
Implicit Regularization in Tensor Factorization

Noam Razin^{1*} Asaf Maman^{1*} Nadav Cohen¹

Abstract

Recent efforts to unravel the mystery of implicit regularization in deep learning have led to a theoretical focus on matrix factorization — matrix completion via linear neural network. As a step further towards practical deep learning, we provide the first theoretical analysis of implicit regularization in tensor factorization — tensor completion via certain type of non-linear neural network. We circumvent the notorious difficulty of tensor problems by adopting a dynamical systems perspective, and characterizing the evolution induced by gradient descent. The characterization suggests a form of greedy low tensor rank search, which we rigorously prove under certain conditions, and empirically demonstrate under others. Motivated by tensor rank capturing the implicit regularization of a non-linear neural network, we empirically explore it as a measure of complexity, and find that it captures the essence of datasets on which neural networks generalize. This leads us to believe that tensor rank may pave way to explaining both implicit regularization in deep learning, and the properties of real-world data translating this implicit regularization to generalization.

1 Introduction

The ability of neural networks to generalize when having far more learnable parameters than training examples, even in the absence of any explicit regularization, is an enigma lying at the heart of deep learning theory. Conventional wisdom is that this generalization stems from an *implicit regularization* — a tendency of gradient-based optimization to fit training examples with predictors whose “complexity” is as low as possible. The fact that “natural” data gives rise to generalization while other types of data (e.g. random) do not, is understood to result from the former being amenable

to fitting by predictors of lower complexity. A major challenge in formalizing this intuition is that we lack definitions for predictor complexity that are both quantitative (*i.e.* admit quantitative generalization bounds) and capture the essence of natural data (in the sense of it being fittable with low complexity). Consequently, existing analyses typically focus on simplistic settings, where a notion of complexity is apparent. A prominent example of such a setting is *matrix completion*.

In matrix completion, we are given a randomly chosen subset of entries from an unknown matrix $W^* \in \mathbb{R}^{d,d'}$, and our goal is to recover unseen entries. This can be viewed as a prediction problem, where the set of possible inputs is $\mathcal{X} = \{1, \dots, d\} \times \{1, \dots, d'\}$, the possible labels are $\mathcal{Y} = \mathbb{R}$, and the label of $(i, j) \in \mathcal{X}$ is $[W^*]_{i,j}$. Under this viewpoint, observed entries constitute the training set, and the average reconstruction error over unobserved entries is the test error, quantifying generalization. A predictor, *i.e.* a function from $\mathcal{X}$ to $\mathcal{Y}$, can then be seen as a matrix, and a natural notion of complexity is its rank. It is known empirically (*cf.* Gunasekar et al. (2017); Arora et al. (2019)) that this complexity measure is oftentimes implicitly minimized by *matrix factorization* — linear neural network¹ trained via gradient descent with small learning rate and near-zero initialization. Mathematically characterizing the implicit regularization in matrix factorization is a highly active area of research. Though initially conjectured to be equivalent to norm minimization (see Gunasekar et al. (2017)), recent studies (Arora et al., 2019; Razin & Cohen, 2020; Li et al., 2021) suggest that this is not the case, and instead adopt a dynamical view, ultimately establishing that (under certain conditions) the implicit regularization in matrix factorization is performing a greedy low rank search.

A central question that arises is the extent to which the study of implicit regularization in matrix factorization is relevant to more practical settings. Recent experiments (see Razin & Cohen (2020)) have shown that the tendency towards low rank extends from matrices (two-dimensional arrays) to *tensors* (multi-dimensional arrays). Namely, in the task of N -dimensional *tensor completion*, which (analogously

^{*}Equal contribution ¹Blavatnik School of Computer Science, Tel Aviv University, Israel. Correspondence to: Noam Razin <noam.raizin@cs.tau.ac.il>, Asaf Maman <asafmaman@mail.tau.ac.il>.

¹That is, parameterization of learned predictor (matrix) as a product of matrices. With such parameterization it is possible to explicitly constrain rank (by limiting shared dimensions of multiplied matrices), but the setting of interest is where rank is unconstrained, meaning all regularization is implicit.

to matrix completion) can be viewed as a prediction problem over N input variables, training a *tensor factorization*² via gradient descent with small learning rate and near-zero initialization tends to produce tensors (predictors) with low *tensor rank*. Analogously to how matrix factorization may be viewed as a linear neural network, tensor factorization can be seen as a certain type of *non-linear* neural network (two layer network with multiplicative non-linearity, cf. Cohen et al. (2016b)), and so it represents a setting much closer to practical deep learning.

In this paper we provide the first theoretical analysis of implicit regularization in tensor factorization. We circumvent the notorious difficulty of tensor problems (see Hillar & Lim (2013)) by adopting a dynamical systems perspective. Characterizing the evolution that gradient descent with small learning rate and near-zero initialization induces on the components of a factorization, we show that their norms are subject to a momentum-like effect, in the sense that they move slower when small and faster when large. This implies a form of greedy low tensor rank search, generalizing phenomena known for the case of matrices. We employ the finding to prove that, with the classic Huber loss from robust statistics (Huber, 1964), arbitrarily small initialization leads tensor factorization to follow a trajectory of rank one tensors for an arbitrary amount of time or distance. Experiments validate our analysis, demonstrating implicit regularization towards low tensor rank in a wide array of configurations.

Motivated by the fact that tensor rank captures the implicit regularization of a non-linear neural network, we empirically explore its potential to serve as a measure of complexity for multivariable predictors. We find that it is possible to fit standard image recognition datasets — MNIST (LeCun, 1998) and Fashion-MNIST (Xiao et al., 2017) — with predictors of extremely low tensor rank, far beneath what is required for fitting random data. This leads us to believe that tensor rank (or more advanced notions such as hierarchical tensor ranks) may pave way to explaining both implicit regularization of contemporary deep neural networks, and the properties of real-world data translating this implicit regularization to generalization.

The remainder of the paper is organized as follows. Section 2 presents the tensor factorization model, as well as its interpretation as a neural network. Section 3 characterizes its dynamics, followed by Section 4 which employs the characterization to establish (under certain conditions) implicit tensor rank minimization. Experiments, demonstrating both the dynamics of learning and the ability of tensor rank to capture the essence of standard datasets, are

²The term “tensor factorization” refers throughout to the classic *CP factorization*; other (more advanced) factorizations will be named differently (see Kolda & Bader (2009); Hackbusch (2012) for an introduction to various tensor factorizations).

given in Section 5. In Section 6 we review related work. Finally, Section 7 concludes. Extension of our results to tensor sensing (more general setting than tensor completion) is discussed in Appendix A.

2 Tensor Factorization

Consider the task of completing an N -dimensional tensor ($N \geq 3$) with axis lengths $d_1, \dots, d_N \in \mathbb{N}$, or, in standard tensor analysis terminology, an *order N* tensor with *modes of dimensions* $d_1, \dots, d_N$. Given a set of observations $\{y_{i_1, \dots, i_N} \in \mathbb{R}\}_{(i_1, \dots, i_N) \in \Omega}$, where Ω is a subset of all possible index tuples, a standard (undetermined) loss function for the task is:

$$\mathcal{L} : \mathbb{R}^{d_1, \dots, d_N} \rightarrow \mathbb{R}_{\geq 0} \quad (1)$$

$$\mathcal{L}(\mathcal{W}) = \frac{1}{|\Omega|} \sum_{(i_1, \dots, i_N) \in \Omega} \ell([W]_{i_1, \dots, i_N} - y_{i_1, \dots, i_N}),$$

where $\ell : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ is differentiable and locally smooth. A typical choice for $\ell(\cdot)$ is $\ell(z) = \frac{1}{2}z^2$, corresponding to ℓ_2 loss. Other options are also common, for example that given in Equation (8), which corresponds to the Huber loss from robust statistics (Huber, 1964) — a differentiable surrogate for ℓ_1 loss.

Performing tensor completion with an R -component tensor factorization amounts to optimizing the following (non-convex) objective:

$$\phi(\{\mathbf{w}_r^n\}_{r=1}^R) := \mathcal{L}(\mathcal{W}_e), \quad (2)$$

defined over *weight vectors* $\{\mathbf{w}_r^n \in \mathbb{R}^{d_n}\}_{r=1}^R$, where:

$$\mathcal{W}_e := \sum_{r=1}^R \mathbf{w}_r^1 \otimes \dots \otimes \mathbf{w}_r^N \quad (3)$$

is referred to as the *end tensor* of the factorization, with $\otimes$ representing outer product.³ The minimal number of components R required in order for $\mathcal{W}_e$ to be able to express a given tensor $\mathcal{W} \in \mathbb{R}^{d_1, \dots, d_N}$, is defined to be the *tensor rank* of $\mathcal{W}$. One may explicitly restrict the tensor rank of solutions produced by the tensor factorization via limiting R . However, since our interest lies in the implicit regularization induced by gradient descent, *i.e.* in the type of end tensors (Equation (3)) it will find when applied to the objective $\phi(\cdot)$ (Equation (2)) with no explicit constraints, we treat the case where R can be arbitrarily large.

In line with analyses of matrix factorization (*e.g.* Gunasekar et al. (2017); Arora et al. (2018; 2019); Eftekhari & Zylakis (2020); Li et al. (2021)), we model small learning rate for gradient descent through the infinitesimal limit,

³For any $\{\mathbf{w}^n \in \mathbb{R}^{d_n}\}_{n=1}^N$, the outer product $\mathbf{w}^1 \otimes \dots \otimes \mathbf{w}^N$, denoted also $\otimes_{n=1}^N \mathbf{w}^n$, is the tensor in $\mathbb{R}^{d_1, \dots, d_N}$ defined by $[\otimes_{n=1}^N \mathbf{w}^n]_{i_1, \dots, i_N} = \prod_{n=1}^N [\mathbf{w}^n]_{i_n}$.

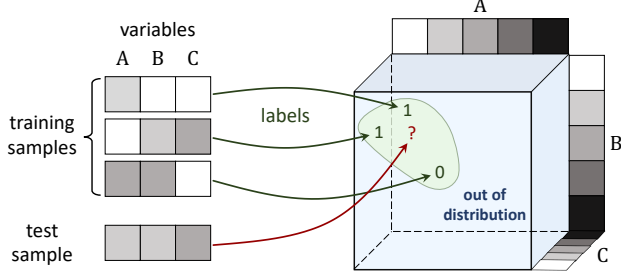


Figure 1: Prediction tasks over discrete variables can be viewed as tensor completion problems. Consider the task of learning a predictor from domain $\mathcal{X} = \{1, \dots, d_1\} \times \dots \times \{1, \dots, d_N\}$ to range $\mathcal{Y} = \mathbb{R}$ (figure assumes $N = 3$ and $d_1 = \dots = d_N = 5$ for the sake of illustration). Each input sample is associated with a location in an order N tensor with mode (axis) dimensions $d_1, \dots, d_N$, where the value of a variable (depicted as a shade of gray) determines the index of the corresponding mode (marked by “A”, “B” or “C”). The associated location stores the label of the sample. Under this viewpoint, training samples are observed entries, drawn according to an unknown distribution from a ground truth tensor. Learning a predictor amounts to completing the unobserved entries, with test error measured by (weighted) average reconstruction error. In many standard prediction tasks (e.g. image recognition), only a small subset of the input domain has non-negligible probability. From the tensor completion perspective this means that observed entries reside in a restricted part of the tensor, and reconstruction error is weighted accordingly (entries outside the support of the distribution are neglected).

i.e. through *gradient flow*:

$$\frac{d}{dt} \mathbf{w}_r^n(t) := - \frac{\partial}{\partial \mathbf{w}_r^n} \phi \left(\{ \mathbf{w}_{r'}^n(t) \}_{r'=1}^R \right) \quad (4)$$

$$, t \geq 0, r = 1, \dots, R, n = 1, \dots, N,$$

where $\{ \mathbf{w}_r^n(t) \}_{r=1}^R$ denote the weight vectors at time t of optimization.

Our aim is to theoretically investigate the prospect of implicit regularization towards low tensor rank, *i.e.* of gradient flow with near-zero initialization learning a solution that can be represented with a small number of components.

2.1 Interpretation as Neural Network

Tensor completion can be viewed as a prediction problem, where each mode corresponds to a discrete input variable. For an unknown tensor $\mathcal{W}^* \in \mathbb{R}^{d_1, \dots, d_N}$, inputs are index tuples of the form $(i_1, \dots, i_N)$, and the label associated with such an input is $[\mathcal{W}^*]_{i_1, \dots, i_N}$. Under this perspective, the training set consists of the observed entries, and the average reconstruction error over unseen entries measures test error. The standard case, in which observations are drawn uniformly across the tensor and reconstruction error weighs all entries equally, corresponds to a data distribution that is uniform, but other distributions are also viable.

Consider for example the task of predicting a continuous la-

bel for a 100-by-100 binary image. This can be formulated as an order 10000 tensor completion problem, where all modes are of dimension 2. Each input image corresponds to a location (entry) in the tensor $\mathcal{W}^*$, holding its continuous label. As image pixels are (typically) not distributed independently and uniformly, locations in the tensor are not drawn uniformly when observations are generated, and are not weighted equally when reconstruction error is computed. See Figure 1 for further illustration of how a general prediction task (with discrete inputs and scalar output) can be formulated as a tensor completion problem.

Under the above formulation, tensor factorization can be viewed as a two layer neural network with multiplicative non-linearity. Given an input, *i.e.* a location in the tensor, the network produces an output equal to the value that the factorization holds at the given location. Figure 2 illustrates this equivalence between solving tensor completion with a tensor factorization and solving a prediction problem with a non-linear neural network. A major drawback of matrix factorization as a theoretical surrogate for modern deep learning is that it misses the critical aspect of non-linearity. Tensor factorization goes beyond the realm of linear predictors — a significant step towards practical neural networks.

3 Dynamical Characterization

In this section we derive a dynamical characterization for the norms of individual components in the tensor factorization. The characterization implies that with small learning rate and near-zero initialization, components tend to be learned incrementally, giving rise to a bias towards low tensor rank solutions. This finding is used in Section 4 to prove (under certain conditions) implicit tensor rank minimization, and is demonstrated empirically in Section 5.⁴

Hereafter, unless specified otherwise, when referring to a norm we mean the standard Frobenius (Euclidean) norm, denoted by $\|\cdot\|$.

The following lemma establishes an invariant of the dynamics, showing that the differences between squared norms of vectors in the same component are constant through time.

Lemma 1. *For all $r \in \{1, \dots, R\}$ and $n, \bar{n} \in \{1, \dots, N\}$:*

$$\|\mathbf{w}_r^n(t)\|^2 - \|\mathbf{w}_r^{\bar{n}}(t)\|^2 = \|\mathbf{w}_r^n(0)\|^2 - \|\mathbf{w}_r^{\bar{n}}(0)\|^2, t \geq 0.$$

Proof sketch (for proof see Lemma 9 in Subappendix C.2.2).

The claim readily follows by showing that under gradient flow $\frac{d}{dt} \|\mathbf{w}_r^n(t)\|^2 = \frac{d}{dt} \|\mathbf{w}_r^{\bar{n}}(t)\|^2$ for all $t \geq 0$. $\square$

⁴We note that all results in this section apply even if the tensor completion loss $\mathcal{L}(\cdot)$ (Equation (1)) is replaced by any differentiable and locally smooth function. The proofs in Appendix C already account for this more general setting.

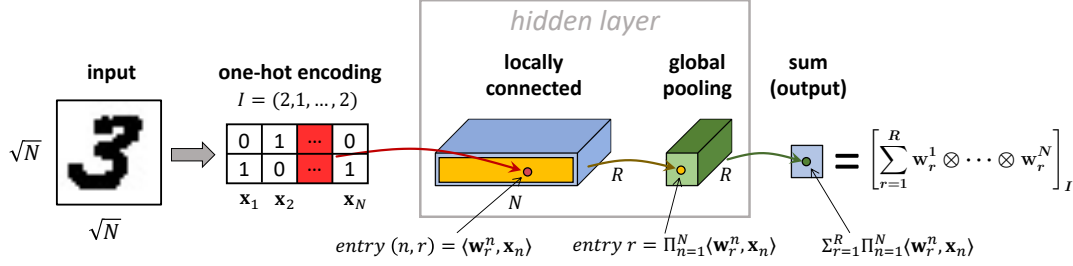


Figure 2: Tensor factorizations correspond to a class of non-linear neural networks. Figure 1 illustrates how a prediction task can be viewed as a tensor completion problem. The current figure extends this correspondence, depicting an equivalence between solving tensor completion via tensor factorization, and learning a predictor using the non-linear neural network portrayed above. The input to the network is a tuple $I = (i_1, \dots, i_N) \in \{1, \dots, d_1\} \times \dots \times \{1, \dots, d_N\}$, encoded via one-hot vectors $(\mathbf{x}_1, \dots, \mathbf{x}_N) \in \mathbb{R}^{d_1} \times \dots \times \mathbb{R}^{d_N}$. For example, in the diagram, I stands for a binary image with N pixels (in which case $d_1 = \dots = d_N = 2$). The one-hot representations are passed through a hidden layer consisting of: (i) locally connected linear operator with R channels, the r 'th one computing $\langle \mathbf{w}_r^1, \mathbf{x}_1 \rangle, \dots, \langle \mathbf{w}_r^N, \mathbf{x}_N \rangle$ with filters (learnable weights) $\{\mathbf{w}_r^n\}_{n=1}^N$; and (ii) channel-wise global product pooling (multiplicative non-linearity). The resulting activations are then reduced through summation to a scalar — the output of the network. All in all, given input tuple $I = (i_1, \dots, i_N)$, the network outputs the I 'th entry of $\sum_{r=1}^R \mathbf{w}_r^1 \otimes \dots \otimes \mathbf{w}_r^N$. Notice that the number of components R and the weight vectors $\{\mathbf{w}_r^n\}_{r,n}$ in the factorization correspond to the width and the learnable filters of the network, respectively.

Lemma 1 naturally leads to the definition below.

Definition 1. The *unbalancedness magnitude* of the weight vectors $\{\mathbf{w}_r^n \in \mathbb{R}^{d_n}\}_{r=1}^R \}_{n=1}^N$ is defined to be:

$$\max_{r \in \{1, \dots, R\}, n, \bar{n} \in \{1, \dots, N\}} \left| \|\mathbf{w}_r^n\|^2 - \|\mathbf{w}_r^{\bar{n}}\|^2 \right|.$$

By Lemma 1, the unbalancedness magnitude is constant during optimization, and thus, is determined at initialization. When weight vectors are initialized near the origin — regime of interest — the unbalancedness magnitude is small, approaching zero as initialization scale decreases.

Theorem 1 below provides a dynamical characterization for norms of individual components in the tensor factorization.

Theorem 1. Assume unbalancedness magnitude $\epsilon \geq 0$ at initialization, and denote by $\mathcal{W}_e(t)$ the end tensor (Equation (3)) at time $t \geq 0$ of optimization. Then, for any $r \in \{1, \dots, R\}$ and time $t \geq 0$ at which $\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| > 0$.⁵

- If $\gamma_r(t) := \langle -\nabla \mathcal{L}(\mathcal{W}_e(t)), \otimes_{n=1}^N \hat{\mathbf{w}}_r^n(t) \rangle \geq 0$, then:

$$\begin{aligned} \frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| &\leq N\gamma_r(t) (\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{\frac{2}{N}} + \epsilon)^{N-1} \\ \frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| &\geq N\gamma_r(t) \cdot \frac{\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^2}{\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{\frac{2}{N}} + \epsilon}, \end{aligned} \quad (5)$$

- otherwise, if $\gamma_r(t) < 0$, then:

$$\begin{aligned} \frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| &\geq N\gamma_r(t) (\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{\frac{2}{N}} + \epsilon)^{N-1} \\ \frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| &\leq N\gamma_r(t) \cdot \frac{\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^2}{\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{\frac{2}{N}} + \epsilon}, \end{aligned} \quad (6)$$

⁵When $\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|$ is zero it may not be differentiable.

where $\hat{\mathbf{w}}_r^n(t) := \mathbf{w}_r^n(t) / \|\mathbf{w}_r^n(t)\|$ for $n = 1, \dots, N$.

Proof sketch (for proof see Subappendix C.3). Differentiating a component's norm with respect to time, we obtain $\frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| = \gamma_r(t) \cdot \sum_{n=1}^N \prod_{n' \neq n} \|\mathbf{w}_r^{n'}(t)\|^2$. The desired bounds then follow from using conservation of unbalancedness magnitude (as implied by Lemma 1), and showing that $\|\mathbf{w}_r^{n'}(t)\|^2 \leq \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{2/N} + \epsilon$ for all $t \geq 0$ and $n' \in \{1, \dots, N\}$. $\square$

Theorem 1 shows that when unbalancedness magnitude at initialization (denoted ϵ) is small, the evolution rates of component norms are roughly proportional to their size exponentiated by $2 - 2/N$, where N is the order of the tensor factorization. Consequently, component norms are subject to a momentum-like effect, by which they move slower when small and faster when large. This suggests that when initialized near zero, components tend to remain close to the origin, and then, upon reaching a critical threshold, quickly grow until convergence, creating an incremental learning effect that yields implicit regularization towards low tensor rank. This phenomenon is used in Section 4 to formally prove (under certain conditions) implicit tensor rank minimization, and is demonstrated empirically in Section 5.

When the unbalancedness magnitude at initialization is exactly zero, our dynamical characterization takes on a particularly lucid form.

Corollary 1. Assume unbalancedness magnitude zero at initialization. Then, with notations of Theorem 1, for any $r \in \{1, \dots, R\}$, the norm of the r 'th component evolves by:

$$\frac{d}{dt} \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| = N\gamma_r(t) \cdot \|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{2-\frac{2}{N}}, \quad (7)$$

where by convention $\hat{\mathbf{w}}_r^n(t) = 0$ if $\mathbf{w}_r^n(t) = 0$.

Proof sketch (for proof see Subappendix C.4). If the time t is such that $\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| > 0$, Equation (7) readily follows from applying Theorem 1 with $\epsilon = 0$. For the case where $\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\| = 0$, we show that the component $\otimes_{n=1}^N \mathbf{w}_r^n(t)$ must be identically zero throughout, hence both sides of Equation (7) are equal to zero. $\square$

It is worthwhile highlighting the relation to matrix factorization. There, an implicit bias towards low rank emerges from incremental learning dynamics similar to above, with singular values standing in place of component norms. In fact, the dynamical characterization given in Corollary 1 is structurally identical to the one provided by Theorem 3 in Arora et al. (2019) for singular values of a matrix factorization. We thus obtained a generalization from matrices to tensors, notwithstanding the notorious difficulty often associated with the latter (cf. Hillar & Lim (2013)),

4 Implicit Tensor Rank Minimization

In this section we employ the dynamical characterization derived in Section 3 to theoretically establish implicit regularization towards low tensor rank. Specifically, we prove that under certain technical conditions, arbitrarily small initialization leads tensor factorization to follow a trajectory of rank one tensors for an arbitrary amount of time or distance. As a corollary, we obtain that if the tensor completion problem admits a rank one solution, and all rank one trajectories uniformly converge to it, tensor factorization with infinitesimal initialization will converge to it as well. Our analysis generalizes to tensor factorization recent results developed in Li et al. (2021) for matrix factorization. As typical in transitioning from matrices to tensors, this generalization entails significant challenges necessitating use of fundamentally different techniques.

For technical reasons, our focus in this section lies on the Huber loss from robust statistics (Huber, 1964), given by:

$$\ell_h : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}, \ell_h(z) := \begin{cases} \frac{1}{2}z^2 & , |z| < \delta_h \\ \delta_h(|z| - \frac{1}{2}\delta_h) & , \text{otherwise} \end{cases}, \quad (8)$$

where $\delta_h > 0$, referred to as the transition point of the loss, is predetermined. Huber loss is often used as a differentiable surrogate for ℓ_1 loss, in which case δ_h is chosen to be small. We will assume it is smaller than observed tensor entries:⁶

Assumption 1. $\delta_h < |y_{i_1, \dots, i_N}|, \forall (i_1, \dots, i_N) \in \Omega$.

We will consider an initialization $\{\mathbf{a}_r^n \in \mathbb{R}^{d_n}\}_{r=1}^R \}_{n=1}^N$ for the weight vectors of the tensor factorization, and will scale this initialization towards zero. In line with infinitesimal initializations being captured by unbalancedness magnitude zero (cf. Section 3), we assume that this is the case:

⁶Note that this entails assumption of non-zero observations.

Assumption 2. The initialization $\{\mathbf{a}_r^n\}_{r=1}^R \}_{n=1}^N$ has unbalancedness magnitude zero.

We further assume that within $\{\mathbf{a}_r^n\}_{r,n}$ there exists a leading component (subset $\{\mathbf{a}_{\bar{r}}^n\}_n$), in the sense that it is larger than others, while having positive projection on the attracting force at the origin, i.e. on minus the gradient of the loss $\mathcal{L}(\cdot)$ (Equation (1)) at zero:

Assumption 3. There exists $\bar{r} \in \{1, \dots, R\}$ such that:

$$\begin{aligned} \langle -\nabla \mathcal{L}(0), \otimes_{n=1}^N \hat{\mathbf{a}}_{\bar{r}}^n \rangle &> 0, \\ \|\mathbf{a}_{\bar{r}}^n\| &> \|\mathbf{a}_r^n\| \cdot \left(\frac{\|\nabla \mathcal{L}(0)\|}{\langle -\nabla \mathcal{L}(0), \otimes_{n=1}^N \hat{\mathbf{a}}_{\bar{r}}^n \rangle} \right)^{1/(N-2)}, \forall r \neq \bar{r}, \end{aligned} \quad (9)$$

where $\hat{\mathbf{a}}_{\bar{r}}^n := \mathbf{a}_{\bar{r}}^n / \|\mathbf{a}_{\bar{r}}^n\|$ for $n = 1, \dots, N$.

Let $\alpha > 0$, and suppose we run gradient flow on the tensor factorization (see Section 2) starting from the initialization $\{\mathbf{a}_r^n\}_{r,n}$ scaled by α . That is, we set:

$$\mathbf{w}_r^n(0) = \alpha \cdot \mathbf{a}_r^n, \quad r = 1, \dots, R, n = 1, \dots, N,$$

and let $\{\mathbf{w}_r^n(t)\}_{r,n}$ evolve per Equation (4). Denote by $\mathcal{W}_e(t)$, $t \geq 0$, the trajectory induced on the end tensor (Equation (3)). We will study the evolution of this trajectory through time. A hurdle that immediately arises is that, by the dynamical characterization of Section 3, when the initialization scale α tends to zero (regime of interest), the time it takes $\mathcal{W}_e(t)$ to escape the origin grows to infinity.⁷ We overcome this hurdle by considering a *reference sphere* — a sphere around the origin with sufficiently small radius:

$$\mathcal{S} := \{\mathcal{W} \in \mathbb{R}^{d_1, \dots, d_N} : \|\mathcal{W}\| = \rho\}, \quad (10)$$

where $\rho \in (0, \min_{(i_1, \dots, i_N) \in \Omega} |y_{i_1, \dots, i_N}| - \delta_h)$ can be chosen arbitrarily. With the reference sphere $\mathcal{S}$ at hand, we define a time-shifted version of the trajectory $\mathcal{W}_e(t)$, aligning $t = 0$ with the moment at which $\mathcal{S}$ is reached:

$$\overline{\mathcal{W}}_e(t) := \mathcal{W}_e(t + \inf\{t' \geq 0 : \mathcal{W}_e(t') \in \mathcal{S}\}), \quad (11)$$

where by definition $\inf\{t' \geq 0 : \mathcal{W}_e(t') \in \mathcal{S}\} = 0$ if $\mathcal{W}_e(t)$ does not reach $\mathcal{S}$. Unlike the original trajectory $\mathcal{W}_e(t)$, the shifted one $\overline{\mathcal{W}}_e(t)$ disregards the process of escaping the origin, and thus admits a concrete meaning to the time elapsing from optimization commencement.

We will establish proximity of $\overline{\mathcal{W}}_e(t)$ to trajectories of rank one tensors. We say that $\mathcal{W}_1(t) \in \mathbb{R}^{d_1, \dots, d_N}$, $t \geq 0$, is a *rank one trajectory*, if it coincides with some trajectory

⁷To see this, divide both sides of Equation (7) from Corollary 1 by $\|\otimes_{n=1}^N \mathbf{w}_r^n(t)\|^{2-2/N}$, and integrate with respect to t . It follows that the norm of a component at any fixed time tends to zero as initialization scale α decreases. This implies that for any $D > 0$, when taking $\alpha \rightarrow 0$, the time required for a component to reach norm D grows to infinity.

of an end tensor in a one-component factorization, *i.e.* if there exists an initialization for gradient flow over a tensor factorization with $R = 1$ components, leading the induced end tensor to evolve by $\mathcal{W}_1(t)$. If the latter initialization has unbalancedness magnitude zero (*cf.* Definition 1), we further say that $\mathcal{W}_1(t)$ is a *balanced rank one trajectory*.⁸

We are now in a position to state our main result, by which arbitrarily small initialization leads tensor factorization to follow a (balanced) rank one trajectory for an arbitrary amount of time or distance.

Theorem 2. *Under Assumptions 1, 2 and 3, for any distance from origin $D > 0$, time duration $T > 0$, and degree of approximation $\epsilon \in (0, 1)$, if initialization scale α is sufficiently small,⁹ then: (i) $\mathcal{W}_e(t)$ reaches the reference sphere $\mathcal{S}$; and (ii) there exists a balanced rank one trajectory $\mathcal{W}_1(t)$ emanating from $\mathcal{S}$, such that $\|\bar{\mathcal{W}}_e(t) - \mathcal{W}_1(t)\| \leq \epsilon$ at least until $t \geq T$ or $\|\bar{\mathcal{W}}_e(t)\| \geq D$.*

Proof sketch (for proof see Subappendix C.5). Using the dynamical characterization from Section 3 (Lemma 1 and Corollary 1), and the fact that $\nabla \mathcal{L}(\cdot)$ is locally constant around the origin, we establish that (i) $\mathcal{W}_e(t)$ reaches the reference sphere $\mathcal{S}$; and (ii) at that time, the norm of the $\bar{r}$ 'th component is of constant scale (independent of α), while the norms of all other components are $\mathcal{O}(\alpha^N)$. Thus, taking α towards zero leads $\mathcal{W}_e(t)$ to arrive at $\mathcal{S}$ while being arbitrarily close to the initialization of a balanced rank one trajectory — $\mathcal{W}_1(t)$. Since the objective is locally smooth, this ensures $\bar{\mathcal{W}}_e(t)$ is within distance ϵ from $\mathcal{W}_1(t)$ for an arbitrary amount of time or distance. That is, if α is sufficiently small, $\|\bar{\mathcal{W}}_e(t) - \mathcal{W}_1(t)\| \leq \epsilon$ at least until $t \geq T$ or $\|\bar{\mathcal{W}}_e(t)\| \geq D$. $\square$

As an immediate corollary of Theorem 2, we obtain that if all balanced rank one trajectories uniformly converge to a global minimum, tensor factorization with infinitesimal initialization will do so too. In particular, its implicit regularization will direct it towards a solution with tensor rank one.

Corollary 2. *Assume the conditions of Theorem 2 (Assumptions 1, 2 and 3), and in addition, that all balanced rank one trajectories emanating from $\mathcal{S}$ converge to a tensor $\mathcal{W}^* \in \mathbb{R}^{d_1, \dots, d_N}$ uniformly, in the sense that they are all confined to some bounded domain, and for any $\epsilon > 0$, there exists a time T after which they are all within distance ϵ from $\mathcal{W}^*$. Then, for any $\epsilon > 0$, if initialization scale α is sufficiently small, there exists a time T for which $\|\mathcal{W}_e(T) - \mathcal{W}^*\| \leq \epsilon$.*

⁸Note that the definitions of rank one trajectory and balanced rank one trajectory allow for $\mathcal{W}_1(t)$ to have rank zero (*i.e.* to be equal to zero) at some or all times $t \geq 0$.

⁹Hiding problem-dependent constants, an initialization scale of $\epsilon D^{-1} \exp(-\mathcal{O}(D^2 T))$ suffices. Exact constants are specified at the beginning of the proof in Subappendix C.5.

Proof sketch (for proof see Subappendix C.6). Let $T' > 0$ be a time at which all balanced rank one trajectories that emanated from $\mathcal{S}$ are within distance $\epsilon/2$ from $\mathcal{W}^*$. By Theorem 2, if α is sufficiently small, $\bar{\mathcal{W}}_e(t)$ is guaranteed to be within distance $\epsilon/2$ from a balanced rank one trajectory that emanated from $\mathcal{S}$, at least until time T' . Recalling that $\bar{\mathcal{W}}_e(t)$ is a time-shifted version of $\mathcal{W}_e(t)$, the desired result follows from the triangle inequality. $\square$

5 Experiments

In this section we present our experiments. Subsection 5.1 corroborates our theoretical analyses (Sections 3 and 4), evaluating tensor factorization (Section 2) on synthetic low (tensor) rank tensor completion problems. Subsection 5.2 explores tensor rank as a measure of complexity, examining its ability to capture the essence of standard datasets. For brevity, we defer a description of implementation details, as well as some experiments, to Appendix B.

5.1 Dynamics of Learning

Recently, Razin & Cohen (2020) empirically showed that, with small learning rate and near-zero initialization, gradient descent over tensor factorization exhibits an implicit regularization towards low tensor rank. Our theory (Sections 3 and 4) explains this implicit regularization through a dynamical analysis — we prove that the movement of component norms is attenuated when small and enhanced when large, thus creating an incremental learning effect which becomes more potent as initialization scale decreases. Figure 3 demonstrates this phenomenon empirically on synthetic low (tensor) rank tensor completion problems. Figures 5, 6 and 7 in Subappendix B.1 extend the experiment, corroborating our analyses in a wide array of settings.

5.2 Tensor Rank as Measure of Complexity

Implicit regularization in deep learning is typically viewed as a tendency of gradient-based optimization to fit training examples with predictors whose “complexity” is as low as possible. The fact that “natural” data gives rise to generalization while other types of data (*e.g.* random) do not, is understood to result from the former being amenable to fitting by predictors of lower complexity. A major challenge in formalizing this intuition is that we lack definitions for predictor complexity that are both quantitative (*i.e.* admit quantitative generalization bounds) and capture the essence of natural data (types of data on which neural networks generalize in practice), in the sense of it being fittable with low complexity.

As discussed in Subsection 2.1, learning a predictor with multiple discrete input variables and a continuous output can be viewed as a tensor completion problem. Specifically,

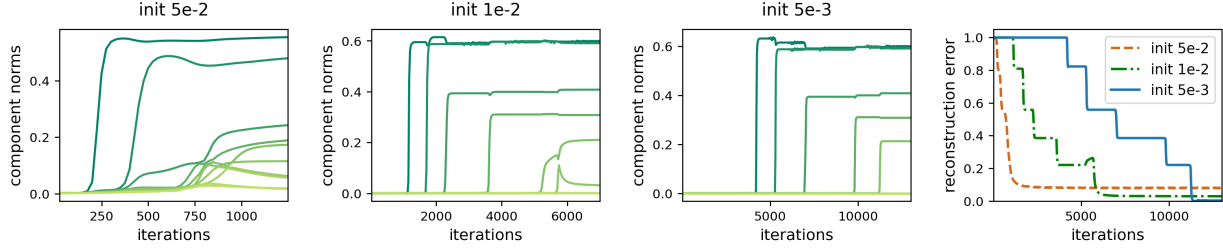


Figure 3: Dynamics of gradient descent over tensor factorization — incremental learning of components yields low tensor rank solutions. Presented plots correspond to the task of completing a (tensor) rank 5 ground truth tensor of size 10-by-10-by-10-by-10 (order 4) based on 2000 observed entries chosen uniformly at random without repetition (smaller sample sizes led to solutions with tensor rank lower than that of the ground truth tensor). In each experiment, the ℓ_2 loss (more precisely, Equation (1) with $\ell(z) := z^2$) was minimized via gradient descent over a tensor factorization with $R = 1000$ components (large enough to express any tensor), starting from (small) random initialization. First (left) three plots show (Frobenius) norms of the ten largest components under three standard deviations for initialization — 0.05, 0.01, and 0.005. Further reduction of initialization scale yielded no noticeable change. The rightmost plot compares reconstruction errors (Frobenius distance from ground truth) from the three runs. To facilitate more efficient experimentation, we employed an adaptive learning rate scheme (see Subappendix B.2 for details). Notice that, in accordance with the theoretical analysis of Section 3, component norms move slower when small and faster when large, creating an incremental process in which components are learned one after the other. This effect is enhanced as initialization scale is decreased, producing low tensor rank solutions that accurately reconstruct the low (tensor) rank ground truth tensor. In particular, even though the factorization consists of 1000 components, when initialization is sufficiently small, only five (tensor rank of the ground truth tensor) substantially depart from zero. Appendix B provides further implementation details, as well as similar experiments with: (i) Huber loss (see Equation (8)) instead of ℓ_2 loss; (ii) ground truth tensors of different orders and (tensor) ranks; and (iii) tensor sensing (see Appendix A).

with $N \in \mathbb{N}$, $d_1, \dots, d_N \in \mathbb{N}$, learning a predictor from domain $\mathcal{X} = \{1, \dots, d_1\} \times \dots \times \{1, \dots, d_N\}$ to range $\mathcal{Y} = \mathbb{R}$ corresponds to completion of an order N tensor with mode (axis) dimensions $d_1, \dots, d_N$. Under this correspondence, any predictor can simply be thought of as a tensor, and vice versa. We have shown that solving tensor completion via tensor factorization amounts to learning a predictor through a certain neural network (Subsection 2.1), whose implicit regularization favors solutions with low tensor rank (Sections 3 and 4). Motivated by these connections, the current subsection empirically explores tensor rank as a measure of complexity for predictors, by evaluating the extent to which it captures natural data, *i.e.* allows the latter to be fit with low complexity predictors.

As representatives of natural data, we chose the classic MNIST dataset (LeCun, 1998) — perhaps the most common benchmark for demonstrating ideas in deep learning — and its more modern counterpart Fashion-MNIST (Xiao et al., 2017). A hurdle posed by these datasets is that they involve classification into multiple categories, whereas the equivalence to tensors applies to predictors whose output is a scalar. It is possible to extend the equivalence by equating a multi-output predictor with multiple tensors, in which case the predictor is associated with multiple tensor ranks. However, to facilitate a simple presentation, we avoid this extension and simply map each dataset into multiple one-vs-all binary classification problems. For each problem, we associate the label 1 with the active category and 0 with all the rest, and then attempt to fit training examples with predictors of low tensor rank, reporting the resulting mean squared error, *i.e.* the residual of the fit. This is compared

against residuals obtained when fitting two types of random data: one generated via shuffling labels, and the other by replacing inputs with noise.

Both MNIST and Fashion-MNIST comprise 28-by-28 grayscale images, with each pixel taking one of 256 possible values. Tensors associated with predictors are thus of order 784, with dimension 256 in each mode (axis).¹⁰ A general rank one tensor can then be expressed as an outer product between 784 vectors of dimension 256 each, and accordingly has roughly $784 \cdot 256$ degrees of freedom. This significantly exceeds the number of training examples in the datasets (60000), hence it is no surprise that we could easily fit them, as well as their random variants, with a predictor whose tensor rank is one. To account for the comparatively small training sets, and render their fit more challenging, we quantized pixels to hold one of two values, *i.e.* we reduced images from grayscale to black and white. Following the quantization, tensors associated with predictors have dimension two in each mode, and the number of degrees of freedom in a general rank one tensor is roughly $784 \cdot 2$ — well below the number of training examples. We may thus expect to see a difference between the tensor ranks needed for fitting original datasets and those required by the random

¹⁰In practice, when associating predictors with tensors, it is often beneficial to modify the representation of the input (*cf.* Cohen et al. (2016b)). For example, in the context under discussion, rather than having the discrete input variables hold pixel intensities, they may correspond to small image patches, where each patch is represented by the index of a centroid it is closest to, with centroids determined via clustering applied to all patches across all images in the dataset. For simplicity, we did not transform representations in our experiments, and simply operated over raw image pixels.



Figure 4: Evaluation of tensor rank as measure of complexity — standard datasets can be fit accurately with predictors of extremely low tensor rank (far beneath what is required by random datasets), suggesting it may capture the essence of natural data. Left and right plots show results of fitting MNIST and Fashion-MNIST datasets, respectively, with predictors of increasing tensor rank. Original datasets are compared against two random variants: one generated by replacing images with noise (“rand image”), and the other via shuffling labels (“rand label”). As described in the text (Subsection 5.2), for simplicity of presentation, each dataset was mapped into multiple (ten) one-vs-all prediction tasks (label 1 for active category, 0 for the rest), with fit measured via mean squared error. Separately for each one-vs-all prediction task and each value $k \in \{1, \dots, 15\}$ for the tensor rank, we applied an approximate numerical method (see Subappendix B.2.2 for details) to find the predictor of tensor rank k (or less) with which the mean squared error over training examples is minimal. We report this mean squared error, as well as that obtained by the predictor on the test set (to mitigate impact of outliers, large squared errors over test samples were clipped — see Subappendix B.2.2 for details). Plots show, for each value of k , mean (as marker) and standard deviation (as error bar) of these errors taken over the different one-vs-all prediction tasks. Notice that the original datasets are fit accurately (low train error) by predictors of tensor rank as low as one, whereas random datasets are not (with tensor rank one, residuals of their fit are close to trivial, *i.e.* to the variance of the label). This suggests that tensor rank as a measure of complexity for predictors has potential to capture the essence of natural data. Notice also that, as expected, accurate fit with low tensor rank coincides with accurate prediction on test set, *i.e.* with generalization. For further details, as well as an experiment showing that linear predictors are incapable of accurately fitting the datasets, see Appendix B.

ones. This is confirmed by Figure 4, displaying the results of the experiment.

Figure 4 shows that with predictors of low tensor rank, MNIST and Fashion-MNIST can be fit much more accurately than the random datasets. Moreover, as one would presume, accurate fit with low tensor rank coincides with accurate prediction on unseen data (test set), *i.e.* with generalization. Combined with the rest of our results, we interpret this finding as an indication that tensor rank may shed light on both implicit regularization of neural networks, and the properties of real-world data translating this implicit regularization to generalization.

6 Related Work

Theoretical analysis of implicit regularization induced by gradient-based optimization in deep learning is a highly active area of research. Works along this line typically focus on simplified settings, delivering results such as: characterizations of dynamical or statistical aspects of learning (Du et al., 2018; Gidel et al., 2019; Arora et al., 2019; Brutzkus & Globerson, 2020; Gissin et al., 2020; Chou et al., 2020); solutions for test error when data distribution is known (Advani & Saxe, 2017; Goldt et al., 2019; Lampinen & Ganguli, 2019); and proofs of complexity measures being implicitly minimized in certain situations, either exactly or approximately.¹¹ The latter type of results is perhaps

the most common, covering complexity measures based on: frequency content of input-output mapping (Rahaman et al., 2019; Xu, 2018); curvature of training objective (Mullayoff & Michaeli, 2020); and norm or margin of weights or input-output mapping (Soudry et al., 2018; Gunasekar et al., 2018a,b; Jacot et al., 2018; Ji & Telgarsky, 2019b; Mei et al., 2019; Wu et al., 2020; Naason et al., 2019; Ji & Telgarsky, 2019a; Oymak & Soltanolkotabi, 2019; Ali et al., 2020; Woodworth et al., 2020; Chizat & Bach, 2020; Yun et al., 2021). An additional complexity measure, arguably the most extensively studied, is matrix rank.

Rank minimization in matrix completion (or sensing) is a classic problem in science and engineering (*cf.* Davenport & Romberg (2016)). It relates to deep learning when solved via linear neural network, *i.e.* through matrix factorization. The literature on matrix factorization for rank minimization is far too broad to cover here — we refer to Chi et al. (2019) for a recent review. Notable works proving rank minimization via matrix factorization trained by gradient descent with no explicit regularization are Tu et al. (2016); Ma et al. (2018); Li et al. (2018). Gunasekar et al. (2017) conjectured that this implicit regularization is equivalent to norm minimization, but the recent studies Arora et al. (2019); Razin & Cohen (2020); Li et al. (2021) argue otherwise, and instead adopt a dynamical view, ultimately establishing that (under certain conditions) the implicit regularization in matrix factorization is performing a greedy low rank search. These studies are relevant to ours in the sense that we generalize some of their results to tensor factorization. As typical in transitioning from matrices to tensors (see Hillar & Lim

¹¹Recent results of Vardi & Shamir (2021) imply that under certain conditions, implicit minimization of a complexity measure must be approximate (cannot be exact).

(2013)), this generalization entails significant challenges necessitating use of fundamentally different techniques.

Recovery of low (tensor) rank tensors from incomplete observations via tensor factorizations is a setting of growing interest (*cf.* Acar et al. (2011); Narita et al. (2012); Anandkumar et al. (2014); Jain & Oh (2014); Yokota et al. (2016); Karlsson et al. (2016); Xia & Yuan (2017); Zhou et al. (2017); Cai et al. (2019) and the survey Song et al. (2019)).¹² However, the experiments of Razin & Cohen (2020) comprise the only evidence we are aware of for successful recovery under gradient-based optimization with no explicit regularization (in particular without imposing low tensor rank through a factorization).¹³ The current paper provides the first theoretical support for this implicit regularization.

We note that the equivalence between tensor factorizations and different types of neural networks has been studied extensively, primarily in the context of expressive power (see, *e.g.*, Cohen et al. (2016b); Cohen & Shashua (2016); Sharir et al. (2016); Cohen & Shashua (2017); Cohen et al. (2017); Sharir & Shashua (2018); Levine et al. (2018b); Cohen et al. (2018); Levine et al. (2018a); Balda et al. (2018); Khrulkov et al. (2018); Levine et al. (2019); Khrulkov et al. (2019); Levine et al. (2020)). Connections between tensor analysis and generalization in deep learning have also been made (*cf.* Li et al. (2020)), but to the best of our knowledge, the notion of quantifying the complexity of predictors through their tensor rank (supported empirically in Subsection 5.2) is novel to this work.

7 Conclusion

In this paper we provided the first theoretical analysis of implicit regularization in tensor factorization. To circumvent the notorious difficulty of tensor problems (see Hillar & Lim (2013)), we adopted a dynamical systems perspective, and characterized the evolution that gradient descent (with small learning rate and near-zero initialization) induces on the components of a factorization. The characterization suggests a form of greedy low tensor rank search, rigorously proven under certain conditions. Experiments demonstrated said phenomena.

A major challenge in mathematically explaining generalization in deep learning is to define measures for predictor complexity that are both quantitative (*i.e.* admit quantitative generalization bounds) and capture the essence of “natural”

data (types of data on which neural networks generalize in practice), in the sense of it being fittable with low complexity. Motivated by the fact that tensor factorization is equivalent to a certain non-linear neural network, and by our analysis implying that the implicit regularization of this network minimizes tensor rank, we empirically explored the potential of the latter to serve as a measure of predictor complexity. We found that it is possible to fit standard image recognition datasets (MNIST and Fashion-MNIST) with predictors of extremely low tensor rank (far beneath what is required for fitting random data), suggesting that it indeed captures aspects of natural data.

The neural network to which tensor factorization is equivalent entails multiplicative non-linearity. It was shown in Cohen & Shashua (2016) that more prevalent non-linearities, for example rectified linear unit (ReLU), can be accounted for by considering *generalized tensor factorizations*. Studying the implicit regularization in generalized tensor factorizations (both empirically and theoretically) is regarded as a promising direction for future work.

There are two drawbacks to tensor factorization when applied to high-dimensional prediction problems. The first is technical, and relates to numerical stability — an order N tensor factorization involves products of N numbers, thus is susceptible to arithmetic underflow or overflow if N is large. Care should be taken to avoid this pitfall, for example by performing computations in log domain (as done in Cohen & Shashua (2014); Cohen et al. (2016a); Sharir et al. (2016)). The second limitation is more fundamental, arising from the fact that tensor rank — the complexity measure implicitly minimized — is oblivious to the ordering of tensor modes (axes). This means that the implicit regularization does not take into account how predictor inputs are arranged (*e.g.*, in the context of image recognition, it does not take into account spatial relationships between pixels). A potentially promising path for overcoming this limitation is introduction of *hierarchy* into the tensor factorization, equivalent to adding depth to the corresponding neural network (*cf.* Cohen et al. (2016b)). It may then be the case that a *hierarchical tensor rank* (see Hackbusch (2012)), which does account for mode ordering, will be implicitly minimized. We hypothesize that hierarchical tensor ranks may be key to explaining both implicit regularization of contemporary deep neural networks, and the properties of real-world data translating this implicit regularization to generalization.

Acknowledgements

This work was supported by a Google Research Scholar Award, a Google Research Gift, the Yandex Initiative in Machine Learning, Len Blavatnik and the Blavatnik Family Foundation, and Amnon and Anat Shashua.

¹²It stands in contrast to inferring representations for fully observed low (tensor) rank tensors via tensor factorizations (*cf.* Wang et al. (2020)) — a setting where implicit regularization (as conventionally defined in deep learning) is not applicable.

¹³In a work parallel to ours, Milanese et al. (2021) provides further empirical evidence for such implicit regularization.

References

- Acar, E., Dunlavy, D. M., Kolda, T. G., and Mørup, M. Scalable tensor factorizations for incomplete data. *Chemometrics and Intelligent Laboratory Systems*, 106(1):41–56, 2011.
- Advani, M. S. and Saxe, A. M. High-dimensional dynamics of generalization error in neural networks. *arXiv preprint arXiv:1710.03667*, 2017.
- Ali, A., Dobriban, E., and Tibshirani, R. J. The implicit regularization of stochastic gradient flow for least squares. In *International Conference on Machine Learning (ICML)*, 2020.
- Anandkumar, A., Ge, R., Hsu, D., Kakade, S. M., and Telgarsky, M. Tensor decompositions for learning latent variable models. *Journal of Machine Learning Research*, 15:2773–2832, 2014.
- Arora, S., Cohen, N., and Hazan, E. On the optimization of deep networks: Implicit acceleration by overparameterization. In *International Conference on Machine Learning (ICML)*, pp. 244–253, 2018.
- Arora, S., Cohen, N., Hu, W., and Luo, Y. Implicit regularization in deep matrix factorization. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 7413–7424, 2019.
- Balda, E. R., Behboodi, A., and Mathar, R. A tensor analysis on dense connectivity via convolutional arithmetic circuits. *Preprint*, 2018.
- Brutzkus, A. and Globerson, A. On the inductive bias of a cnn for orthogonal patterns distributions. *arXiv preprint arXiv:2002.09781*, 2020.
- Cai, C., Li, G., Poor, H. V., and Chen, Y. Nonconvex low-rank tensor completion from noisy data. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 1863–1874, 2019.
- Chi, Y., Lu, Y. M., and Chen, Y. Nonconvex optimization meets low-rank matrix factorization: An overview. *IEEE Transactions on Signal Processing*, 67(20):5239–5269, 2019.
- Chizat, L. and Bach, F. Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss. In *Conference on Learning Theory (COLT)*, pp. 1305–1338, 2020.
- Chou, H.-H., Gieshoff, C., Maly, J., and Rauhut, H. Gradient descent for deep matrix factorization: Dynamics and implicit bias towards low rank. *arXiv preprint arXiv:2011.13772*, 2020.
- Cohen, N. and Shashua, A. Simnets: A generalization of convolutional networks. *Advances in Neural Information Processing Systems (NeurIPS), Deep Learning Workshop*, 2014.
- Cohen, N. and Shashua, A. Convolutional rectifier networks as generalized tensor decompositions. *International Conference on Machine Learning (ICML)*, 2016.
- Cohen, N. and Shashua, A. Inductive bias of deep convolutional networks through pooling geometry. *International Conference on Learning Representations (ICLR)*, 2017.
- Cohen, N., Sharir, O., and Shashua, A. Deep simnets. *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016a.
- Cohen, N., Sharir, O., and Shashua, A. On the expressive power of deep learning: A tensor analysis. *Conference On Learning Theory (COLT)*, 2016b.
- Cohen, N., Sharir, O., Levine, Y., Tamari, R., Yakira, D., and Shashua, A. Analysis and design of convolutional networks via hierarchical tensor decompositions. *Intel Collaborative Research Institute for Computational Intelligence (ICRI-CI) Special Issue on Deep Learning Theory*, 2017.
- Cohen, N., Tamari, R., and Shashua, A. Boosting dilated convolutional networks with mixed tensor decompositions. *International Conference on Learning Representations (ICLR)*, 2018.
- Davenport, M. A. and Romberg, J. An overview of low-rank matrix recovery from incomplete observations. *IEEE Journal of Selected Topics in Signal Processing*, 10(4):608–622, 2016.
- Du, S. S., Hu, W., and Lee, J. D. Algorithmic regularization in learning deep homogeneous models: Layers are automatically balanced. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 384–395, 2018.
- Eftekhari, A. and Zygalakis, K. Implicit regularization in matrix sensing: A geometric view leads to stronger results. *arXiv preprint arXiv:2008.12091*, 2020.
- Gidel, G., Bach, F., and Lacoste-Julien, S. Implicit regularization of discrete gradient dynamics in linear neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 3196–3206, 2019.
- Gissin, D., Shalev-Shwartz, S., and Daniely, A. The implicit bias of depth: How incremental learning drives generalization. *International Conference on Learning Representations (ICLR)*, 2020.
- Goldt, S., Advani, M., Saxe, A. M., Krzakala, F., and Zdeborová, L. Dynamics of stochastic gradient descent for two-layer neural networks in the teacher-student setup. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 6979–6989, 2019.
- Gunasekar, S., Woodworth, B. E., Bhojanapalli, S., Neyshabur, B., and Srebro, N. Implicit regularization in matrix factorization. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 6151–6159, 2017.
- Gunasekar, S., Lee, J., Soudry, D., and Srebro, N. Characterizing implicit bias in terms of optimization geometry. In *Proceedings of the 35th International Conference on Machine Learning (ICML)*, volume 80, pp. 1832–1841, 2018a.
- Gunasekar, S., Lee, J. D., Soudry, D., and Srebro, N. Implicit bias of gradient descent on linear convolutional networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 9461–9471, 2018b.
- Hackbusch, W. *Tensor spaces and numerical tensor calculus*, volume 42. Springer, 2012.
- Hillar, C. J. and Lim, L.-H. Most tensor problems are np-hard. *Journal of the ACM (JACM)*, 60(6):1–39, 2013.
- Huber, P. J. Robust estimation of a location parameter. *Annals of Mathematical Statistics*, 35:73–101, 1964.

- Ibrahim, S., Fu, X., and Li, X. On recoverability of randomly compressed tensors with low cp rank. *IEEE Signal Processing Letters*, 27:1125–1129, 2020.
- Jacot, A., Gabriel, F., and Hongler, C. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems (NeurIPS)*, pp. 8571–8580, 2018.
- Jain, P. and Oh, S. Provable tensor factorization with missing data. In *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 1431–1439, 2014.
- Ji, Z. and Telgarsky, M. Gradient descent aligns the layers of deep linear networks. *International Conference on Learning Representations (ICLR)*, 2019a.
- Ji, Z. and Telgarsky, M. The implicit bias of gradient descent on nonseparable data. In *Conference on Learning Theory (COLT)*, pp. 1772–1798, 2019b.
- Karlsson, L., Kressner, D., and Uschmajew, A. Parallel algorithms for tensor completion in the cp format. *Parallel Computing*, 57: 222–234, 2016.
- Khrulkov, V., Novikov, A., and Oseledets, I. Expressive power of recurrent neural networks. *International Conference on Learning Representations (ICLR)*, 2018.
- Khrulkov, V., Hrinchuk, O., and Oseledets, I. Generalized tensor models for recurrent neural networks. *International Conference on Learning Representations (ICLR)*, 2019.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Kolda, T. G. and Bader, B. W. Tensor decompositions and applications. *SIAM review*, 51(3):455–500, 2009.
- Lampinen, A. K. and Ganguli, S. An analytic theory of generalization dynamics and transfer learning in deep linear networks. *International Conference on Learning Representations (ICLR)*, 2019.
- LeCun, Y. The mnist database of handwritten digits. <http://yann.lecun.com/exdb/mnist/>, 1998.
- Levine, Y., Sharir, O., and Shashua, A. Benefits of depth for long-term memory of recurrent networks. *International Conference on Learning Representations (ICLR) Workshop*, 2018a.
- Levine, Y., Yakira, D., Cohen, N., and Shashua, A. Deep learning and quantum entanglement: Fundamental connections with implications to network design. *International Conference on Learning Representations (ICLR)*, 2018b.
- Levine, Y., Sharir, O., Cohen, N., and Shashua, A. Quantum entanglement in deep learning architectures. *To appear in Physical Review Letters*, 2019.
- Levine, Y., Wies, N., Sharir, O., Bata, H., and Shashua, A. Limits to depth efficiencies of self-attention. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- Li, J., Sun, Y., Su, J., Suzuki, T., and Huang, F. Understanding generalization in deep learning via tensor methods. In *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, pp. 504–515. PMLR, 2020.
- Li, Y., Ma, T., and Zhang, H. Algorithmic regularization in over-parameterized matrix sensing and neural networks with quadratic activations. In *Proceedings of the 31st Conference On Learning Theory (COLT)*, pp. 2–47, 2018.
- Li, Z., Luo, Y., and Lyu, K. Towards resolving the implicit bias of gradient descent for matrix factorization: Greedy low-rank learning. *International Conference on Learning Representations (ICLR)*, 2021.
- Ma, C., Wang, K., Chi, Y., and Chen, Y. Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval and matrix completion. In *International Conference on Machine Learning (ICML)*, pp. 3351–3360, 2018.
- Mei, S., Misiakiewicz, T., and Montanari, A. Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit. In *Conference on Learning Theory (COLT)*, pp. 2388–2464, 2019.
- Milanesi, P., Kadri, H., Ayache, S., and Artières, T. Implicit regularization in deep tensor factorization. In *International Joint Conference on Neural Networks (IJCNN)*, 2021.
- Mulayoff, R. and Michaeli, T. Unique properties of wide minima in deep networks. In *International Conference on Machine Learning (ICML)*, 2020.
- Nacson, M. S., Lee, J., Gunasekar, S., Savarese, P. H. P., Srebro, N., and Soudry, D. Convergence of gradient descent on separable data. In *Proceedings of Machine Learning Research*, volume 89, pp. 3420–3428, 2019.
- Narita, A., Hayashi, K., Tomioka, R., and Kashima, H. Tensor factorization using auxiliary information. *Data Mining and Knowledge Discovery*, 25(2):298–324, 2012.
- Oymak, S. and Soltanolkotabi, M. Overparameterized nonlinear learning: Gradient descent takes the shortest path? In *International Conference on Machine Learning (ICML)*, pp. 4951–4960, 2019.
- Paszke, A., Gross, S., Chintala, S., Chanan, G., Yang, E., DeVito, Z., Lin, Z., Desmaison, A., Antiga, L., and Lerer, A. Automatic differentiation in pytorch. In *NIPS-W*, 2017.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., and Duchesnay, E. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12:2825–2830, 2011.
- Rahaman, N., Arpit, D., Baratin, A., Draxler, F., Lin, M., Hamprecht, F. A., Bengio, Y., and Courville, A. On the spectral bias of deep neural networks. In *International Conference on Machine Learning (ICML)*, pp. 5301–5310, 2019.
- Rauhut, H., Schneider, R., and Stojanac, Ž. Low rank tensor recovery via iterative hard thresholding. *Linear Algebra and its Applications*, 523:220–262, 2017.
- Razin, N. and Cohen, N. Implicit regularization in deep learning may not be explainable by norms. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.

- Recht, B., Fazel, M., and Parrilo, P. A. Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization. *SIAM review*, 52(3):471–501, 2010.
- Sharir, O. and Shashua, A. On the expressive power of overlapping architectures of deep learning. *International Conference on Learning Representations (ICLR)*, 2018.
- Sharir, O., Tamari, R., Cohen, N., and Shashua, A. Tensorial mixture models. *arXiv preprint*, 2016.
- Song, Q., Ge, H., Caverlee, J., and Hu, X. Tensor completion algorithms in big data analytics. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 13(1):1–48, 2019.
- Soudry, D., Hoffer, E., Nacson, M. S., Gunasekar, S., and Srebro, N. The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research*, 19(1):2822–2878, 2018.
- Teschl, G. *Ordinary differential equations and dynamical systems*, volume 140. American Mathematical Soc., 2012.
- Tu, S., Boczar, R., Simchowitz, M., Soltanolkotabi, M., and Recht, B. Low-rank solutions of linear matrix equations via procrustes flow. In *International Conference on Machine Learning (ICML)*, pp. 964–973, 2016.
- Vardi, G. and Shamir, O. Implicit regularization in relu networks with the square loss. In *Conference on Learning Theory (COLT)*, 2021.
- Wang, X., Wu, C., Lee, J. D., Ma, T., and Ge, R. Beyond lazy training for over-parameterized tensor decomposition. 2020.
- Woodworth, B., Gunasekar, S., Lee, J. D., Moroshko, E., Savarese, P., Golan, I., Soudry, D., and Srebro, N. Kernel and rich regimes in overparametrized models. In *Conference on Learning Theory (COLT)*, pp. 3635–3673, 2020.
- Wu, X., Dobriban, E., Ren, T., Wu, S., Li, Z., Gunasekar, S., Ward, R., and Liu, Q. Implicit regularization and convergence for weight normalization. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- Xia, D. and Yuan, M. On polynomial time methods for exact low rank tensor completion. *arXiv preprint arXiv:1702.06980*, 2017.
- Xiao, H., Rasul, K., and Vollgraf, R. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- Xu, Z. J. Understanding training and generalization in deep learning by fourier analysis. *arXiv preprint arXiv:1808.04295*, 2018.
- Yokota, T., Zhao, Q., and Cichocki, A. Smooth parafac decomposition for tensor completion. *IEEE Transactions on Signal Processing*, 64(20):5423–5436, 2016.
- Yun, C., Krishnan, S., and Mobahi, H. A unifying view on implicit bias in training linear neural networks. *International Conference on Learning Representations (ICLR)*, 2021.
- Zhou, P., Lu, C., Lin, Z., and Zhang, C. Tensor factorization for low-rank tensor completion. *IEEE Transactions on Image Processing*, 27(3):1152–1163, 2017.