

# Deep Networks Always Grok and Here is Why

Ahmed Imtiaz Humayun<sup>1</sup> Randall Balestriero<sup>2</sup> Richard Baraniuk<sup>1</sup>

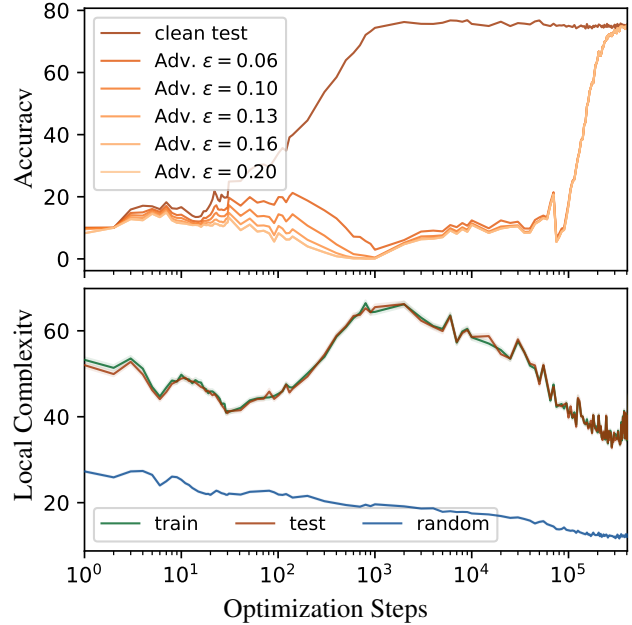
## Abstract

Grokking, or *delayed generalization*, is a phenomenon where generalization in a deep neural network (DNN) occurs long after achieving near zero training error. Previous studies have reported the occurrence of grokking in specific controlled settings, such as DNNs initialized with large-norm parameters or transformers trained on algorithmic datasets. We demonstrate that grokking is actually much more widespread and materializes in a wide range of practical settings, such as training of a convolutional neural network (CNN) on CIFAR10 or a Resnet on Imagenette. We introduce the new concept of *delayed robustness*, whereby a DNN groks adversarial examples and becomes robust, long after interpolation and/or generalization. We develop an analytical explanation for the emergence of both delayed generalization and delayed robustness based on the *local complexity* of a DNN’s input-output mapping. Our *local complexity* measures the density of so-called “linear regions” (aka, spline partition regions) that tile the DNN input space and serves as a utile progress measure for training. We provide the first evidence that, for classification problems, the linear regions undergo a phase transition during training whereafter they migrate away from the training samples (making the DNN mapping smoother there) and towards the decision boundary (making the DNN mapping less smooth there). Grokking occurs post phase transition as a robust partition of the input space thanks to the linearization of the DNN mapping around the training points. Web: [bit.ly/grok-adversarial](http://bit.ly/grok-adversarial).

## 1. Introduction

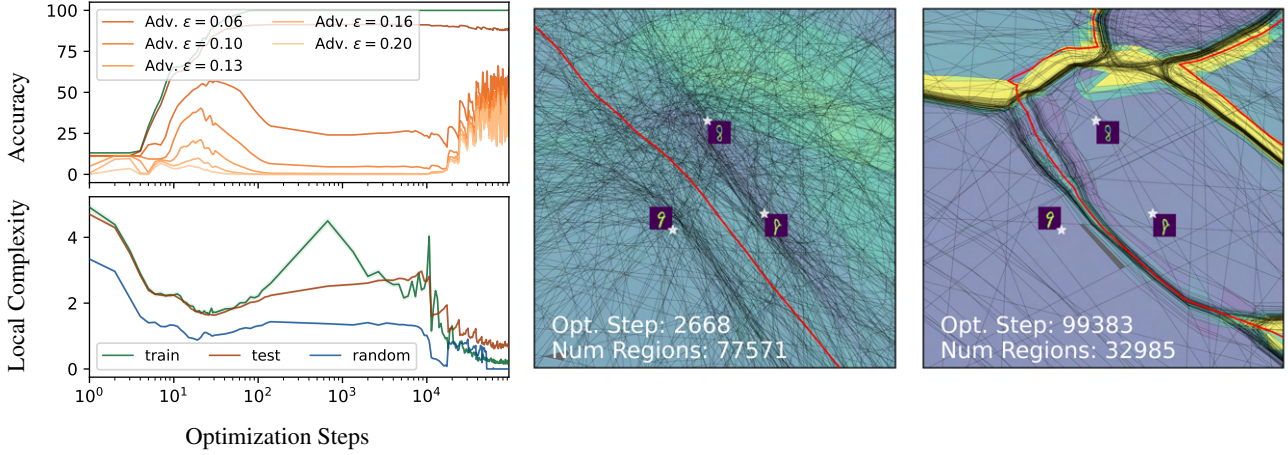
Grokking is a surprising phenomenon related to representation learning in Deep Neural Networks (DNNs) whereby

<sup>1</sup>Rice University <sup>2</sup>Brown University. Correspondence to: Ahmed Imtiaz Humayun <imtiaz@rice.edu>.



**Figure 1. Deep Neural Networks grok robustness.** When training a ResNet18 on CIFAR10, without any controlled initialization as in Liu et al. (2022), the network starts grokking adversarial examples generated using Projected Gradient Descent (Madry et al., 2017) after 10<sup>4</sup> optimization steps (top) and attains almost equal robustness and generalization performance after 2 × 10<sup>5</sup> steps. We see that, prior to grokking, the network undergoes a phase change during training in the *local complexity*, i.e., the local density of spline partition regions in the input space (bottom). After test accuracy converges, the network starts *migrating* its non-linearities away from the data points and closer to the decision boundary (see Figure 2), eventually reducing the complexity of the learned function around the data points. This increase and subsequent decrease in local non-linearity is a phenomenon visible for a wide variety of networks and training settings (see Figure 6). In this paper, we show that this particular training dynamic always results in delayed generalization or robustness.

DNNs may learn generalizing solutions to a task long after interpolating the training dataset, i.e., reaching near zero training error. It was first demonstrated by (Power et al., 2022) on simple Transformer architectures performing modular addition or division. Subsequently, multiple studies have reported instances of grokking for settings outside of modular addition, e.g., DNNs initialized with large weight



**Figure 2. Emergence of Robust Partition.** We train a 4-layer ReLU Multi Layer Perceptron (MLP) of 200 width, on  $1K$  samples from MNIST for  $10^5$  optimization steps, with batch size 200. We see that the network starts grokking adversarial examples after approximately  $10^4$  optimization steps (top-left). The local complexity around data points (bottom-left) follows a double descent curve with the final descent starting approximately after  $10^4$  optimization steps as well. *Where do the non-linearities migrate to?* In the **middle** and **right** images we present analytically computed visualizations of the DNN input space partition (Humayun et al., 2023a). The partition or *linear regions* are visualized across a 2D domain in the input space, that intersects three training samples. We see that during the final descent in local complexity, a unique structure emerges in the DNN partition geometry, where a large number of non-linearities (black lines) therefore linear regions, have concentrated around the decision boundary (red line). We dub this phenomenon *Region Migration*. Animation for an entire training run in [bit.ly/grok-splinecam](https://bit.ly/grok-splinecam).

norms for MNIST, IMDB (Liu et al., 2022), or XOR cluster data (Xu et al., 2023). For all the reported instances, DNNs that grok show a standard behavior in the training loss/accuracy curves approaching zero error as training progresses. The test error however, remains high even long after training error reaches zero. After a large number of training iterations, the DNN starts grokking—or generalizing—to the test data. This paper concerns the following question:

**Question.** *How subjective is the onset of grokking on the test data? When grokking does not manifest as a measurable change in the test set performance, could there exist an alternate set of samples for which grokking would occur?*

To find an answer to the question, we look past the test dataset towards progressively generated adversarial samples, i.e., we generate adversarial samples after each training update by using PGD (Madry et al., 2017) attacks on the test data and monitor accuracy on adversarial samples. Note that it is not guaranteed that robustness towards adversarial samples would emerge with generalization, quite the contrary has been demonstrated in previous papers. For example, Tsipras et al. (2018) introduced the generalization-robustness trade-off, Ilyas et al. (2019) demonstrated that robust networks learn fundamentally different representations. On the other hand, Li et al. (2022) introduced the notion of ‘robust generalization’ and provided theoretical proof of its existence under linear separability conditions, indicating that robustness may be achieved alongside generalization. We report the following observation:

**Observation.** *For a number of training settings, with standard initialization with or without weight decay, DNNs grok adversarial samples long after generalizing on the test dataset. We dub this novel, previously unreported form of grokking **delayed robustness**.*

We make this observation for a number of training settings including for fully connected networks trained on MNIST (Figure 2), Convolutional Neural Networks (CNNs) trained on CIFAR10 and CIFAR100 (Figure 6), ResNet18 without batch-normalization, trained on CIFAR10 (Figure 1) and Imagenette (Figure 6), and a GPT-based Architecture trained on Shakespeare Text (Figure 9). We generate adversarial examples after each training step using  $\ell_\infty$ -PGD with varying  $\epsilon \in \{0.03, 0.06, 0.10, 0.13, 0.16, 0.20\}$ ,  $\alpha = 0.0156$  and 10 (100 for MNIST) PGD steps. This observation answers our initial question: indeed there can exist a dataset other than the test dataset for which grokking manifests in classification accuracy. Moreover, we observe that the same phenomenon occurs when test set grokking is induced via initialization scaling (Figure 7) or when training transformers on Modular Addition (Figure 8).

**Question.** *How can we explain both delayed generalization and delayed robustness?*

It has previously been established that both robustness and generalization are a function of the expressivity (Xu & Manor, 2012; Li et al., 2022) as well as the local linearity (Qin et al., 2019; Balestrieri & LeCun, 2023; Humayun et al.,

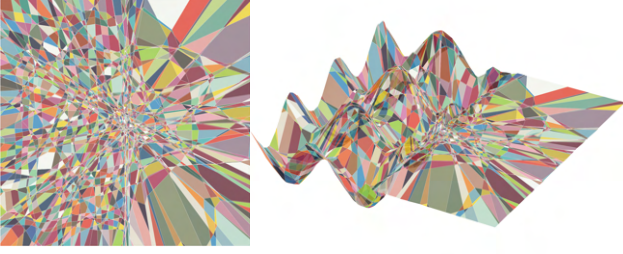


Figure 3. **Curvature and complexity.** Visual depiction of Equation (2) with a toy affine spline  $S : \mathbb{R}^2 \rightarrow \mathbb{R}$ , obtained by training an MLP to regress the piecewise function  $f(x_1, x_2) = \{\sin(x_1) + \cos(x_2)\} \mathbb{1}_{x_1 < 0}$ . Regions in the input space partition  $\Omega$  (left) and the graph of the affine spline function (right) are randomly colored. The spline partition has significantly higher density of non-linearities for  $x_1 < 0$ , i.e., the local complexity is higher where the learned function has more curvature.

2023c) of a DNN. To explain grokking, we propose a novel complexity measure based on the local non-linearity of the DNN. Our novel measure does not rely on the dataset, labels, or loss function that is used during training. It behaves as a progress measure (Barak et al., 2022; Nanda et al., 2023) that exhibits dynamics correlating with the onset of *both delayed generalization and robustness*, opening new avenues to study grokking and DNN training dynamics. We show that DNNs undergo a phase change in the local complexity (LC) averaged over data points. Based on these dynamics, we come to the following conclusion:

**Claim.** *Grokking occurs due to the emergence of a robust input space partition by a DNN, through a linearization of the DNN function around training points as a consequence of the training dynamics. This leads to larger linear regions around training points, and accumulation of non-linearities/linear regions around the decision boundary.*

We summarize our contributions as follows:

- We observe for the first time **delayed robustness**, a novel form of grokking for DNNs that occurs for a wide range of training settings and co-occurs with delayed generalization.
- We develop a novel *progress measure* (Barak et al., 2022) for DNNs based on the local complexity of a DNN’s input space partition. Our proposed measure is a proxy for the DNN’s expressivity, it is task agnostic yet informative of training dynamics. Using our measure, we detect three phases in training: two descent phases and an ascent phase. This is the first time that such dynamics in a DNN’s partition are reported. We crucially observe that a DNN’s partition regions

concentrate around the decision boundary long after interpolation, a phenomenon we term *region migration*.

- We pinpoint the origin of grokking via the spline viewpoint of DNNs (Balestrierio & Baraniuk, 2018), connect it with the circuits viewpoint (Olah et al., 2020), and show that grokking always occurs during region migration.
- Through a number of ablation studies we connect the training phases with DNN design parameters and study their changes during memorization/generalization.

We organize the rest of the paper as follows. In Section. 2 we overview the spline interpretation of deep networks and introduce our proposed local complexity measure. We also draw contrasts with common interpretability frameworks, e.g., the commonly used notion of circuits (Olah et al., 2020) in mechanistic interpretability. In Section. 3 we introduce the double descent characteristics of local complexity and connect region migration, i.e., the final phase of the double descent LC dynamics with grokking. We also present results showing that grokking does not happen when using batch normalization and provide theoretical justification. We present results connecting grokking with parameterization and memorization. Finally we draw conclusions from our results and discuss the limitations of our analysis.

## 2. Local Complexity: A New Progress Measure

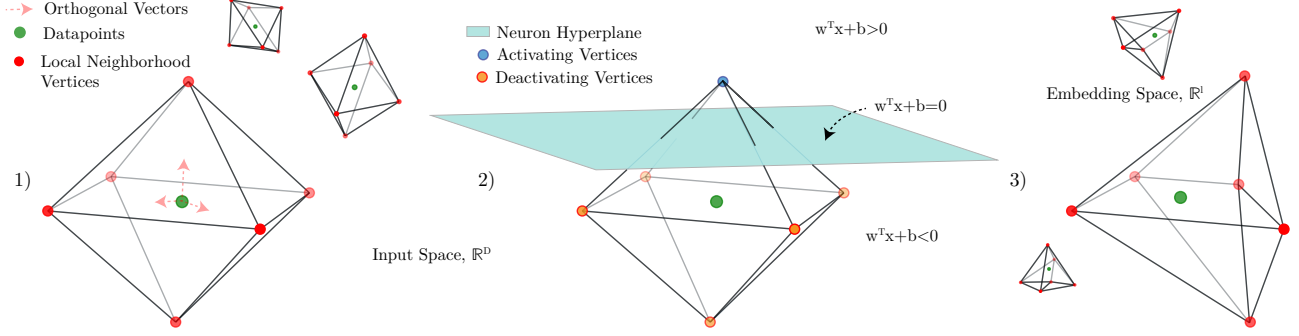
Barak et al. (2022) introduced the notion of *progress measures* for DNN training, as scalar quantities that are causally linked with the training state of a network. The spline framework enables us to introduce our proposed progress measure, the local complexity of a DNN’s partition. In later sections we show that local complexity dynamics are directly linked to grokking and present results showing its dependence on training and architectural parameters.

### 2.1. Deep Networks are Affine Spline Operators

DNNs primarily perform a sequential mapping of an input vector  $\mathbf{x}$  through  $L$  nonlinear transformations, i.e., layers, as in

$$f_{\theta}(\mathbf{x}) \triangleq \mathbf{W}^{(L)} \dots \mathbf{a} \left( \mathbf{W}^{(2)} \mathbf{a} \left( \mathbf{W}^{(1)} \mathbf{x} + \mathbf{b}^{(1)} \right) + \mathbf{b}^{(2)} \right) \dots + \mathbf{b}^{(L)}, \quad (1)$$

starting with some input  $\mathbf{x}$ . For any layer  $\ell \in \{1, \dots, L\}$ , the  $\mathbf{W}^{(\ell)}$  weight matrix, and the  $\mathbf{b}^{(\ell)}$  bias vector can be parameterized to control the type of operation for that layer, e.g., a circulant matrix as  $\mathbf{W}^{(\ell)}$  results in a convolutional layer. The operator  $\mathbf{a}$  is an element-wise nonlinearity, e.g., ReLU, and  $\theta$  is the set of all parameters of the network. According to Balestrierio & Baraniuk (2018), for any  $\mathbf{a}$  that is a continuous piecewise linear function,  $f_{\theta}$  is a continuous piecewise affine spline operator. That is, there exists a parti-



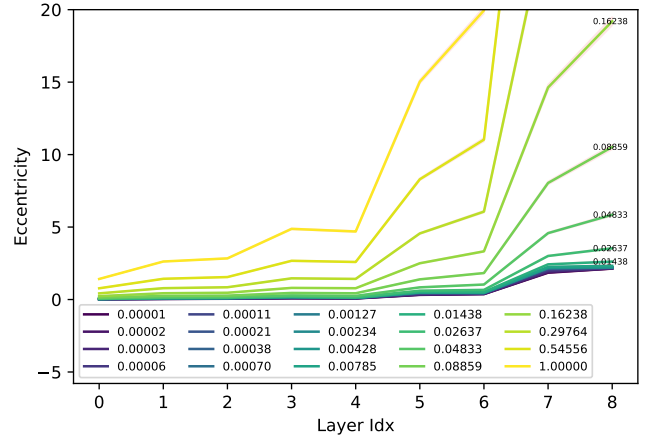
**Figure 4. Local Complexity Approximation.** 1) Given a point in the input space  $x \in \mathbb{R}^D$ , we start by sampling  $P$  orthonormal vectors  $\{v_1, v_2, \dots, v_P\}$  to obtain cross-polytopal frame  $V_x = \{x \pm r * v_p \forall p\}$  centered on  $x$ , where  $r$  is a radius parameter. We consider the convex hull  $\text{conv}(V_x)$  as the local neighborhood of  $x$ . 2) If any neuron hyperplane intersects the neighborhood  $\text{conv}(V_x)$  then the pre-activation sign will be different for the different vertices. We can therefore count the number neurons for a given layer, which results in sign changes in the pre-activation of  $V_x$  to quantify local complexity  $x$  for that layer. 3) By embedding  $V_x$  to the input of the next layer, we can obtain a coarse approximation of the local neighborhood of  $x$  and continue computing local complexity in a layerwise fashion.

tion  $\Omega$  of the DNN’s input space  $\mathbb{R}^D$  (for example, Figure 3 left) comprised of non-overlapping regions that span the entire input space. On any region of the partition  $\omega \in \Omega$ , the DNN’s input-output mapping is a simple affine mapping with parameters  $(A_\omega, b_\omega)$ . In short, we can express  $f_\theta$  as

$$f_\theta(x) = \sum_{\omega \in \Omega} (A_\omega x + b_\omega) \mathbb{1}_{\{x \in \omega\}}, \quad (2)$$

where,  $\mathbb{1}_{\{x \in \omega\}}$  is an indicator function that is non-zero for  $x \in \omega$ .

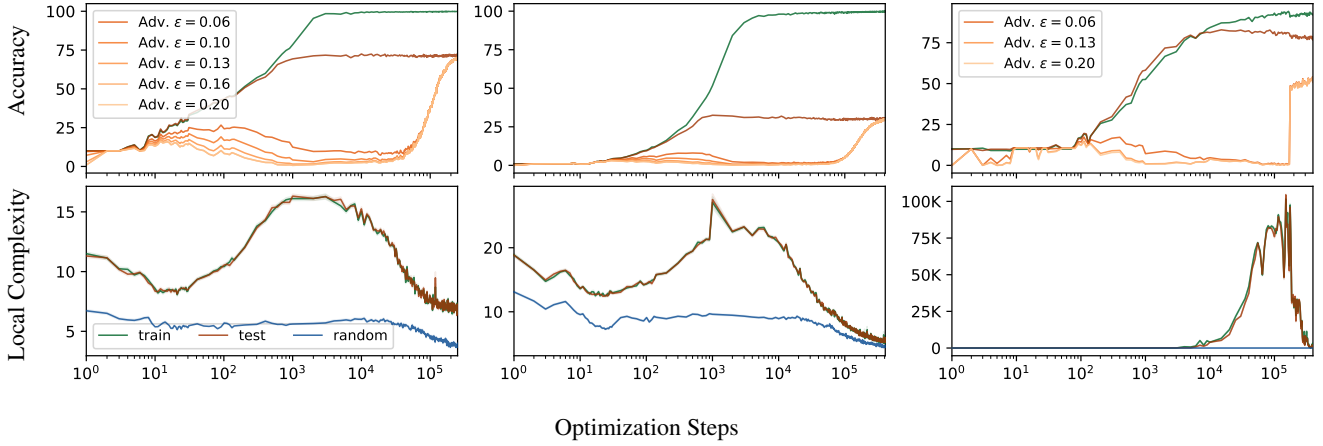
**Curvature and Linear Regions.** Formulations like that in Equation (2) that represent DNNs as continuous piecewise affine splines, have previously been employed to make theoretical studies amenable to actual DNNs, e.g. in generative modeling (Humayun et al., 2022), network pruning (You et al., 2021), and OOD detection (Ji et al., 2022). Empirical estimates of the density of linear regions in the spline partition have also been employed in sensitivity analysis (Novak et al., 2018), quantifying non-linearity (Gamba et al., 2022), quantifying expressivity (Raghu et al., 2017) or to estimate the complexity of spline functions (Hanin & Rolnick, 2019). We demonstrate the relationship between function curvature and linear region density through a toy example in Figure 3. In Figure 3-left and Figure 2-(middle,right), any contiguous line is a non-linearity in the input space, corresponding to a single neuron of the network. All the non-linearities re-orient themselves during training to be able to obtain the target function (Figure 3-right). Therefore, in Figure 3, we see that DNN partitions have higher density of linear regions/non-linearities/knots in the spline partition, where the target function curvature is non-zero.



**Figure 5. Deformation with depth.** Change of average eccentricity (Xu et al., 2021) of the input space neighborhoods  $V_x$  by different layers of a CNN trained on the CIFAR10 dataset, for different radius  $r$ . We see that, for larger radius, the deformation increases with depth almost exponentially. For  $r \leq 0.014$  deformation is low, indicating that smaller radius neighborhoods are reliable for LC computation on deeper networks. Values are averaged over neighborhoods sampled for 1000 training points from CIFAR10. For ResNet18, see Figure 23.

## 2.2. Measuring Local Complexity using the Deep Network Spline Partition

Suppose a domain is specified as the convex hull of a set of vertices  $V = [v_1, \dots, v_p]^T$  in the DNN’s input space. We wish to compute the local complexity or smoothness (Hanin & Rolnick, 2019) for neighborhood  $\mathcal{V} = \text{conv}(V)$ . Consider a single hidden layer of a network. Let’s denote the DNN layer weight as  $W^{(\ell)} \triangleq [w_1^{(\ell)}, \dots, w_{D^{(\ell)}}^{(\ell)}]$ ,  $b^{(\ell)}$  where  $\ell$  is the layer index,  $w_i^{(\ell)}$  is the  $i$ -th row of  $W^{(\ell)}$  or weight of the  $i$ -th neuron, and  $D^{(\ell)}$  is the output space



**Figure 6. Grokking across datasets and architectures.** From left to right, examples of delayed robustness emerging late in training for a CNN trained on CIFAR10, CNN trained on CIFAR100, and ResNet18 trained on the Imagenette<sup>2</sup> dataset. Clear double descent behavior visible in the local complexity of CNN with CIFAR10 and CIFAR100. The ResNet18 trained with Imagenette obtains a very high local complexity during the ascent phase of training. To compute local complexity we consider 25 dimensional neighborhoods centered on 1024 train, test or random samples. We use  $r = 0.005$  for CNN and  $r = 10^{-4}$  for ResNet18.

dimension of layer  $\ell$ . The forward pass through this layer for  $\mathbf{V}$  can be considered an inner product with each row of the weight matrix  $W^{(\ell)}$  followed by a continuous piecewise linear activation function. Without loss of generality, let's consider ReLU as the activation function in our network. The partition at the input space of layer  $\ell$  can therefore be expressed as the set of all hyperplane equations formed via the neuron weights such as:

$$\partial\Omega = \bigcup_{i=1}^{D^{(\ell)}} \mathcal{H}_i^{(\ell)} \quad (3)$$

$$\mathcal{H}_i^{(\ell)} = \left\{ \mathbf{x} \in \mathbb{R}^{D^{(\ell-1)}} : \langle \mathbf{w}_i^{(\ell)}, \mathbf{x} \rangle + b_i^{(\ell)} = 0 \right\}, \quad (4)$$

which is also the set of layer  $\ell$  non-linearities. Let,  $\Phi = f_{1:\ell-1}(\mathcal{V})$  be the embedded representation of the neighborhood  $\mathcal{V}$  by layer  $\ell - 1$  of the network. Therefore, approximating the local complexity of  $\mathcal{V}$  induced by layer  $\ell$ , would be equivalent to counting the number of linear regions in,

$$\Phi \cap \partial\Omega = \bigcup_{i=1}^{D^{(\ell)}} \Phi \cap \mathcal{H}_i^{(\ell)}. \quad (5)$$

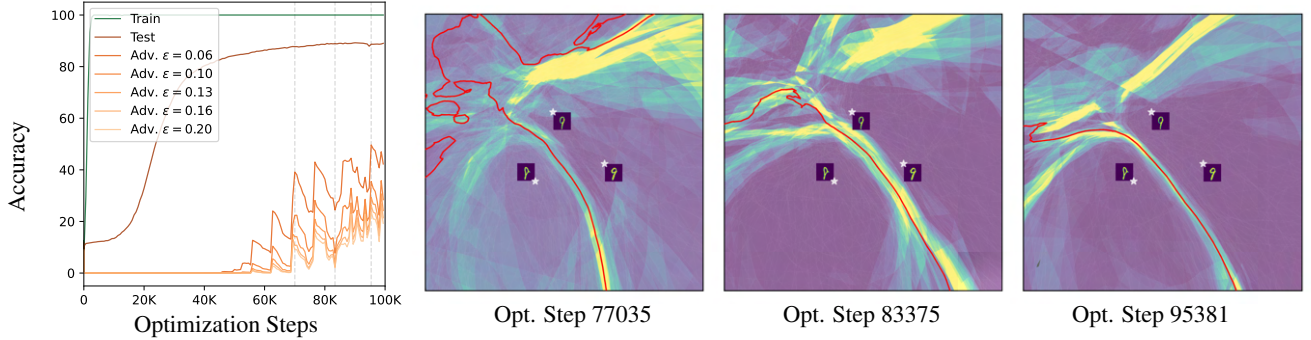
The local partition inside  $\Phi$  results from an arrangement of hyperplanes; therefore the number of regions is of the order  $\mathcal{N}^{D^{(\ell-1)}}$  (Toth et al., 2017), where

$$\mathcal{N} = |\{i : i = 1, 2, \dots, D^{(\ell)} \text{ and } \mathcal{H}_i^{(\ell)} \cap \Phi \neq \emptyset\}|, \quad (6)$$

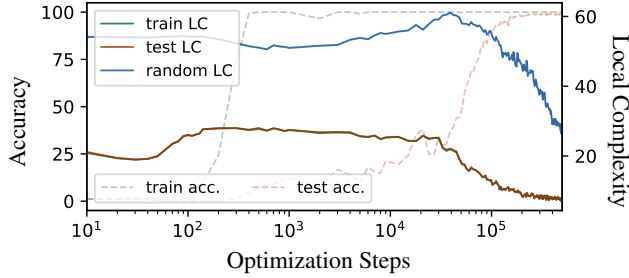
is the number of hyperplanes from layer  $\ell$  intersecting  $\Phi$ . We consider  $\mathcal{N}$  as a proxy for the local complexity of any neighborhood  $\Phi$ . To make computation tractable, let,  $\Phi \approx \hat{\Phi} = \text{conv}(f_{1:\ell-1}(\mathbf{V}))$ . Therefore, for  $\hat{\Phi}$ , any

sign changes in layer  $\ell$  pre-activations is due to the corresponding neuron hyperplanes intersecting  $\text{conv}(\mathbf{V})$ . For a single layer, the local complexity (LC) for a sample in the input space can be approximated by the number of neuron hyperplanes that intersect  $\mathbf{V}$  embedded to that layers input space. If we consider input space neighborhoods with the same volume, then our approximation method measures the un-normalized density of non-linearity in an input space locality, which we consider a proxy for local complexity. We highlight that this is tied to the VC-dimension of (ReLU) DNN (Bartlett et al., 2019) where the more regions are present the more expressive the decision boundary can be (Montufar et al., 2014). In Figure 4, we provide a visual explanation of our method for local complexity approximation through a cartoon schematic diagram. To summarize, we consider randomly oriented  $P$  dimensional  $\ell_1$  norm balls with radius  $r$ , i.e., cross-polytopes centered on any given data point  $x$  as a frame defining the neighborhood. We therefore follow the steps entailed in Figure 4 in a layer-wise fashion, to approximate the local complexity in the prescribed neighborhood for a given layer.

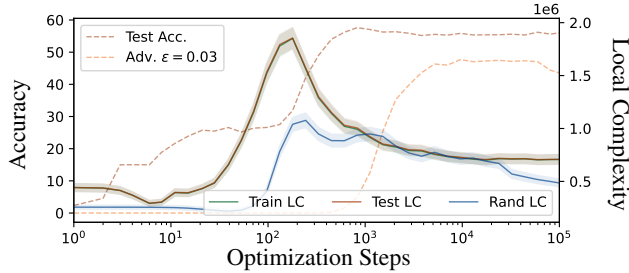
**Sensitivity of approximation to  $P$  and  $r$ .** One of the possible limitations of local complexity measure is the deformation of the local neighborhood when it is passed through a network from layer to layer, as shown in Figure 4. For different radius  $r$  of the input space neighborhood  $\mathbf{V}_x$  centered on any arbitrary data point  $x$ , we compute the graph eccentricity (Xu et al., 2021) of  $\mathbf{V}_x$  after being embedded by different layers of a CNN. We present the results in Figure 5 for 1000 different training data points for a CNN trained on CIFAR10. The higher the change of eccentricity compared to the input space (index 0), the more likely the neighbor-



**Figure 7. Grokking visualized.** We induce grokking by randomly initializing a 4 depth 200 width ReLU MLP and scaling the initialized parameters by eight following (Liu et al., 2022). In the leftmost figure, we can see that the grokking is visible for both the test samples as well as adversarial examples generated using the test set. We see that the network robustness, periodically increases. By visualization the partition and curvature of the function across a 2D slice of the input space (Humayun et al., 2023a), we see that the network periodically increases the concentration of non-linearity around its decision boundary, making the boundary sharper at each robustness peak. This occurs even when the network doesn’t undergo delayed generalization (Figure 2). As the local complexity around the decision boundary increases, the local complexity around data points farther from the decision boundary decreases (Figure 26).



**Figure 8. Region migration in modular addition.** By measuring the local complexity for the GeLU activated fully connected layers of a Transformer architecture, we see that here as well, region migration occurs during grokking.



**Figure 9. Delayed robustness in LLMs.** Grokking observed in a GPT architecture with 12 heads and 12 layers trained on next character prediction using the Shakespeare Text Dataset. We see that the second local complexity descent starts prior to the test acc. peak, and descent continues while the network groks  $\epsilon = 0.03$   $\ell_\infty$ -PGD adversarial examples, generated in the token embedding space. Approximate input space partition visualized in Figure 19 and Figure 20.

hood gets deformed, leading to less reliable approximation. We see that below a certain radius value, deformation by

the CNN is limited and does not exponentially increase. In subsequent experiments however, e.g., Figure 27, we have observed that the dynamics of local complexity is similar between large and small  $r$  neighborhoods. We present more validation experiments in Appendix A. Our proposed method can also be used to approximate the input space partition formed by a neural network. In Figure 25 we compare the partition approximated via LC computations on a grid, with analytically computed partition via (Humayun et al., 2023b). In Figure 19 and Figure 20 we present the input space partition approximated for a GPT architecture before and after delayed robustness occurs.

**Experimental Setup.** For all experiments we sample 1024 train test and random points for local complexity (LC) computation, except for the MNIST experiments, where we use 1000 training points (all of the training set where applicable) and 10000 test and random points for LC computation. We use  $r = 0.005$  and  $P = 25$  unless specified otherwise and except for the ResNet18 experiments with Imagenette where we use  $r = 10^{-4}$ . For training, we use the Adam optimizer and a weight decay of 0 for all the experiments except for the MNIST-MLP experiments where we use a weight decay of 0.01. Unless specified, we use CNNs with 5 convolutional layers and two linear layers. For the ResNet18 experiments with CIFAR10, we use a pre-activation architecture with width 16. For the Imagenette experiments, we use the standard torchvision Resnet architecture. For all settings we do not use Batch Normalization, as reasoned in Appendix B. In all our plots, we denote training accuracy/LC using green, test accuracy/LC using orange and random LC using blue colors. We also color curves for adversarial examples using different shades of orange. All local complexity plots show the 99% confidence interval.

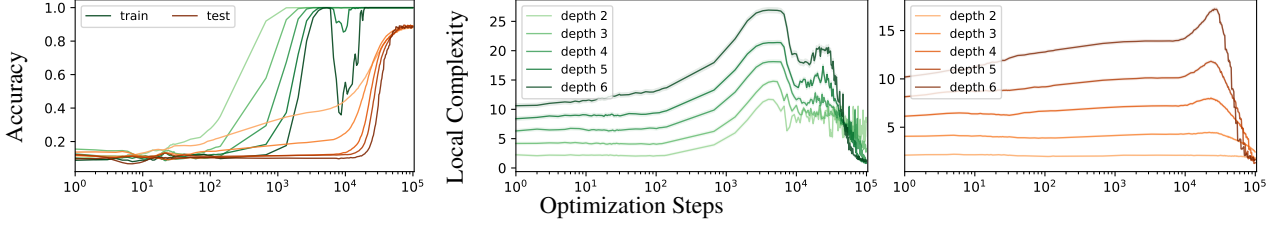


Figure 10. **Local complexity across depths.** From left to right, accuracy, local complexity around training and local complexity around test data points, for an MLP trained on MNIST with width 200 and varying depth. As depth is increased the max LC during ascent phase becomes larger. We can also see a distinct second peak right before the descent phase.

### 3. Local Complexity Training Dynamics and Grokking

#### 3.1. Emergence of a Robust Partition

We start our exploration of the training dynamics of deep neural networks by formalizing the phases of local complexity observed during training. In all our experiments either involving delayed generalization or robustness, we see three distinct phases in the dynamics of local complexity:

- *The first descent*, when the local complexity start by descending after initialization. This phase is subject to the network parameterization as well as initialization, e.g., when grokking is induced in the MLP-MNIST case with scaled initialization, we do not see the first descent (Figure 29, Figure 10).
- *The ascent phase*, when the local complexity accumulates around both training and test data points. The ascent phase happens ubiquitously, and the local complexity generally keeps ascending until training interpolation is reached (e.g., Figure 6, Figure 1). During the ascent phase, the training local complexity may be higher for training data points than for test data points, indicating an accumulation of non-linearities around training data compared to test data Figure 2.
- *The second descent phase* or region migration phase, during which the network moves the linear regions or non-linearities away from the training and test data points. Focusing on Figure 2-bottom-left and Figure 29 for the MLP-MNIST setting, one perplexing observation that we make is that the local complexity around random points – uniformly sampled from the domain of the data – also decreases during the final descent phase. This would mean that the non-linearities are not randomly moving away from the training data, but systematically reorganizing where we do not have our LC approximation probes. To better understand the phenomenon, we consider a square domain  $\mathbb{D}$  that passes through three MNIST training points, and use Splinecam (Humayun et al., 2023a) to analytically compute the input space partition on  $\mathbb{D}$ . In short, Splinecam uses the weights of the network to exactly compute the input space representation of each neuron’s zero-level set on  $\mathbb{D}$  (black lines in

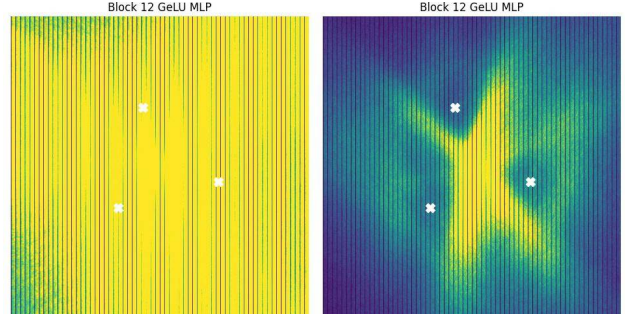
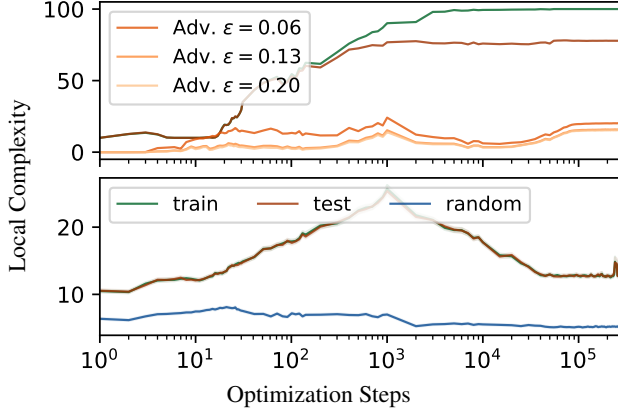


Figure 11. Token embedding space partition formed by the GeLU activated MLP of Block 12 of the GPT model mentioned in Figure 9. The partition is approximated by computing LC on a  $512 \times 512$  grid on a 2D subspace in the token embedding space. Note that after grokking (right), the three random samples used to determine the 2D subspace, has visibly lower local complexity in its immediate neighborhood.

Figure 2). We present Splinecam visualizations for different optimization steps in Figure 2, Figure 7, and Figure 26. Through these visualizations, we see clear evidence that *during the second descent phases of training, linear regions or the non-linearities of the network, migrate close to the decision boundary creating a robust partition in the input space*. The robust partition contains large linear regions around the training data, as suggested by papers in literature as a precursor for robustness (Qin et al., 2019). Moreover, during region migration, *the network intends to lower the local complexity around training points*, resulting in a decrease in local complexity around training even compared to test data points.

**Local complexity as a progress measure.** While we don’t quite understand why the network goes from accumulation to repelling of non-linearities around the training data between the ascent and second descent phases, we see that the second descent always precedes the onset of delayed generalization or delayed robustness. In Figure 7-middle and right, we present splinecam visualizations for a network during grokking. The colors denote the norm of the slope parameter  $A_\omega$  for each region  $\omega$  computed obtained via

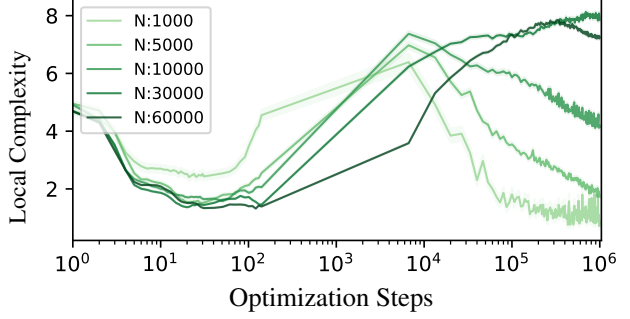


**Figure 12. Batch-norm removes grokking.** Training a CNN with an identical setting as in Figure 6-left, except the CNN now has Batch Normalization layers after every convolution. With the presence of batchnorm, the LC values increase, the initial descent gets removed and most importantly, grokking does not occur for adversarial samples.

SplineCam. We see that while a network groks, the regions start concentrating around the decision boundary where the network has the highest norm. This is intuitive because in such classification settings, an increase of local complexity around the decision boundary allows the function to sharply transition from one class to another. Therefore, the more the non-linearities converge towards the decision boundary, the higher the function norm can be while smoothly transitioning as well. We have provided an animation showing the evolution of partition geometry and emergence of the robust partition during training here<sup>3</sup>. In the animation, we can see that the partition periodically switches between robust configurations during region migration. As time progresses we see increasing accumulation of the non-linearities around the decision boundary. These results undoubtedly show that the local non-linearity or local complexity dynamics is directly tied to the partition geometry and emergence of delayed generalization/robustness.

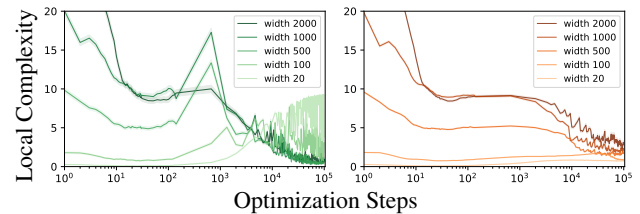
**Relationship with Circuits.** A common theme in mechanistic interpretability, especially when it comes to explaining the grokking phenomenon, is the idea of ‘circuit’ formation during training (Nanda et al., 2023; Varma et al., 2023; Olah et al., 2020). A circuit is loosely defined as a subgraph of a deep neural network containing neurons (or linear combination of neurons) as nodes, and weights of the network as edges. Recall that Equation (2) expresses the operation of the network in a region-wise fashion, i.e., for all input vectors  $\{x : x \in \omega\}$ , the network performs the same affine operation using parameters  $(A_\omega, b_\omega)$  while mapping  $x$  to the output. The affine parameters for any given region,

<sup>3</sup>[bit.ly/grok-splinecam](https://bit.ly/grok-splinecam)



**Figure 13. Memorization requirement delays grok.** When training an MLP on varying number of randomly labeled MNIST samples, we see that with increase in the number of samples, the local complexity dynamics get delayed, especially the ascent phase gets elongated. This shows that with increased demand for memorization the network takes longer to complete ascent and later undergo region migration.

are a function of the active neurons in the network as was shown by Humayun et al. (2023a) (Lemma 1). Therefore for each region, we necessarily have a circuit or subgraph of the network performing the linear operation. Between two neighboring regions, only one node of the circuit changes. From this perspective, our local complexity measure can be interpreted as a way to measure the density of unique circuits formed in a locality of the input space as well. While in practice this would result in an exponential number of circuits, the emergence of a robust partition show that towards the end of training, the number of unique circuits get drastically reduced. This is especially true for sub-circuits corresponding to deeper layers only. In Figure 24, we show the robust partition in a layerwise fashion. We can see that for deeper layers, there exists large regions, i.e., embedding regions with only one circuit operation through the layer. This result, matches with the intuition provided by Nanda et al. (2023) on the cleanup phase of circuit formation late in training.



**Figure 14. Increasing width hastens region migration.** LC dynamics while training an MLP with varying width on MNIST. For the peak LC achieved around training points during the ascent phase, we see an initial increase and then decrease as the network gets overparameterized. For test and random samples, we see the LC during ascent phase saturating as we increase width.

Table 1. Summary of all the experiments showing the relationship between delayed generalization and training/model hyperparameters.

| Training Intervention                       | Dataset          | Model    | Description                                    | Effect on Adv. Grokking              |
|---------------------------------------------|------------------|----------|------------------------------------------------|--------------------------------------|
| Increasing Batchsize                        | CIFAR10          | Resnet18 | Batchsize Increased from 64 to 512             | Expedited (Figure 17)                |
| Increasing Parameters: Depth                | MNIST            | MLP      | Increasing Depth from 2 to 6                   | Expedited (Figure 27)                |
| Increasing Parameters: Depth                | Shakespeare Text | GPT      | Increasing number of layers from 6 to 12       | Expedited                            |
| Increasing Parameters: Width                | MNIST            | MLP      | Increasing Width from 20 to 2000               | Expedited (Figure 14)                |
| Increasing Parameters: Width                | CIFAR10          | Resnet18 | Increasing Width from 16 filters to 64 filters | Expedited                            |
| Increasing Regularization                   | MNIST            | MLP      | Weight Decay increased from 0. to 1            | Delayed (Figure 30)                  |
| Increasing Training Data                    | MNIST            | MLP      | Training data increased from 1K to 60K         | Expedited (Figure 32)                |
| Increasing Training Data (Randomly Labeled) | MNIST            | MLP      | Training data increased from 1K to 60K         | Delayed (Figure 13)                  |
| No BatchNorm -> BatchNorm                   | CIFAR10          | Resnet18 | Adding BatchNorm Layer after convolution       | Does not occur/significantly delayed |
| No BatchNorm -> BatchNorm                   | CIFAR10          | CNN      | Adding BatchNorm Layer after convolution       | Does not occur/significantly delayed |

#### 4. What Affects the Progress Measure?

**Parameterization.** In Figures 14, 27 and 29, we see that increasing the number of parameters either by increasing depth, or by increasing width of the network in our MNIST-MLP experiments, hastens region migration, therefore makes grokking happen earlier.

**Weight Decay** regularizes a neural network by reducing the norm of the network weights, therefore reducing the per region slope norm as well. We train a CNN with depth 5 and width 32 on CIFAR10 with varying weight decay. In Figure 30 we present the train, test and random LC for our experiments for neighborhoods of different radius. Weight decay does not seem to have a monotonic behavior as it both delays and hastens region migration, based on the amount of weight decay.

**Batch Normalization.** Batch normalization removes grokking. In Appendix B, we show that at each layer  $\ell$  of a DN, BN explicitly adapts the partition so that the partition boundaries are as close to the training data as possible. This is confirmed by our experiments in Figure 12 where we see that grokking adversarial examples ceases to occur compared to the non-batchnorm setting in Figure 6. BN also removes the first descent, monotonically increasing the local complexity around the data manifold and after a while undergoing a phase change and decreasing. The degree of region migration is reduced during this phase, as can be seen in the higher LC when we use batch normalization. While training a ResNet18 with Batch Norm on Imagenet Full (Figure 22), we see that the local complexity keeps increasing indefinitely, removing any signs of region migration.

**Activation function** While most of our experiments use ReLU activated networks, in Figure 34 we present results for a GeLU activated MLP, as well as in Figure 8 we present results for a GeLU activated Transformer. For both settings we see similar training dynamics as is observed for ReLU.

**Effect of Training Data.** We control the training dataset to either induce higher generalization on higher memorization. Recall that in our MNIST experiments, we use 1k training samples. We increase the number of samples in our

dataset to monitor the effect of grokking Figure 28 and LC Figure 32. We see that increasing the size of the dataset hastens grokking. On the other hand we also sweep the dataset size for a random label memorization task Figure 13, Figure 36. We see that in this case, increasing dataset size results in more memorization requirement, therefore it delays the region migration phase.

#### 5. Conclusions and Limitations

We have pursued a thorough empirical study of grokking, both on the test dataset and adversarial examples generated using the test dataset. We obtained new observations hinting that grokking is a common phenomenon in deep learning that is not restricted to particular tasks or DNN initialization. Upon this discovery, we delved into DNNs geometry to isolate the root cause of both delayed generalization and robustness which we attributed to the DNN’s linear region migration that occurs in the latest phase of training. Again, the observation of such migration of the DNN partition is a new discovery of its own right. We hope that our analysis has provided novel insights into DNNs training dynamics from which grokking naturally emerges. While we empirically study the local complexity dynamics, a theoretical justification behind the double descent behavior is lacking. At a high level, it is clear that the classification function being learned has its curvature concentrated at the decision boundary and approximation theory would normally dictate a free-form spline to therefore concentrate its partition regions around the decision boundary to minimize approximation error. However, it is not clear why that migration occurs so late in the training process, and we hope to study that in future research. We also see empirical evidence of region migration while using Adam as the optimizer. The training dynamics of stochastic gradient descent, as well as sharpness aware minimization (Andriushchenko & Flammarion, 2022) can also be studied using our framework. There can be possible connections between region migration and neural collapse (Papayan et al., 2020) which are not explored in this paper. The spline viewpoint of deep neural networks may provide strong geometric insights to assist in mechanistic understanding in future works as well.

## Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. One takeaway of this work is that training Deep Neural Networks longer may lead to increased robustness. Training networks especially foundation models for longer may have potential societal consequences in terms of carbon emissions.

## Acknowledgements

Humayun and Baraniuk were supported by NSF grants CCF1911094, IIS-1838177, and IIS-1730574; ONR grants N00014-18-12571, N00014-20-1-2534, and MURI N00014-20-1-2787; AFOSR grant FA9550-22-1-0060; and a Vannevar Bush Faculty Fellowship, ONR grant N00014-18-1-2047.

## References

- Andriushchenko, M. and Flammarion, N. Towards understanding sharpness-aware minimization. In *International Conference on Machine Learning*, pp. 639–668. PMLR, 2022.
- Balestrieri, R. and Baraniuk, R. A spline theory of deep networks. In *Proc. ICML*, pp. 374–383, 2018.
- Balestrieri, R. and Baraniuk, R. G. Batch normalization explained. *arXiv preprint arXiv:2209.14778*, 2022.
- Balestrieri, R. and LeCun, Y. Police: Provably optimal linear constraint enforcement for deep neural networks. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 1–5. IEEE, 2023.
- Barak, B., Edelman, B., Goel, S., Kakade, S., Malach, E., and Zhang, C. Hidden progress in deep learning: Sgd learns parities near the computational limit. *Advances in Neural Information Processing Systems*, 35:21750–21764, 2022.
- Bartlett, P. L., Harvey, N., Liaw, C., and Mehrabian, A. Nearly-tight vc-dimension and pseudodimension bounds for piecewise linear neural networks. *The Journal of Machine Learning Research*, 20(1):2285–2301, 2019.
- Croce, F. and Hein, M. Reliable evaluation of adversarial robustness with an ensemble of diverse parameter-free attacks. In *International conference on machine learning*, pp. 2206–2216. PMLR, 2020.
- Gamba, M., Chmielewski-Anders, A., Sullivan, J., Azizpour, H., and Bjorkman, M. Are all linear regions created equal? In *AISTATS*, pp. 6573–6590, 2022.
- Garbin, C., Zhu, X., and Marques, O. Dropout vs. batch normalization: an empirical study of their impact to deep learning. *Multimedia Tools and Applications*, 79:12777–12815, 2020.
- Hanin, B. and Rolnick, D. Complexity of linear regions in deep networks. *arXiv preprint arXiv:1901.09021*, 2019.
- Humayun, A. I., Balestrieri, R., and Baraniuk, R. Polarity sampling: Quality and diversity control of pre-trained generative networks via singular values. In *CVPR*, pp. 10641–10650, 2022.
- Humayun, A. I., Balestrieri, R., Balakrishnan, G., and Baraniuk, R. G. Splinecam: Exact visualization and characterization of deep network geometry and decision boundaries. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 3789–3798, June 2023a.
- Humayun, A. I., Balestrieri, R., Balakrishnan, G., and Baraniuk, R. G. Splinecam: Exact visualization and characterization of deep network geometry and decision boundaries. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 3789–3798, 2023b.
- Humayun, A. I., Casco-Rodriguez, J., Balestrieri, R., and Baraniuk, R. Provable instance specific robustness via linear constraints. In *2nd AdvML Frontiers Workshop at International Conference on Machine Learning 2023*, 2023c.
- Ilyas, A., Santurkar, S., Tsipras, D., Engstrom, L., Tran, B., and Madry, A. Adversarial examples are not bugs, they are features. *Advances in neural information processing systems*, 32, 2019.
- Ioffe, S. and Szegedy, C. Batch normalization: Accelerating deep network training by reducing internal covariate shift. *arXiv preprint arXiv:1502.03167*, 2015.
- Ji, X., Pascanu, R., Hjelm, R. D., Lakshminarayanan, B., and Vedaldi, A. Test sample accuracy scales with training sample density in neural networks. In *Conference on Lifelong Learning Agents*, pp. 629–646. PMLR, 2022.
- Kubo, M., Banno, R., Manabe, H., and Minoji, M. Implicit regularization in over-parameterized neural networks. *arXiv preprint arXiv:1903.01997*, 2019.
- Li, B., Jin, J., Zhong, H., Hopcroft, J., and Wang, L. Why robust generalization in deep learning is difficult: Perspective of expressive power. *Advances in Neural Information Processing Systems*, 35:4370–4384, 2022.
- Liu, Z., Michaud, E. J., and Tegmark, M. Omnigrok: Grokking beyond algorithmic data. *arXiv preprint arXiv:2210.01117*, 2022.
- Madry, A., Makelov, A., Schmidt, L., Tsipras, D., and Vladu, A. Towards deep learning models resistant to

- adversarial attacks. *arXiv preprint arXiv:1706.06083*, 2017.
- Montufar, G. F., Pascanu, R., Cho, K., and Bengio, Y. On the number of linear regions of deep neural networks. In *NeurIPS*, pp. 2924–2932, 2014.
- Nanda, N., Chan, L., Lieberum, T., Smith, J., and Steinhart, J. Progress measures for grokking via mechanistic interpretability. *arXiv preprint arXiv:2301.05217*, 2023.
- Novak, R., Bahri, Y., Abolafia, D. A., Pennington, J., and Sohl-Dickstein, J. Sensitivity and generalization in neural networks: an empirical study. *arXiv preprint arXiv:1802.08760*, 2018.
- Olah, C., Cammarata, N., Schubert, L., Goh, G., Petrov, M., and Carter, S. Zoom in: An introduction to circuits. *Distill*, 5(3):e00024–001, 2020.
- Papayan, V., Han, X., and Donoho, D. L. Prevalence of neural collapse during the terminal phase of deep learning training. *Proceedings of the National Academy of Sciences*, 117(40):24652–24663, 2020.
- Power, A., Burda, Y., Edwards, H., Babuschkin, I., and Misra, V. Grokking: Generalization beyond overfitting on small algorithmic datasets. *arXiv preprint arXiv:2201.02177*, 2022.
- Qin, C., Martens, J., Gowal, S., Krishnan, D., Dvijotham, K., Fawzi, A., De, S., Stanforth, R., and Kohli, P. Adversarial robustness through local linearization. *Advances in Neural Information Processing Systems*, 32, 2019.
- Raghu, M., Poole, B., Kleinberg, J., Ganguli, S., and Dickstein, J. S. On the expressive power of deep neural networks. In *ICML*, pp. 2847–2854, 2017.
- Tan, J., LeJeune, D., Mason, B., Javadi, H., and Baraniuk, R. G. A blessing of dimensionality in membership inference through regularization. In *International Conference on Artificial Intelligence and Statistics*, pp. 10968–10993. PMLR, 2023.
- Toth, C. D., O’Rourke, J., and Goodman, J. E. *Handbook of discrete and computational geometry*. CRC press, 2017.
- Tsipras, D., Santurkar, S., Engstrom, L., Turner, A., and Madry, A. Robustness may be at odds with accuracy. *arXiv preprint arXiv:1805.12152*, 2018.
- Varma, V., Shah, R., Kenton, Z., Kramár, J., and Kumar, R. Explaining grokking through circuit efficiency. *arXiv preprint arXiv:2309.02390*, 2023.
- Xu, H. and Mannor, S. Robustness and generalization. *Machine learning*, 86:391–423, 2012.
- Xu, K., Ilić, A., Iršič, V., Klavžar, S., and Li, H. Comparing wiener complexity with eccentric complexity. *Discrete Applied Mathematics*, 290:7–16, 2021.
- Xu, Z., Wang, Y., Frei, S., Vardi, G., and Hu, W. Benign overfitting and grokking in relu networks for xor cluster data. *arXiv preprint arXiv:2310.02541*, 2023.
- You, H., Balestrieri, R., Lu, Z., Kou, Y., Shi, H., Zhang, S., Wu, S., Lin, Y., and Baraniuk, R. Max-affine spline insights into deep network pruning. *arXiv preprint arXiv:2101.02338*, 2021.

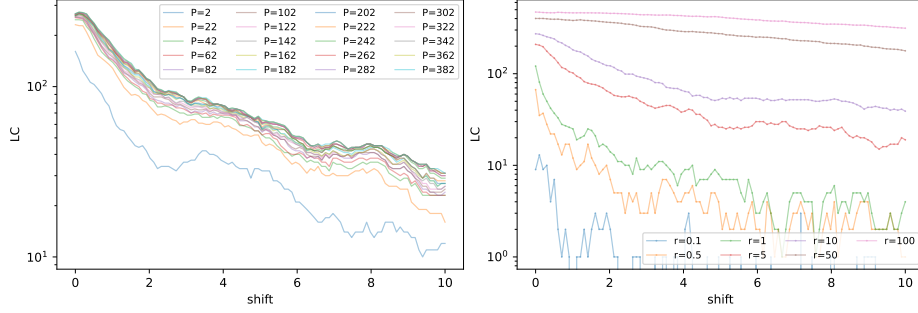


Figure 15. LC for a  $P$  dimensional neighborhood with radius  $r$  while being shifted from the origin  $[0]^d$  to vector  $[10]^d$ . In **left**, we vary  $P$  with fixed  $r = 5$  while on **right** we vary  $r$  for fixed  $P = 20$ . We see that for all the settings, shifting away from the origin reduces LC. The increase of LC with the neighborhood dimensionality  $P$  gets saturated as we increase  $P$ , showing that lower dimensional neighborhoods can be good enough for approximating LC. Increasing  $r$  on the other hand, increases LC and reduces LC variations between shifts, since the neighborhood becomes larger and LC becomes less local.

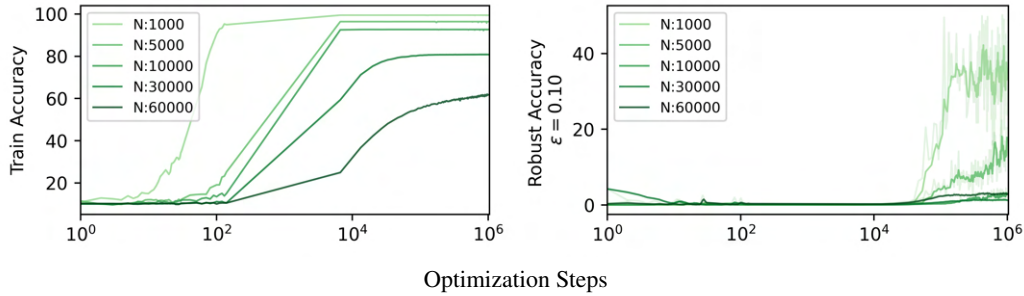


Figure 16. Training accuracy and robust accuracy for the networks trained on randomly labeled MNIST samples presented in Figure 13.

## A. Empirical analysis of our proposed method

Computing the exact number of linear regions or piecewise-linear hyperplane intersections for an deep network with  $N$ -dimensional input space neighborhood has combinatorial complexity and therefore is intractable. This is one of the key motivations behind our approximation method.

**MLP with zero bias.** To validate our method, we start with a toy experiment with a linear MLP with width 400, depth 50, 784 dimensional input space, initialized with zero bias and random weights. In such a setting all the layerwise hyperplanes intersect the origin at their input space. We compute the LC around the input space origin using our method, for neighborhoods of varying radius  $r = \{0.0001, 0.001, 0.01, 0.1, 1, 10\}$  and dimensionality  $P = \{2, 10, 25, 50, 100, 200\}$ . For all the trials, our method recovers all the layerwise hyperplane intersections, even with a neighborhood dimensionality of  $P = 2$ .

**Non-Zero Bias Random MLP with shifting neighborhood.** For a randomly initialized MLP, we expect to see lower local complexity as we move away from the origin (Hanin & Rolnick, 2019). For this experiment we take a width 100 depth 18 MLP with input dimensionality  $d = 784$ , Leaky-ReLU activation with negative slope 0.01. We start by computing LC at the origin  $[0]^d$ , and linearly shift towards the vector  $[10]^d$ . We see that for all the settings, shifting away from the origin reduces LC. LC gets saturated with increasing  $P$ , showing that lower dimensional neighborhoods can be good enough for approximating LC. Increasing  $r$  on the other hand, increases LC and reduces LC variations between shifts, since the neighborhood becomes larger and LC becomes less local.

**Trained MLP comparison with SplineCam.** For non-linear MLPs, we compare with the exact computation method Splinecam (Humayun et al., 2023a). We take a depth 3 width 200 MLP and train it on MNIST for 100K training steps. For 20 different training checkpoints, we compute the local complexity in terms of the number of linear regions computed via SplineCam and number of hyperplane intersections via our proposed method. We compute the local complexity for 500 different training samples. For both our method and SplineCam we consider a radius of 0.001. For our method, we consider a neighborhood with dimensionality  $P = 25$ . We present the LC trajectories in Fig. 35. We can see that for both methods

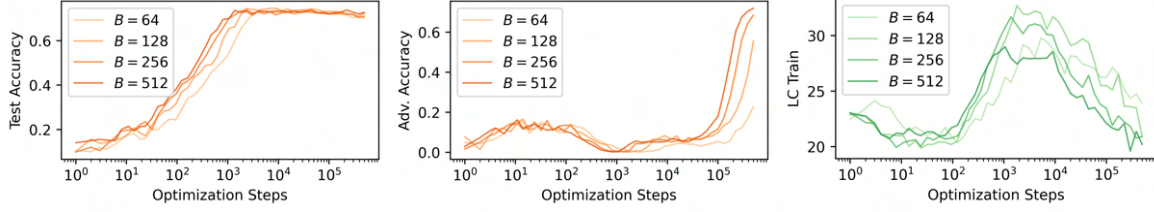


Figure 17. Increasing the batch-size expedites grokking. This indicates that reduced SGD noise allows region migration to occur earlier in training.

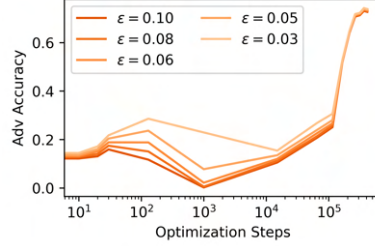


Figure 18. **Grokking stronger adversarial attacks.** We see that during delayed generalization, the robustness to Auto-Attack (Croce & Hein, 2020) also increases. This shows the universality of delayed robustness.

the local complexity follows a similar trend with a double descent behavior.

**Deformation of neighborhood by deep networks.** As mentioned in Appendix A, we compute the local complexity in a layerwise fashion by embedding a neighborhood  $\text{conv}(V)$  into the input space for any layer and computing the number of hyperplane intersections with  $\text{conv}(V^\ell)$ , where  $V^\ell$  is the embedded vertices at the input space of layer  $\ell$ . The approximation of local complexity is therefore subject to the deformation induced by each layer to  $\text{conv}(V)$ . To measure deformation by layers 1 to  $\ell - 1$ , we consider the undirected graph formed by the vertices  $V^\ell$  and compute the average eccentricity and diameter of the graphs (Xu et al., 2021). Eccentricity for any vertex  $v$  of a graph, is denoted by the maximum shortest path distance between  $v$  and all the connected vertices in the graph. The diameter is the maximum eccentricity over vertices of a graph. Recall from Appendix A that  $\text{conv}(V)$  where  $V = \{x \pm rv_p : p = 1 \dots P\}$  for an input space point  $x$ , is a cross-polytope of dimensionality  $P$ , where only two vertices are sampled from any of the orthogonal directions  $v_p$ . Therefore, all vertices share edges with each other except for pairs  $\{(x + rv_p, x - rv_p) : p = 1 \dots P\}$ . Given such connectivity, we compute the average eccentricity and diameter of neighborhoods  $\text{conv}(V^\ell)$  around 1000 training points from CIFAR10 for a trained CNN (Fig. 21). We see that for larger  $r$  both of the deformation metrics exponentially increase, where as for  $r \leq 0.014$  the deformation is lower and more stable. This shows that for lower  $r$  our LC approximation for deeper CNN networks would be better since the neighborhood does not get deformed significantly.

## B. Understanding Batch Normalization and its effect on the partition

Suppose the usual layer mapping is

$$z_{\ell+1} = a(W_\ell z_\ell + c_\ell), \quad \ell = 0, \dots, L - 1 \quad (7)$$

While a host of different DNN architectures have been developed over the past several years, modern, high-performing DNNs nearly universally employ *batch normalization* (BN) (Ioffe & Szegedy, 2015) to center and normalize the entries of the feature maps using four additional parameters  $\mu_\ell, \sigma_\ell, \beta_\ell, \gamma_\ell$ . Define  $z_{\ell,k}$  as  $k^{\text{th}}$  entry of feature map  $z_\ell$  of length  $D_\ell$ ,  $w_{\ell,k}$  as the  $k^{\text{th}}$  row of the weight matrix  $W_\ell$ , and  $\mu_{\ell,k}, \sigma_{\ell,k}, \beta_{\ell,k}, \gamma_{\ell,k}$  as the  $k^{\text{th}}$  entries of the BN parameter vectors  $\mu_\ell, \sigma_\ell, \beta_\ell, \gamma_\ell$ , respectively. Then we can write the BN-equipped layer  $\ell$  mapping extending (1) as

$$z_{\ell+1,k} = a\left(\frac{\langle w_{\ell,k}, z_\ell \rangle - \mu_{\ell,k}}{\sigma_{\ell,k}} \gamma_{\ell,k} + \beta_{\ell,k}\right), k = 1, \dots, D_\ell. \quad (8)$$

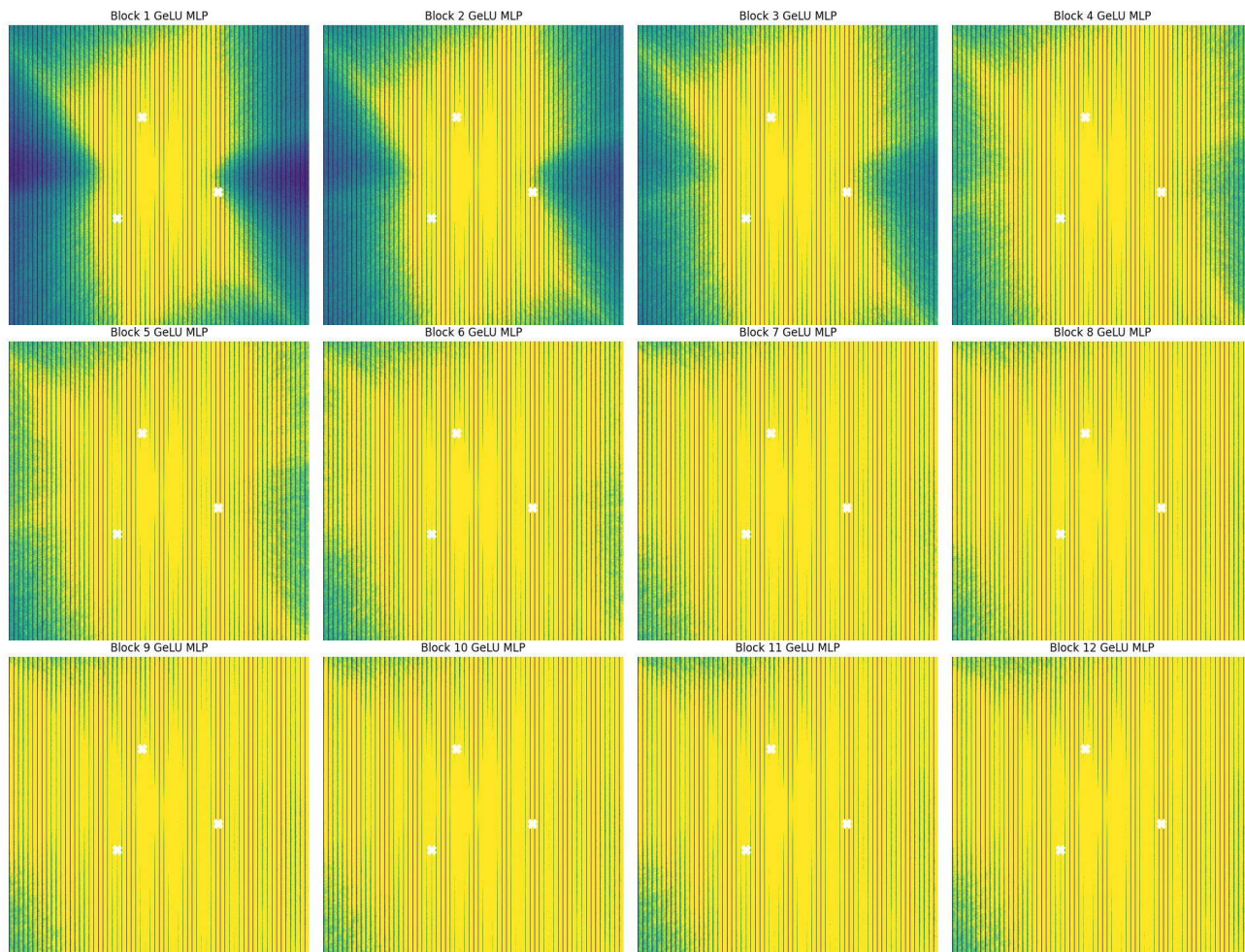


Figure 19. Token embedding space LC on a 2D subspace intersecting three random training points. We visualize the LC layerwise for the GeLU activated MLP layers inside each of the 12 blocks of the LLM for which we present training dynamics in Figure 9. The LC is computed after 197 optimization steps, during the peak ascent. We see that LC on this subspace is very high especially close to the data points. LC values are clamped to a maximum of 150.

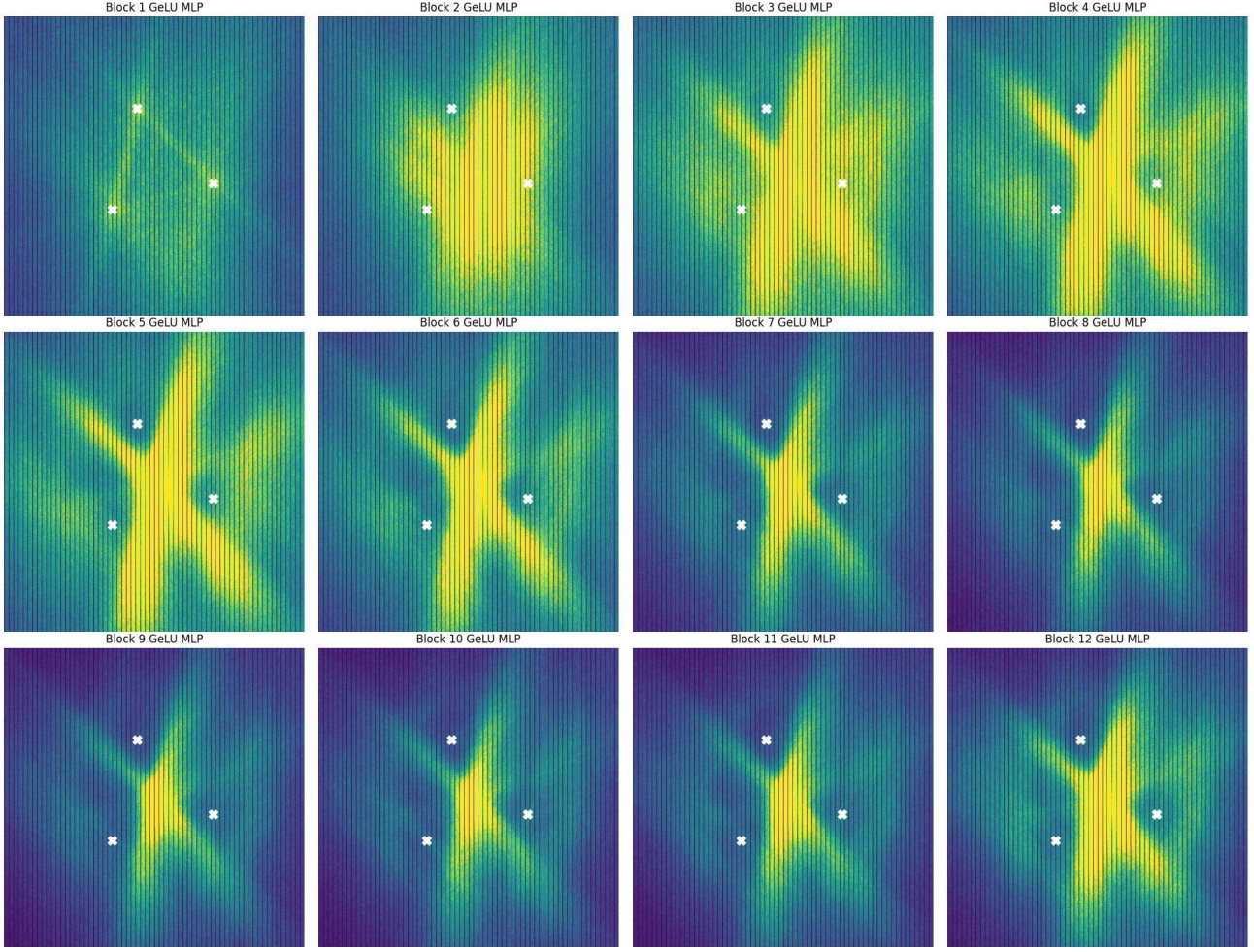


Figure 20. Token embedding space LC on a 2D subspace intersecting three random training points. We visualize the LC layerwise for the GeLU activated MLP layers inside each of the 12 blocks of the LLM for which we present training dynamics in Figure 9. The LC is computed after 372759 optimization steps, therefore during the second LC descent. We see that LC on this subspace is concentrated away from the training points, especially for the deeper layers. This indicates that region migration occurs in LLMs as well, leading to delayed robustness. LC values are clamped to a maximum of 150.

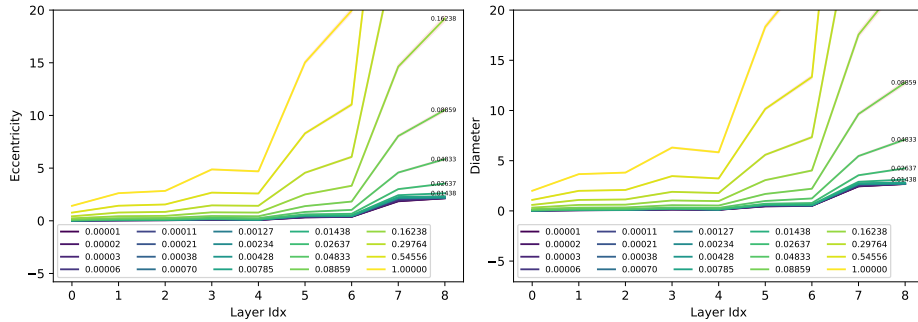


Figure 21. Change of avg. eccentricity and diameter (Xu et al., 2021) of the input space neighborhood by different layers of a CNN trained on the CIFAR10 dataset. For different sampling radius  $r$  of the sampled input space neighborhood  $V$ , the change of eccentricity and diameter denotes how much deformation the neighborhood undergoes between layers. Here, layer 0 corresponds to the input space neighborhood. Numbers are averaged over neighborhoods sampled for 1000 training points from CIFAR10. For larger radius the deformation increases with depth exponentially. For  $r \leq 0.014$  deformation is lower, indicating that smaller radius neighborhoods are reliable for LC computation on deeper networks. Confidence interval shown in red, is almost imperceptible.

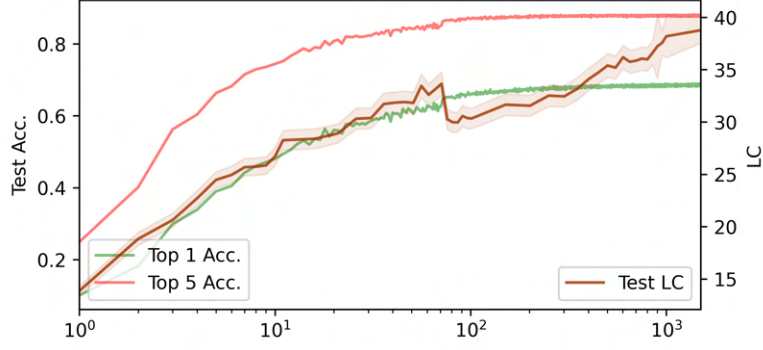


Figure 22. Training a ResNet18 with batchnorm on Imagenet Full. LC is computed only on test points using 1000 test set samples. Computing LC 1000 samples takes approx. 28s on an RTX 8000.

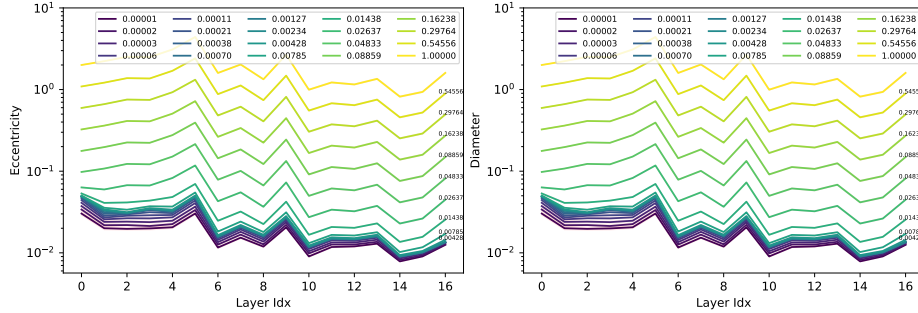


Figure 23. Change of avg. eccentricity and diameter (Xu et al., 2021) of the input space neighborhood by different layers of a ResNet18 trained on the CIFAR10 dataset, similar to the setting of Fig. 21. Resnet deforms the input neighborhood by reducing the avg. eccentricity and diameter of the neighborhood graphs. For  $r \leq 0.014$  deformation is lower, indicating that smaller radius neighborhoods are reliable for LC computation on deeper networks.

The parameters  $\mu_\ell, \sigma_\ell$  are computed as the element-wise mean and standard deviation of  $\mathbf{W}_\ell \mathbf{z}_\ell$  for each mini-batch during training and for the entire training set during testing. The parameters  $\beta_\ell, \gamma_\ell$  are learned along with  $\mathbf{W}_\ell$  via SGD.<sup>4</sup> For each mini-batch  $\mathbb{B}$  during training, the BN parameters  $\mu_\ell, \sigma_\ell$  are *calculated directly* as the mean and standard deviation of the current mini-batch feature maps  $\mathcal{B}_\ell$

$$\mu_\ell \leftarrow \frac{1}{|\mathbb{B}_\ell|} \sum_{\mathbf{z}_\ell \in \mathbb{B}_\ell} \mathbf{W}_\ell \mathbf{z}_\ell, \quad \sigma_\ell \leftarrow \sqrt{\frac{1}{|\mathbb{B}_\ell|} \sum_{\mathbf{z}_\ell \in \mathbb{B}_\ell} (\mathbf{W}_\ell \mathbf{z}_\ell - \mu_\ell)^2}, \quad (9)$$

where the right-hand side square is taken element-wise. After SGD learning is complete, a final fixed “test time” mean  $\bar{\mu}_\ell$  and standard deviation  $\bar{\sigma}_\ell$  are computed using the above formulae over all of the training data,<sup>5</sup> i.e., with  $\mathbb{B}_\ell = \mathbb{X}_\ell$ .

The Euclidean distance from a point  $\mathbf{v}$  in layer  $\ell$ ’s input space to the layer’s  $k^{\text{th}}$  hyperplane  $\mathbb{H}_{\ell,k}$  is easily calculated as

$$d(\mathbf{v}, \mathbb{H}_{\ell,k}) = \frac{|\langle \mathbf{w}_{\ell,k}, \mathbf{v} \rangle - \mu_{\ell,k}|}{\|\mathbf{w}_{\ell,k}\|_2} \quad (10)$$

as long as  $\|\mathbf{w}_{\ell,k}\| > 0$ .

Then, the average squared distance between  $\mathbb{H}_{\ell,k}$  and a collection of points  $\mathbb{V}$  in layer  $\ell$ ’s input space is given by

$$\mathbf{L}_k(\mu_{\ell,k}, \mathbb{V}) = \frac{1}{|\mathbb{V}|} \sum_{\mathbf{v} \in \mathbb{V}} d(\mathbf{v}, \mathbb{H}_{\ell,k})^2 = \frac{\sigma_{\ell,k}^2}{\|\mathbf{w}_{\ell,k}\|_2^2}, \quad (11)$$

### C. What affects the robust partition? *Reprise*

**Depth.** In Figure 27 we plot LC during training on MNIST for Fully Connected Deep Networks with depth in  $\{2, 3, 4, 5\}$  and width 200. In each plot, we show both LC as well as train-test accuracy. For all the depths, the accuracy on both the train and test sets peak during the first descent phase. During the ascent phase, we see that the train LC has a sharp ascent while the test and random LC do not.

The difference as well as the sharpness of the ascent is reduced when increasing the depth of the network. This is visible for both fine and coarse  $r$  scales. For the shallowest network, we can see a second descent in the coarser scale but not in the finer  $r$  scale. This indicates that for the shallow network some regions closer to the training samples are retained during later stages of training. One thing to note is that during the ascent and second descent phase, there is a clear distinction between the train and test LC. *This is indicative of membership inference fragility especially during latter phases of training.* It has previously been observed in membership inference literature (Tan et al., 2023), where early stopping has been used as a regularizer for membership inference. We believe the LC dynamics can shed a new light towards membership inference and the role of network complexity/capacity.

In Figure 12, we plot the local complexity during training for CNNs trained on CIFAR10 with varying depths with and without batch normalization. The CNN architecture comprises of only convolutional layers except for one fully connected layer before output. Therefore when computing LC, we only take into account the convolutional layers in the network. Contrary to the MNIST experiments, we see that in this setting, the train-test LC are almost indistinguishable throughout training. We can see that the network train and test accuracy peaks during the ascent phase and is sustained during the second descent. It can also be noticed that increasing depth increases the max LC during the ascent phase for CNNs which is contrary to what we saw for fully connected networks on MNIST. The increase of density during ascent is all over the data manifold, contrasting to just the training samples for fully connected networks.

In Appendix, we present layerwise visualization of the LC dynamics. We see that shallow layers have sharper peak during ascent phase, with distinct difference between train and test. For deeper layers however, the train vs test LC difference is negligible.

**Width.** In Figure 14 we present results for a fully connected DNN with depth 3 and width  $\{20, 100, 500, 1000, 2000\}$ . Networks with smaller width start from a low LC at initialization compared to networks that are wider. Therefore for small width networks the initial descent becomes imperceptible. We see that as we increase width from 20 to 1000 the ascent

<sup>4</sup>Note that the DNN bias  $\mathbf{c}_\ell$  from (1) has been subsumed into  $\mu_\ell$  and  $\beta_\ell$ .

<sup>5</sup>or more commonly as an exponential moving average of the training mini-batch values.

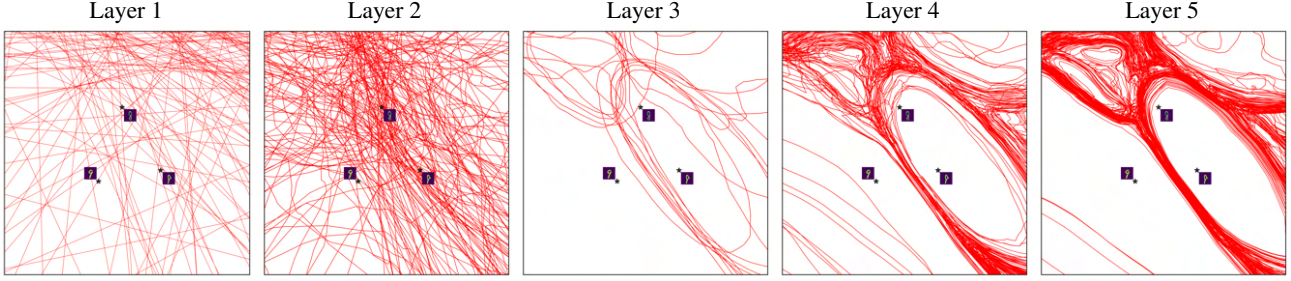


Figure 24. Layerwise visualization of the input space partition for a 2D domain passing through a training set triad, after robust partition formation. The partition is visualized for an MLP with depth 6 and width 200, trained on 1000 samples from MNIST, similar to the setting described in Figure 2. We see that deeper layer neurons partake more in the formation of the robust partition, compared to shallower layers. This is due to the fact that deeper layer neurons can be more localized in the input space due to the non-linearity induced by preceding layers.

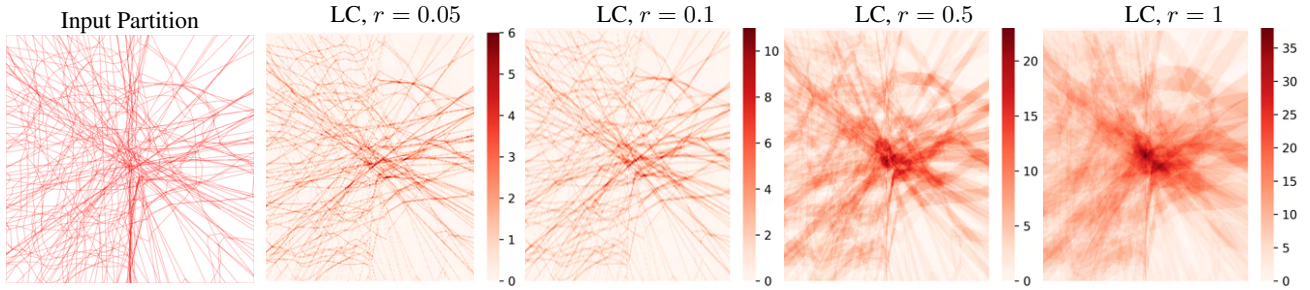


Figure 25. Input space partition computed analytically via SplineCam (Humayun et al., 2023b) for the 2D toy setting presented in Figure 3 (left). Regions are colored by white and knots are colored by red. The partition is computed for the input space domain  $[-10, 10]^2$ , induced by an MLP of depth 5 and width 30. We take a meshgrid of  $300 \times 300$  points over the input domain, and measure the local complexity at each point with radius,  $r \in \{0.05, 0.1, 0.5, 1\}$  (rest). We see that our proposed method can locate the non-linearities for small  $r$ . As  $r$  is increased our method provides a coarser estimate of the local density of non-linearities, i.e., number of non-linearities intersecting the a fixed volume defined by the local neighborhood.

phase starts earlier as well as reaches a higher maximum LC. However overparameterizing the network by increasing the width further to 2000, reduces the max LC during ascent, therefore reducing the crowding of neurons near training samples. *This is a possible indication of how overparameterization performs implicit regularization (Kubo et al., 2019), by reducing non-linearity or local complexity concentration around training samples.*

**Weight Decay** regularizes a neural network by reducing the norm of the network weights, therefore reducing the per region slope norm as well. We train a CNN with depth 5 and width 32 and varying weight decay. In Fig. 30 we present the train and random LC for our experiments. We can see that increasing weight decay also delays or removes the second descent in training LC. Moreover, strong weight decay also reduces the duration of ascent phase, as well as reduces the peak LC during ascent. This is dissimilar from BN, which removes the second descent but increases LC overall.

**Batch Normalization.** It has previously been shown that Batch normalization (BN) regularizes training by dynamically updating the normalization parameters for every mini-batch, therefore increasing the noise in training (Garbin et al., 2020). In fact, we recall that BN replaces the per-layer mapping from Equation (1) by centering and scaling the layer’s pre-activation and adding back the learnable bias  $b^{(\ell)}$ . The centering and scaling statistics are computed for each mini-batch. After learning is complete, a final fixed “test time” mean  $\bar{\mu}^{(\ell)}$  and standard deviation  $\bar{\sigma}^{(\ell)}$  are computed using the training data. Of key interest to our observation is a result tying BN to the position in the input space of the partition region from (Balestriero & Baraniuk, 2022). In particular, it was proved that at each layer  $\ell$  of a DN, BN explicitly adapts the partition so that the partition boundaries are as close to the training data as possible. This is confirmed by our experiments in Fig. 12 we present results for CNN trained on CIFAR10, with and without BN.

## D. Extra Figures

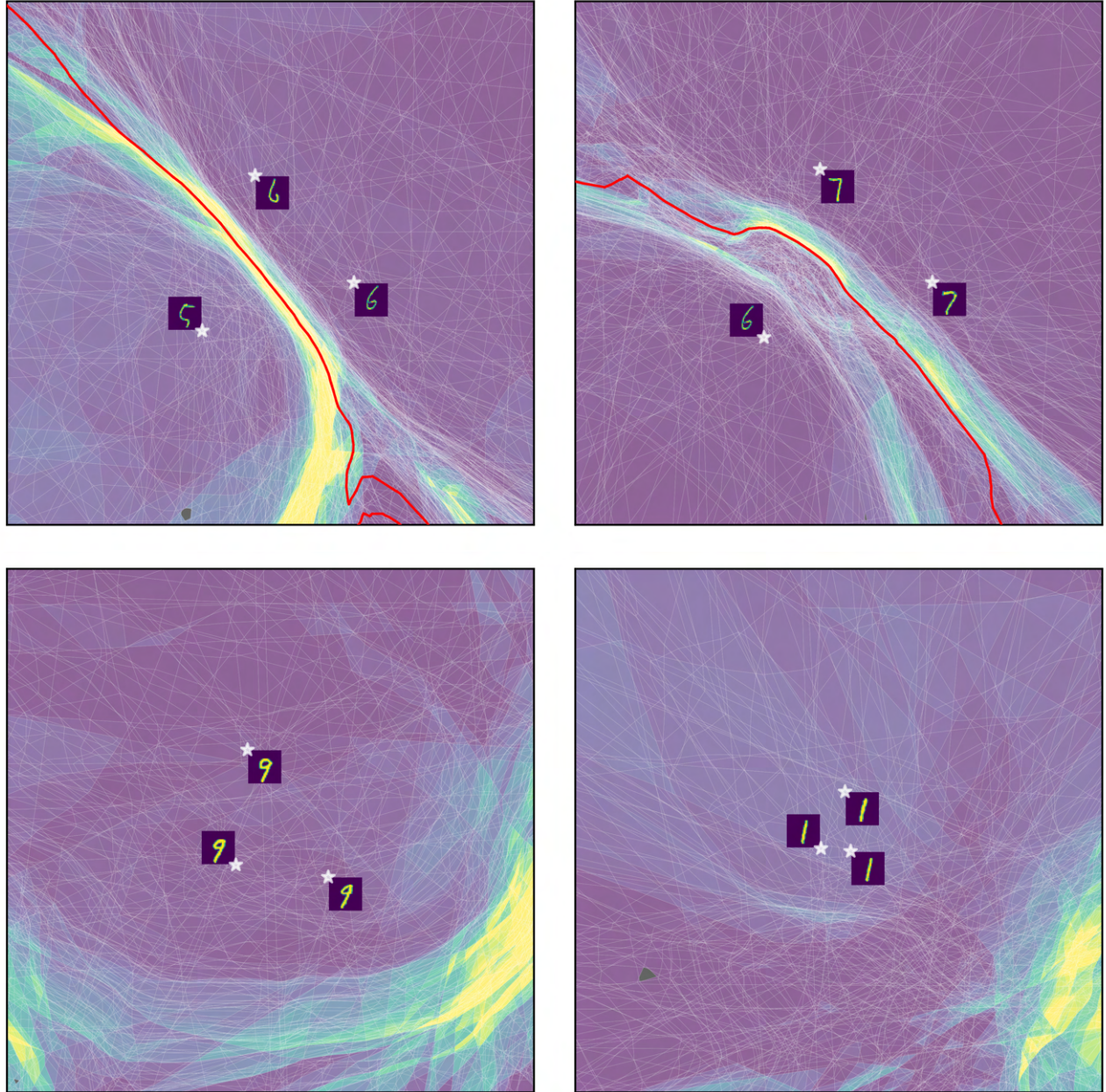


Figure 26. Partition visualization for 2D domains localized around the decision boundary (top) and away from the decision boundary (bottom) for the grokking setup presented in Figure 7. All the plots are show for the optimization step 95381. Number of regions in the partition for top-right, top-left, bottom-right, and bottom-left are 123156, 88362, 33273, and 32018 respectively. The domain used for all of the plots has the same area/volume. Therefore, close to the decision boundary, the region density is much higher compared to away from the decision boundary. This is evidence of region migration happening during the latter phases of training.

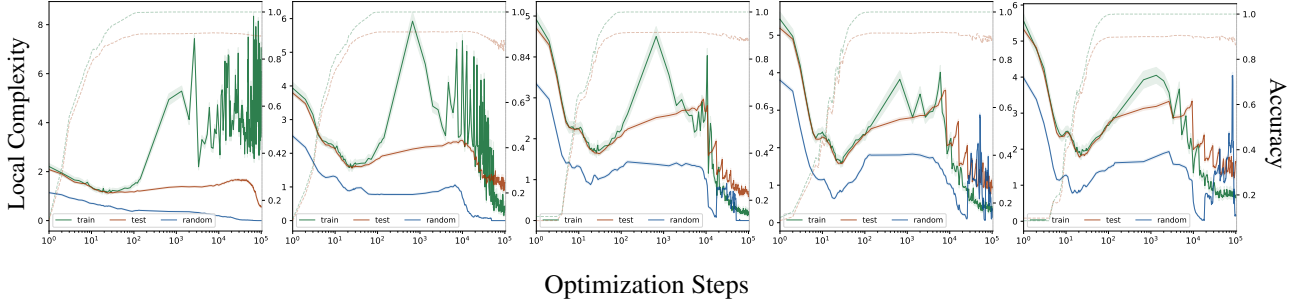


Figure 27. MLP with width 200 and varying depth being trained on 1000 samples from MNIST. Increasing the depth of the network decreases the sharpness of the LC peak during ascend phase. Deeper networks also tend to have a sharper decline in the training LC during region migration.

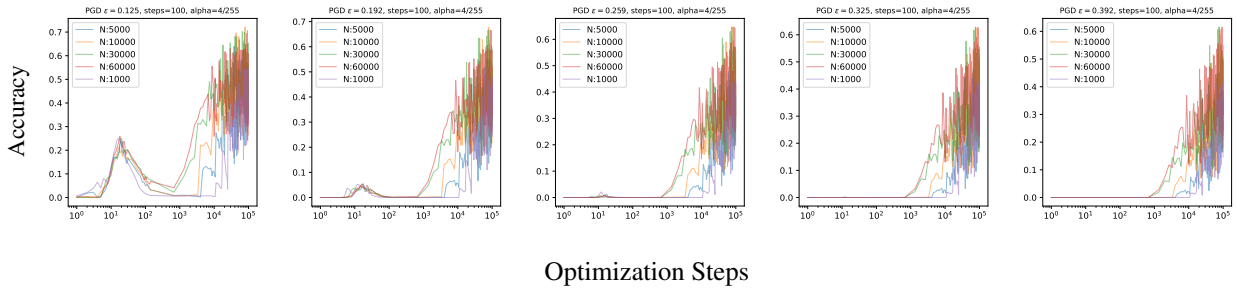


Figure 28. For an MLP with depth 4 and width 200, we train with varying training set sizes and evaluate the adversarial performance after each training iteration. We see that with increasing dataset size, the network groks earlier in time, as can be visible in the adversarial grokking curves for all the different epsilon values.

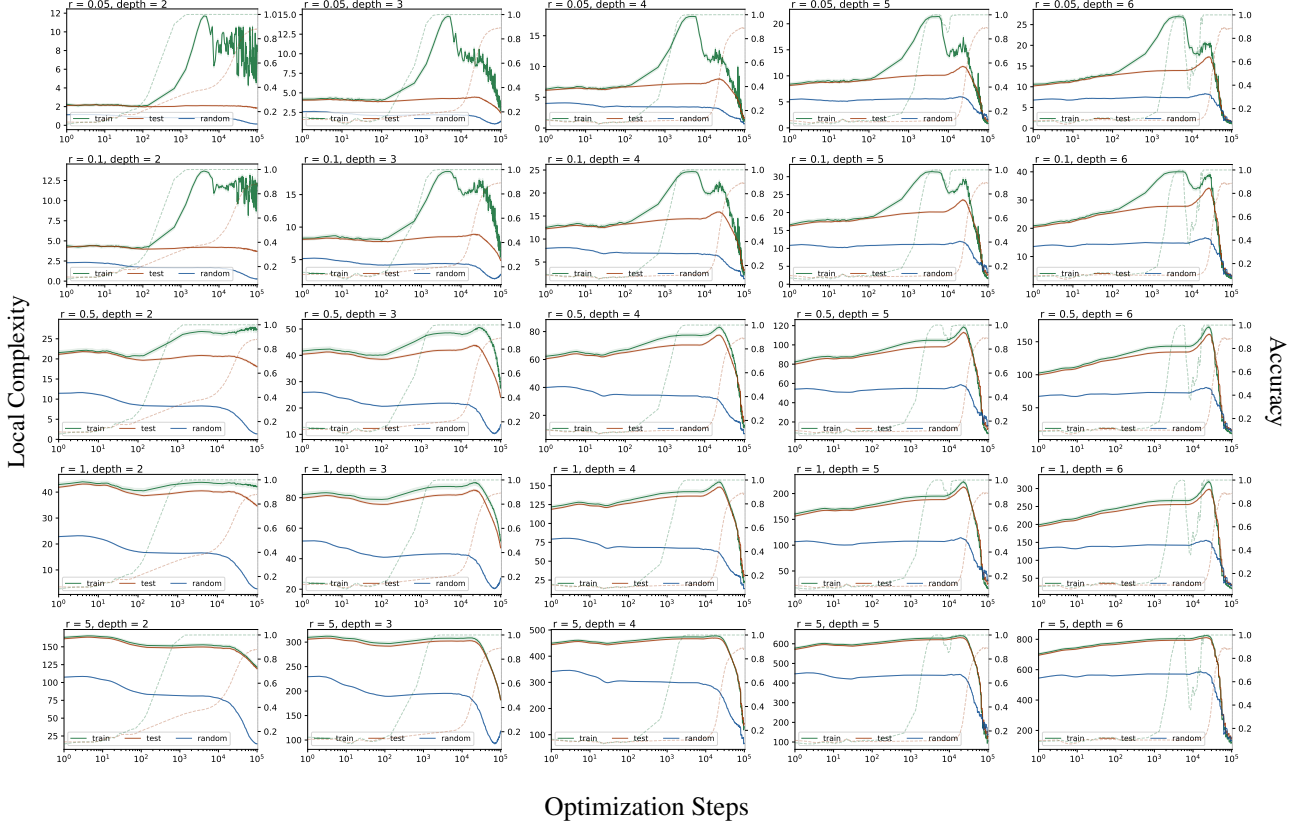


Figure 29. Training a 200 width MLP on MNIST with initialization scaling of 8 and varying depths. Along the row, we consider larger and larger radius neighborhoods for local complexity approximation.

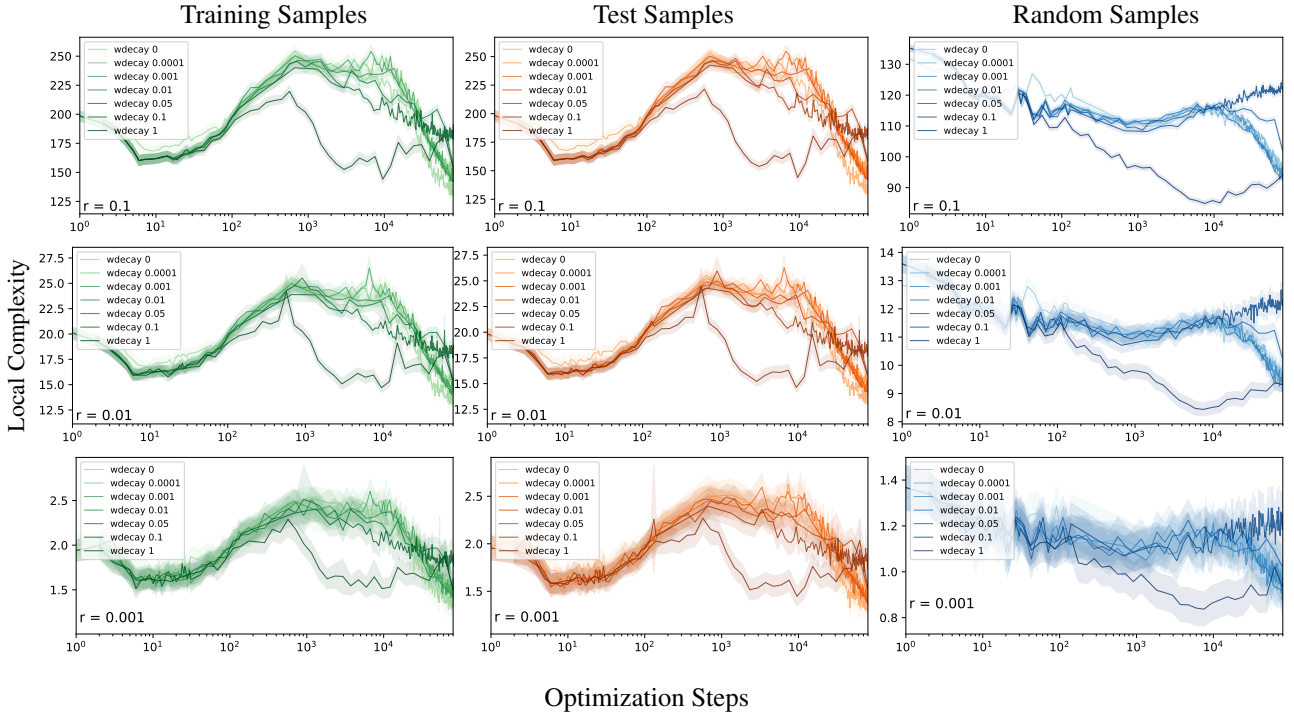


Figure 30. Local complexity dynamics training an MLP on MNIST with weight decay

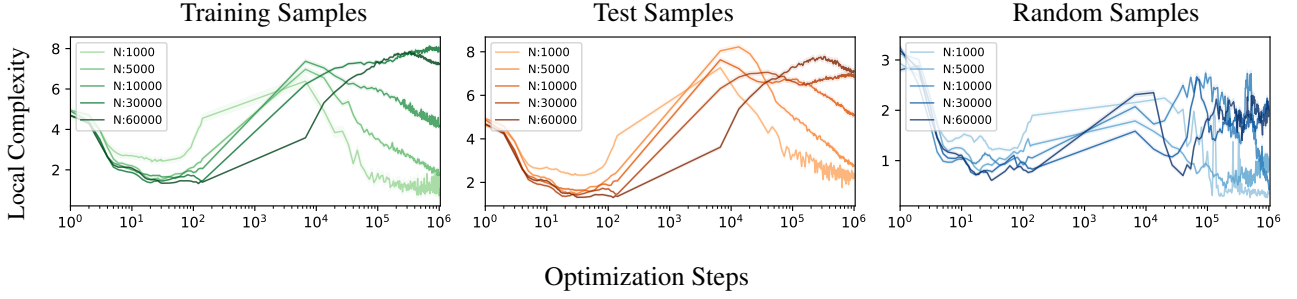


Figure 31. **Increasing the volume of randomly labeled training data.** Continued from Figure 13. Increasing the number of randomly labeled training samples delays the ascent phase of the LC training dynamics for both training and test samples. For random samples the behavior is not affected as much.

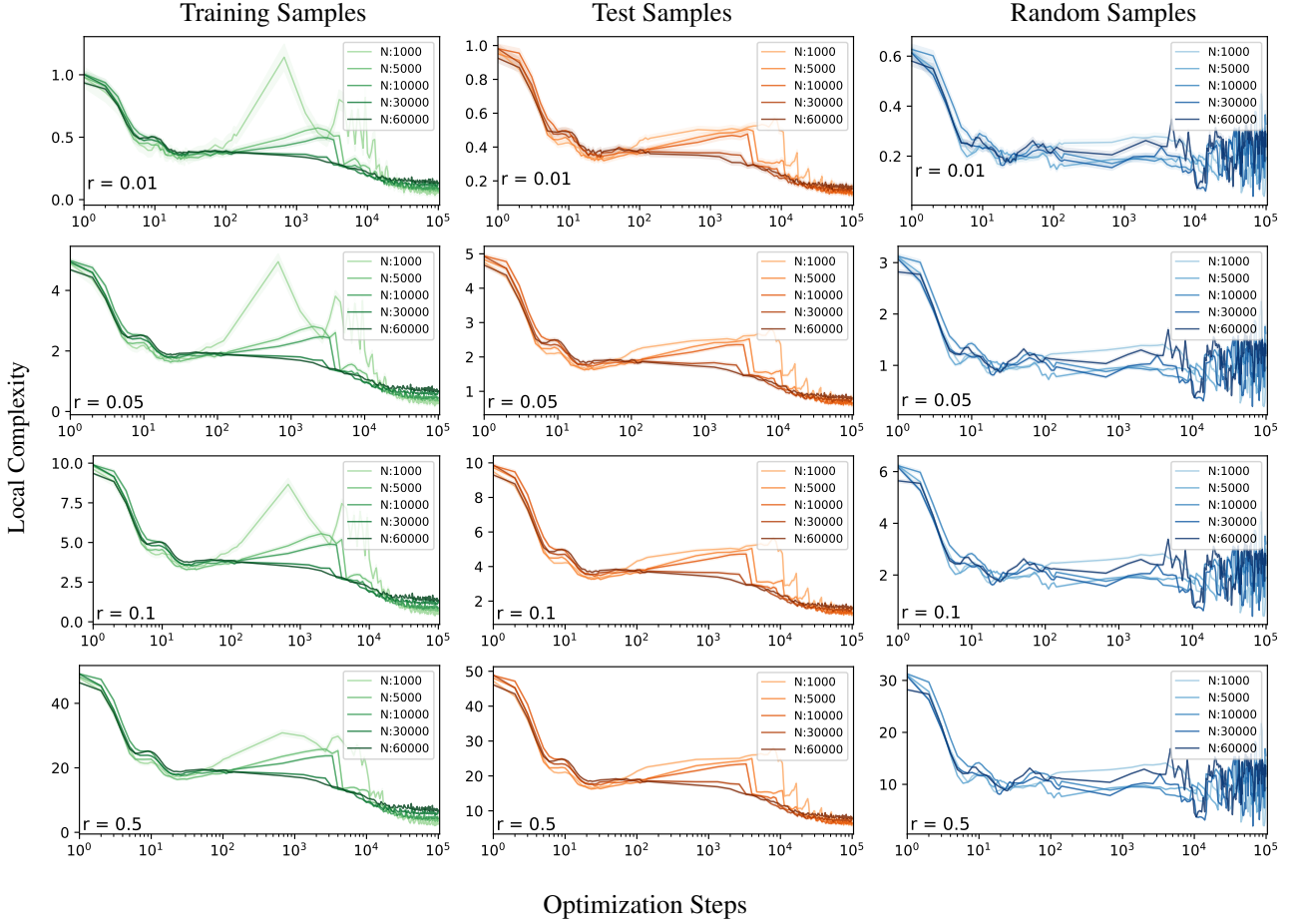


Figure 32. **Increasing training data size expedites region migration.** Local complexity dynamics training an MLP on MNIST with weight decay. Robustness plots presented in Figure 28.

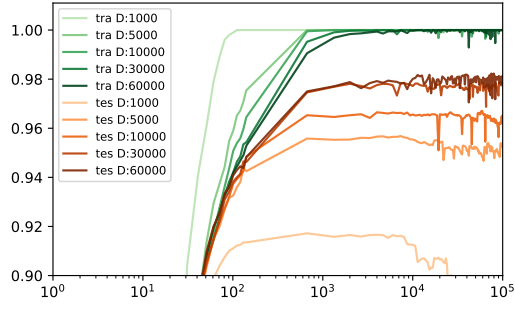


Figure 33. Training and Test accuracy for the different dataset sizes presented in Figure 32.

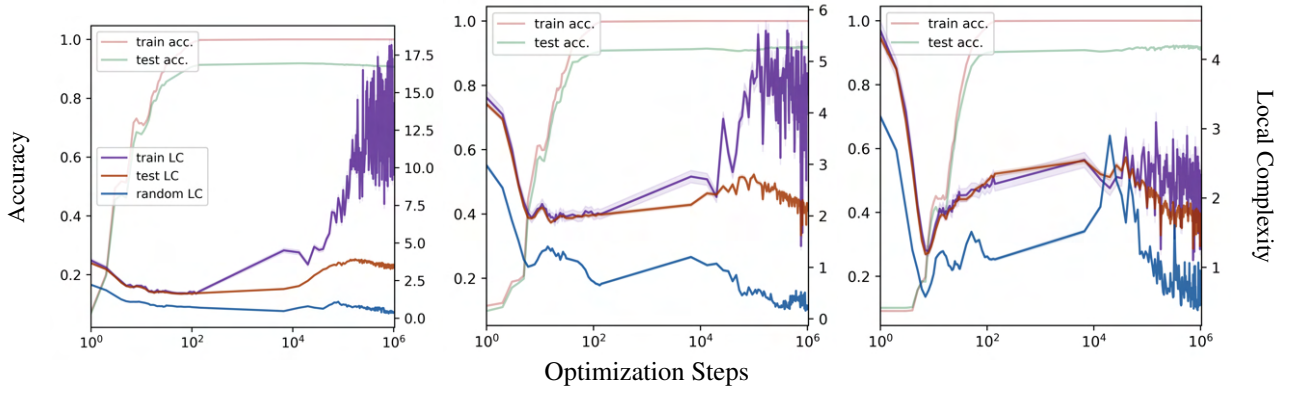


Figure 34. LC dynamics for a GeLU-MLP with width 200 and depth  $\{3, 4, 5\}$  presented from left to right. LC is calculated at 1000 training points and 10000 test and random points during training on MNIST.

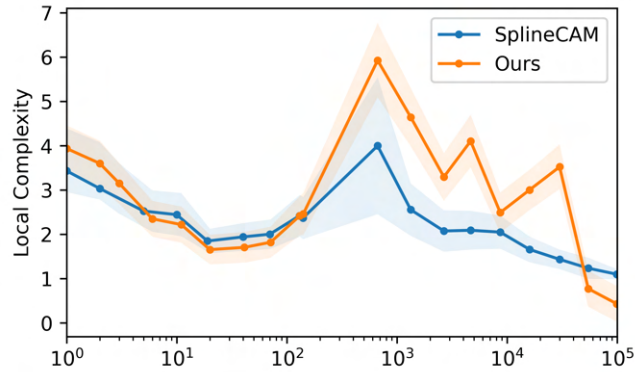


Figure 35. Comparing the local complexity measured in terms of the number of linear regions computed exactly by SplineCAM (Humayun et al., 2023a) and number of hyperplane cuts by our proposed method. Both methods exhibit the double descent behavior.

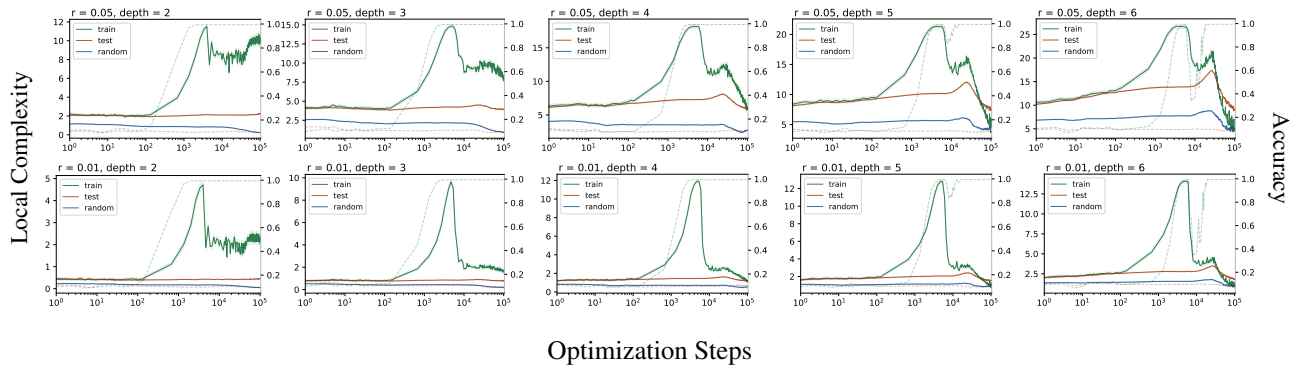


Figure 36. Random label radius and depth Sweep