ReLU activations and vary the slope of the Leaky ReLU (a higher slope means an activation function that is closer to a linear function). The results of this experiment are presented in Figure 2(a) and 2(b). In Figure 2(a), “ReLU(DC)” means that we interleave ReLU activation functions between every diagonal-circulant matrix, whereas ReLU(DCDC) means we interleave a ReLU activation every other block etc. In both Figure 2(a) and Figure 2(b), we observe that reducing the non-linearity of the networks can be used to train deeper networks. This is an interesting result, since we can use this technique to adjust the number of parameters in the network, without facing training difficulties. We obtain a maximum accuracy of 0.56 with one ReLU every three layers and leaky-ReLUs with a slope of 0.5. We hence rely on this setting in the experimental section.

6 Empirical evaluation

This experimental section aims at answering the following questions:

- Q6.1** – How do DCNNs compare to other approaches such as ACDC, LDR or other structured approaches?
- Q6.2** – How do DCNNs compare to other compression based techniques?
- Q6.3** – How do DCNNs perform in the context of large scale real-world machine learning applications?

6.1 Comparison with other structured approaches (Q6.1)

Comparison with ACDC [20]. In Section 2, we have discussed the differences between the ACDC framework and our approach from a theoretical perspective. In this section, we conduct experiments to compare the performance of DCNNs with neural networks based on ACDC layers. We first reproduce the experimental setting from [20], and compare both approaches using only linear networks (*i.e.* networks without any ReLU activations). The synthetic dataset has been created in order to reproduce the experiment on the regression linear problem proposed by [20]. We draw X and W from a uniform distribution between $[-1, +1]$ and ϵ from a normal distribution with mean 0 and variance 0.01. The relationship between X and Y is defined by $Y = XW + \epsilon$. The results are presented in Figure 3(a). On this simple setting, while both architectures demonstrate good performance, we can observe that DCNNs offer a better convergence rate. In Figure 3(b), we compare neural networks with ReLU activations on CIFAR-10.

We found that networks which are based only on ACDC layers are difficult to train and offer poor accuracy on CIFAR-10 (we have tried different initialization schemes including the one from the original paper, and the one we introduce in this paper). Moczulski et al. [20] manage to train a large VGG network however these networks are generally highly redundant and the contribution of the structured layer is difficult to quantify. We also observe that adding a single dense layer improves the convergence rate of ACDC in the linear case, which explains the good results of [20]. However, it is difficult to characterize the true contribution of the ACDC layers when the network has a large number of expressive layers.

In contrast, deep DCNNs can be trained and offer good performance without additional dense layers (these results are in line with our experiments on the *YouTube-8M* dataset). We can conclude that DCNNs are able to model complex relations at a low cost.

Comparison with Dense networks, Toeplitz networks and Low Rank networks. We now compare DCNNs with other state-of-the-art

structured networks by measuring the accuracy on a flattened version of the CIFAR-10 dataset. Our baseline is a dense feed-forward network with a fixed number of weights (9 million weights). We compare with DCNNs and with DTNNs (see below), Toeplitz networks, and Low-Rank networks [30]. We first consider Toeplitz networks which are stacked Toeplitz matrices interleaved with ReLU activations since Toeplitz matrices are closely related to circulant matrices. However, Toeplitz networks have a different structure than DCNNs (they do not include diagonal matrices), therefore, we also experiment using DTNNs, a variant of DCNNs where all the circulant matrices have been replaced by Toeplitz matrices. Finally we conduct experiments using networks based on low-rank matrices as they are also closely related to our work. For each approach, we report the accuracy of several networks with a varying depth ranging from 1 to 40 (DCNNs, Toeplitz networks) and from 1 to 30 (from DTNNs). For low-rank networks, we used a fixed depth network and increased the rank of each matrix from 7 to 40. We also tried to increase the depth of low rank matrices, but we found that deep low-rank networks are difficult to train so we do not report the results here. We compare all the networks based on the number of weights from 21K (0.2% of the dense network) to 370K weights (4% of the dense network) and we report the results in Figure 4(a). First we can see that the size of the networks correlates positively with their accuracy which demonstrate successful training in all cases. We can also see that the DCNNs achieve the maximum accuracy of 56% with 20 layers ($\sim 200K$ weights) which is as good as the dense networks with only 2% of the number of weights. Other approaches also offer good trade-off but they are not able to reach the accuracy of a dense network.

Table 1. LDR networks compared with DCNNs on a flattened version of CIFAR-10. DCNNs outperform all LDR configurations with fewer weights.²

Architectures	#Params	Acc.
<i>Dense</i>	<i>9.4M</i>	<i>0.562</i>
DCNN (5 layers)	49K	0.543
DCNN (2 layers)	21K	0.536
LDR-TD ($r = 2$)	64K	0.511
LDR-TD ($r = 3$)	70K	0.473
Toeplitz-like ($r = 2$)	46K	0.483
Toeplitz-like ($r = 3$)	52K	0.496

Table 2. Two depths scattering on CIFAR-10 followed by LDR or DC layer. Networks with DC layers outperform all LDR configurations with fewer weights.

Architectures	#Params	Acc.
DC (1 layers)	124K	0.757
DC (3 layers)	217K	0.785
Ensemble x5 DC (3 layers)	1.08M	0.811
LDR-SD ($r = 1$)	140K	0.701
LDR-SD ($r = 10$)	420K	0.728
Toeplitz-like ($r = 1$)	110K	0.711
Toeplitz-like ($r = 10$)	388K	0.720

Comparison with LDR networks [27]. We now compare DCNNs with the LDR framework using the network configuration experimented in the original paper: a single LDR structured layer followed

² Remark: the numbers may differ from the original experiments by [27] because we use the original dataset instead of a monochrome version.

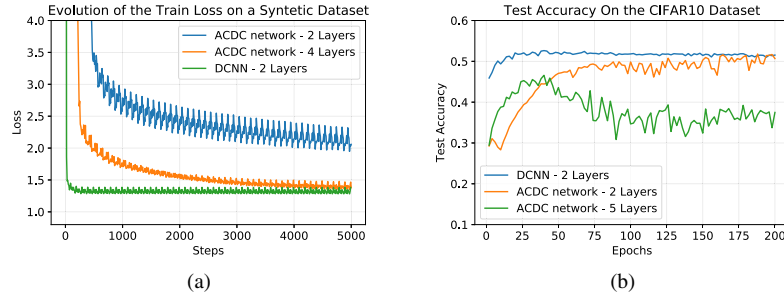


Figure 3. Comparison of DCNNs and ACDC networks on two different tasks. Figure 3(a) shows the evolution of the training loss on a regression task with synthetic data. Figure 3(b) shows the test accuracy on the CIFAR-10 dataset.

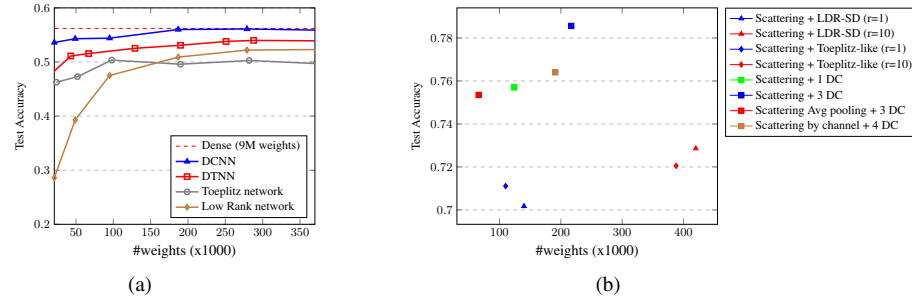


Figure 4. Figure 4(a): network size vs. accuracy compared on Dense networks, DCNNs (our approach), DTNNs (our approach), neural networks based on Toeplitz matrices and neural networks based on Low Rank-based matrices. DCNNs outperforms alternatives structured approaches. Figure 4(b) shows the accuracy of different structured architecture given the number of trainable parameters.

by a dense layer. In the LDR framework, we can change the size of a network by adjusting the rank of the residual matrix, effectively capturing matrices with a structure that is close to a known structure but not exactly (e.g. in the LDR framework, Toeplitz matrices can be encoded with a residual matrix with rank=2, so a matrix that can be encoded with a residual of rank=3 can be seen as Toeplitz-like.). The results are presented in Table 1 and demonstrate that DCNNs outperforms all LDR networks both in terms in size and accuracy.

Exploiting image features. Dense layers and DCNNs are not designed to capture task-specific features such as the translation invariance inherently useful in image classification. We can further improve the accuracy of such general purpose architectures on image classification without dramatically increasing the number of trained parameters by stacking them on top of fixed (i.e. non-trained) transforms such as the scattering transform [18]. In this section we compare the accuracy of various structured networks, enhanced with the scattering transform, on an image classification task, and run comparative experiments on CIFAR-10.

Our test architecture consists of 2 depth scattering on the RGB images followed by a batch norm and LDR or DC layer. To vary the number of parameters of Scattering+LDR architecture, we increase the rank of the matrix (stacking several LDR matrices quickly exhausted the memory). The Figure 4(b) and 2 shows the accuracy of these architectures given the number of trainable parameters.

First, we can see that the DCNN architecture very much benefits from the scattering transform and is able to reach a competitive accuracy over 78%. We can also see that scattering followed by a DC layer systematically outperforms scattering + LDR or scattering + Toeplitz-like with less parameters.

6.2 Comparison with other compression based approaches (Q6.2)

Table 3. Comparison with compression based approaches

Architecture	#Params	Error (%)
<i>LeNet</i> [16]	4 257 674	0.61
DCNN	25 620	1.74
	46 875	2.79
HashNet [6]	78 125	1.99
	46 875	6.32
Dark Knowledge [12]	78 125	2.16

We provide a comparison with other compression based approaches such as HashNet [6], Dark Knowledge [12] and Fast Food Transform (FF) [29]. Table 3 shows the test error of DCNN against other know compression techniques on the MNIST datasets. We can observe that DCNN outperform easily HashNet [6] and Dark Knowledge [12] with fewer number of parameters. The architecture with Fast Food (FF) [29] achieves better performance but with convolutional layers and only 1 Fast Food Layer as the last Softmax layer.

6.3 DCNNs for large-scale video classification on the *YouTube-8M* dataset (Q6.3)

To understand the performance of deep DCNNs on large scale applications, we conducted experiments on the *YouTube-8M* video classification with 3.8 training examples introduced by [2]. Notice that we favour this experiment over ImageNet applications because modern

image classification architectures involve a large number of convolutional layers, and compressing convolutional layers is out of our scope. Also, as mentioned earlier, testing the performance of DCNN architectures mixed with a large number of expressive layers makes little sense. The *YouTube-8M* includes two datasets describing 8 million labeled videos. Both datasets contain audio and video features for each video. In the first dataset (*aggregated*) all audio and video features have been aggregated every 300 frames. The second dataset (*full*) contains the descriptors for all the frames. To compare the models we use the GAP metric (Global Average Precision) proposed by [2]. On the simpler *aggregated* dataset we compared off-the-shelf DCNNs with a dense baseline with 5.7M weights. On the full dataset, we designed three new compact architectures based on the state-of-the-art architecture introduced by [2].

Experiments on the *aggregated* dataset with DCNNs: We compared DCNNs with a dense baseline with 5.7 millions weights. The goal of this experiment is to discover a good trade-off between depth and model accuracy. To compare the models we use the GAP metric (Global Average Precision) following the experimental protocol in [2], to compare our experiments.

Table 4 shows the results of our experiments on the *aggregated YouTube-8M* dataset in terms of number of weights, compression rate and GAP. We can see that the compression ratio offered by the circulant architectures is high. This comes at the cost of a little decrease of GAP measure. The 32 layers DCNN is 46 times smaller than the original model in terms of number of parameters while having a close performance.

Table 4. This table shows the GAP score for the *YouTube-8M* dataset with DCNNs. We can see a large increase in the score with deeper networks.

Architecture	#Weights	GAP@20
<i>original</i>	5.7M	0.773
4 DC	25 410 (0.44)	0.599
32 DC	122 178 (2.11)	0.685
4 DC + 1 FC	4.46M (77)	0.747

Table 5. This table shows the GAP score for the *YouTube-8M* dataset with different layer represented with our DC decomposition.

Architecture	#Weights	GAP@20
<i>original</i>	45M	0.846
DBoF with DC	36M (80)	0.838
FC with DC	41M (91)	0.845
MoE with DC	12M (26)	0.805

Experiments with DCNNs Deep Bag-of-Frames Architecture:

The Deep Bag-of-Frames architecture can be decomposed into three blocks of layers, as illustrated in Figure 5. The first block of layers, composed of the Deep Bag-of-Frames embedding (DBoF), is meant to model an embedding of these frames in order to make a simple representation of each video. A second block of fully connected layers (FC) reduces the dimensionality of the output of the embedding and merges the resulting output with a concatenation operation. Finally, the classification block uses a combination of Mixtures-of-Experts (MoE) [14, 1] and Context Gating [19] to calculate the final class probabilities. Table 5 shows the results in terms of number of weights, size of the model (MB) and GAP on the full dataset, replacing the DBoF block reduces the size of the network without impacting the

accuracy. We obtain the best compression ratio by replacing the MoE block with DCNNs (26%) of the size of the original dataset with a GAP score of 0.805 (95% of the score obtained with the original architecture). We conclude that DCNN are both theoretically sound and of practical interest in real, large scale applications.

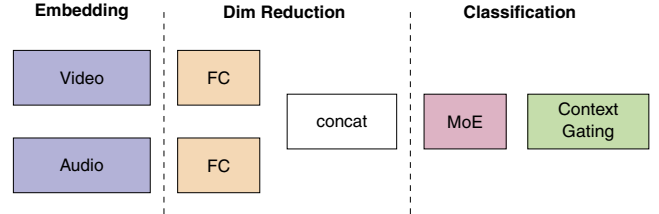


Figure 5. This figure shows the state-of-the-art neural network architecture, initially proposed by [2] and later improved by [19], used in our experiment.

Architectures & Hyper-Parameters: For the first set of our experiments (e.g. experiments on CIFAR-10), we train all networks for 200 epochs, a batch size of 200, Leaky ReLU activation with a different slope. We minimize the Cross Entropy Loss with Adam optimizer and use a piecewise constant learning rate of 5×10^{-5} , 2.5×10^{-5} , 5×10^{-6} and 1×10^{-6} after respectively 40K, 60K and 80K steps. For the *YouTube-8M* dataset experiments, we built a neural network based on the SOTA architecture initially proposed by [2] and later improved by [19]. Remark that no convolution layer is involved in this application since the input vectors are embeddings of video frames processed using state-of-the-art convolutional neural networks trained on ImageNet. We trained our models with the CrossEntropy loss and used Adam optimizer with a 0.0002 learning rate and a 0.8 exponential decay every 4 million examples. All fully connected layers are composed of 512 units. DBoF, NetVLAD and NetFV are respectively 8192, 64 and 64 of cluster size for video frames and 4096, 32, 32 for audio frames. We used 4 mixtures for the MoE Layer. We used all the available 300 frames for the DBoF embedding. In order to stabilize and accelerate the training, we used batch normalization before each non linear activation and gradient clipping.

7 Conclusion

This paper deals with the training of diagonal circulant neural networks. To the best of our knowledge, training such networks with a large number of layers had not been done before. We also endowed this kind of models with theoretical guarantees, hence enriching and refining previous theoretical work from the literature. More importantly, we showed that DCNNs outperform their competing structured alternatives, including the very recent general approach based on LDR networks. Our results suggest that stacking diagonal circulant layers with non linearities improves the convergence rate and the final accuracy of the network. Formally proving these statements constitutes the future directions of this work. We would like to generalize the good results of DCNNs to convolutional neural networks. We also believe that circulant matrices deserve a particular attention in deep learning because of their strong ties with convolutions: a circulant matrix operator is equivalent to the convolution operator with circular paddings. This fact makes any contribution to the area of circulant matrices particularly relevant to the field of deep learning with impacts beyond the problem of designing compact models. As future work, we would like to generalize our results to deep convolutional neural networks.

References

- [1] Sami Abu-El-Haija, Nisarg Kothari, Joonseok Lee, Apostol (Paul) Natsev, George Toderici, Balakrishnan Varadarajan, and Sudheendra Vijayanarasimhan, 'Youtube-8m: A large-scale video classification benchmark', in *arXiv:1609.08675*, (2016).
- [2] Sami Abu-El-Haija, Nisarg Kothari, Joonseok Lee, Paul Natsev, George Toderici, Balakrishnan Varadarajan, and Sudheendra Vijayanarasimhan, 'Youtube-8m: A large-scale video classification benchmark', *arXiv preprint arXiv:1609.08675*, (2016).
- [3] Alexandre Araujo, Benjamin Negrevert, Yann Chevalere, and Jamal Atif, 'Training compact deep learning models for video classification using circulant matrices', in *The 2nd Workshop on YouTube-8M Large-Scale Video Understanding at ECCV 2018*, (2018).
- [4] Sanjeev Arora, Nadav Cohen, Wei Hu, and Yuping Luo, 'Implicit regularization in deep matrix factorization', in *Neurips*, (05 2019).
- [5] Jimmy Ba and Rich Caruana, 'Do deep nets really need to be deep?', in *Advances in neural information processing systems*, pp. 2654–2662, (2014).
- [6] Wenlin Chen, James T. Wilson, Stephen Tyree, Kilian Q. Weinberger, and Yixin Chen, 'Compressing neural networks with the hashing trick', in *Proceedings of the 32nd International Conference on International Conference on Machine Learning - Volume 37, ICML'15*, pp. 2285–2294. JMLR.org, (2015).
- [7] Y. Cheng, F. X. Yu, R. S. Feris, S. Kumar, A. Choudhary, and S. F. Chang, 'An exploration of parameter redundancy in deep networks with circulant projections', in *2015 IEEE International Conference on Computer Vision (ICCV)*, pp. 2857–2865, (Dec 2015).
- [8] Krzysztof Choromanski, Mark Rowland, Vikas Sindhwani, Richard E. Turner, and Adrian Weller, 'Structured evolution with compact architectures for scalable policy optimization', in *ICML*, (2018).
- [9] Misha Denil, Babak Shakhb, Laurent Dinh, Marc Aurelio Ranzato, and Nando de Freitas, 'Predicting parameters in deep learning', in *Advances in Neural Information Processing Systems 26*, eds., C. J. C. Burges, L. Bottou, M. Welling, Z. Ghahramani, and K. Q. Weinberger, 2148–2156, Curran Associates, Inc., (2013).
- [10] S. Goyal, A. Roy Choudhury, and V. Sharma, 'Compression of deep neural networks by combining pruning and low rank decomposition', in *2019 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*, pp. 952–958, (2019).
- [11] Aicke Hinrichs and Jan Vybiral, 'Johnson-lindenstrauss lemma for circulant matrices', *Random Structures & Algorithms*, **39**(3), 391–398, (2011).
- [12] Geoffrey Hinton, Oriol Vinyals, and Jeffrey Dean, 'Distilling the knowledge in a neural network', in *NIPS Deep Learning and Representation Learning Workshop*, (2015).
- [13] Marko Huhtanen and Allan Perämäki, 'Factoring matrices into the product of circulant and diagonal matrices', *Journal of Fourier Analysis and Applications*, **21**(5), 1018–1033, (Oct 2015).
- [14] M. I. Jordan and R. A. Jacobs, 'Hierarchical mixtures of experts and the em algorithm', in *Proceedings of 1993 International Conference on Neural Networks (IJCNN-93-Nagoya, Japan)*, volume 2, pp. 1339–1344 vol.2, (Oct 1993).
- [15] Jakub Konečný, H. Brendan McMahan, Felix X. Yu, Peter Richtarik, Ananda Theertha Suresh, and Dave Bacon, 'Federated learning: Strategies for improving communication efficiency', in *NIPS Workshop on Private Multi-Party Machine Learning*, (2016).
- [16] Yann Lecun, Léon Bottou, Yoshua Bengio, and Patrick Haffner, 'Gradient-based learning applied to document recognition', in *Proceedings of the IEEE*, pp. 2278–2324, (1998).
- [17] Chong Li and C. J. Richard Shi, 'Constrained optimization based low-rank approximation of deep neural networks', in *Computer Vision – ECCV 2018*, eds., Vittorio Ferrari, Martial Hebert, Cristian Sminchisescu, and Yair Weiss, pp. 746–761, Cham, (2018). Springer International Publishing.
- [18] Stéphane Mallat, 'Recursive interferometric representation', in *Proc. of EUSICO conference, Denmark*, (2010).
- [19] Antoine Miech, Ivan Laptev, and Josef Sivic, 'Learnable pooling with context gating for video classification', *CoRR*, **abs/1706.06905**, (2017).
- [20] Marcin Moczulski, Misha Denil, Jeremy Appleyard, and Nando de Freitas, 'Acde: A structured efficient linear layer', *arXiv preprint arXiv:1511.05946*, (2015).
- [21] Jörn Müller-Quade, Harald Aagedal, Th Beth, and Michael Schmid, 'Algorithmic design of diffractive optical systems for information processing', *Physica D: Nonlinear Phenomena*, **120**(1-2), 196–205, (1998).
- [22] Alexander Novikov, Dmitrii Podoprikin, Anton Osokin, and Dmitry P Vetrov, 'Tensorizing neural networks', in *Advances in Neural Information Processing Systems*, pp. 442–450, (2015).
- [23] Tara Sainath and Carolina Parada, 'Convolutional neural networks for small-footprint keyword spotting', in *Interspeech*, (2015).
- [24] Victoria Sanchez, Pedro Garcia, Antonio M Peinado, José C Segura, and Antonio J Rubio, 'Diagonalizing properties of the discrete cosine transforms', *IEEE transactions on Signal Processing*, **43**(11), 2631–2641, (1995).
- [25] Michael Schmid, Rainer Steinwandt, Jörn Müller-Quade, Martin Rötteler, and Thomas Beth, 'Decomposing a matrix into circulant and diagonal factors', *Linear Algebra and its Applications*, **306**(1-3), 131–143, (2000).
- [26] Hanie Sedghi, Vineet Gupta, and Philip Long, 'The singular values of convolutional layers', in *ICLR*, (2018).
- [27] Anna Thomas, Albert Gu, Tri Dao, Atri Rudra, and Christopher Ré, 'Learning compressed transforms with low displacement rank', in *Advances in Neural Information Processing Systems 31*, eds., S. Bengio, H. Wallach, H. Larochelle, K. Grauman, N. Cesa-Bianchi, and R. Garnett, 9066–9078, Curran Associates, Inc., (2018).
- [28] Chiheb Trabelsi, Olexa Bilaniuk, Ying Zhang, Dmitriy Serdyuk, Sandeep Subramanian, João Felipe Santos, Soroush Mehri, Negar Rostamzadeh, Yoshua Bengio, and Christopher J. Pal, 'Deep complex networks', in *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*, (2018).
- [29] Z. Yang, M. Moczulski, M. Denil, N. d. Freitas, A. Smola, L. Song, and Z. Wang, 'Deep fried convnets', in *2015 IEEE International Conference on Computer Vision (ICCV)*, pp. 1476–1483, (Dec 2015).
- [30] X. Yu, T. Liu, X. Wang, and D. Tao, 'On compressing deep models by low rank and sparse decomposition', in *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 67–76, (July 2017).
- [31] Liang Zhao, Siyu Liao, Yanzhi Wang, Zhe Li, Jian Tang, and Bo Yuan, 'Theoretical properties for neural networks with weight matrices of low displacement rank', in *Proceedings of the 34th International Conference on Machine Learning*, eds., Doina Precup and Yee Whye Teh, volume 70 of *Proceedings of Machine Learning Research*, pp. 4082–4090, International Convention Centre, Sydney, Australia, (06–11 Aug 2017). PMLR.