

Generalization in Deep Neural Networks

Peter Bartlett

UC Berkeley
Computer Science and Statistics

October 5, 2017

Generalization in Neural Networks

- What determines the statistical complexity of a deep network?
 - VC theory: Number of parameters
 - Margins analysis: **Size of parameters**
 - Understanding generalization failures

Neural Networks for Classification

Neural network computes $f : \mathbb{R}^d \rightarrow \mathbb{R}$.

Directed acyclic graph with one output node, nodes compute

$$(z_1, \dots, z_m) \mapsto \sigma \left(\sum_{i=1}^m w_i z_i + w_0 \right),$$

where $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ is a nonlinear function, such as

$$\sigma(\alpha) = \frac{1}{1 + e^{-\alpha}}, \quad \text{or} \quad \sigma(\alpha) = \max\{0, \alpha\}.$$

Parameters (w 's) adjusted by gradient descent to minimize an objective function on training examples, such as

$$\sum_{i=1}^n (f(x_i) - y_i)^2.$$

- Assume network maps to $\{-1, 1\}$.
(Threshold its output)
- Data generated by a probability distribution P on $X \times \{-1, 1\}$.
- Want to choose a function f such that with high probability $P(f(x) \neq y)$ is small (near optimal).

Theorem (Vapnik and Chervonenkis)

Suppose $F \subseteq \{-1, 1\}^X$.

For every prob distribution P on $X \times \{-1, 1\}$,
with probability $1 - \delta$ over n iid examples $(x_1, y_1), \dots, (x_n, y_n)$,
every f in F satisfies

$$P(f(x) \neq y) \leq \frac{1}{n} |\{i : f(x_i) \neq y_i\}| + \left(\frac{c}{n} (\text{VCdim}(F) + \log(1/\delta)) \right)^{1/2}.$$

- For uniform bounds (that is, for all distributions and all $f \in F$, proportions are close to probabilities), this inequality is tight within a constant factor.
- For neural networks, VC-dimension:
 - increases with number of parameters
 - depends on nonlinearity and depth

VC-Dimension of Neural Networks

Theorem

Consider the class F of $\{-1, 1\}$ -valued functions computed by a network with L layers, p parameters, and k computation units with the following nonlinearities:

- 1 Piecewise constant (linear threshold units):
$$\text{VCdim}(F) = \tilde{O}(p).$$

(Baum and Haussler, 1989)
- 2 Piecewise linear (ReLU):
$$\text{VCdim}(F) = \tilde{O}(pL).$$

(B., Harvey, Liaw, Mehrabian, 2017)
- 3 Piecewise polynomial:
$$\text{VCdim}(F) = \tilde{O}(pL^2).$$

(B., Maierov, Meir, 1998)
- 4 Sigmoid:
$$\text{VCdim}(F) = \tilde{O}(p^2 k^2).$$

(Karpinsky and MacIntyre, 1994)

Generalization in Neural Networks: Number of Parameters

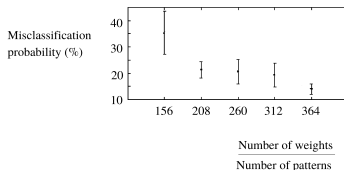
NIPS 1996

Experimental Results

Neural networks with many parameters, trained on small data sets, sometimes generalize well.

Eg: Face recognition (Lawrence *et al*, 1996)

$m = 50$ training patterns.



Generalization in Neural Networks

- What determines the statistical complexity of a deep network?
 - VC theory: Number of parameters
 - **Margins analysis: Size of parameters**
 - Understanding generalization failures

Generalization: Margins and Size of Parameters

Theorem (B., 1996)

1. With high probability over n training examples $(X_1, Y_1), \dots, (X_n, Y_n) \in \mathcal{X} \times \{\pm 1\}$, every $f \in \mathcal{F} \subset \mathbb{R}^{\mathcal{X}}$ has

$$\Pr(\text{sign}(f(X)) \neq Y) \leq \frac{1}{n} \sum_{i=1}^n \mathbb{1}[Y_i f(X_i) \leq \gamma] + \tilde{O} \left(\sqrt{\frac{\text{fat}_{\mathcal{F}}(\gamma)}{n}} \right).$$

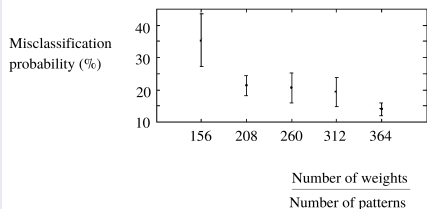
2. If functions in $\mathcal{F}$ are computed by two-layer sigmoid networks with each unit's weights bounded in 1-norm, that is, $\|w\|_1 \leq B$, then

$$\text{fat}_{\mathcal{F}}(\gamma) = \tilde{O}((B/\gamma)^2).$$

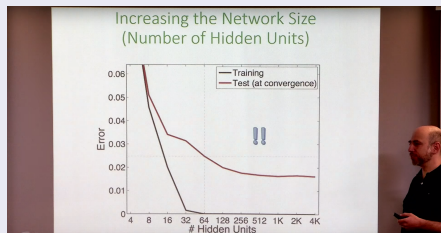
- The bound depends on the margin loss plus an error term.
- Minimizing quadratic loss or cross-entropy loss leads to large margins.
- $\text{fat}_{\mathcal{F}}(\gamma)$ is a scale-sensitive version of VC-dimension. Unlike the VC-dimension, it need not grow with the number of parameters.

Generalization: Margins and Size of Parameters

1996: Sigmoid networks



2017: Deep ReLU networks



simons.berkeley.edu

- Qualitative behavior explained by small weights theorem.

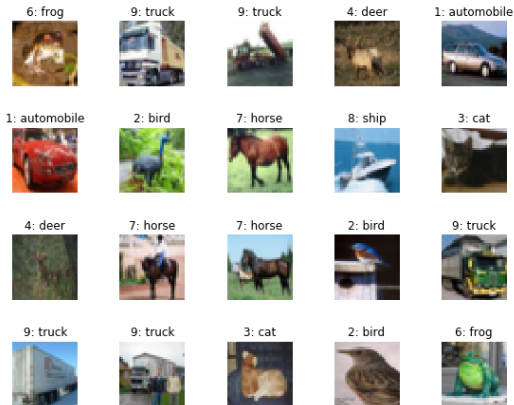
- How to measure the complexity of a ReLU network?

Generalization in Neural Networks

- What determines the statistical complexity of a deep network?
 - VC theory: Number of parameters
 - Margins analysis: Size of parameters
 - **Understanding generalization failures**

Explaining Generalization Failures

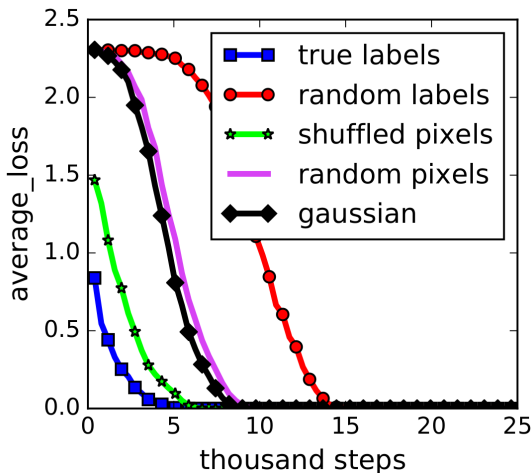
CIFAR10



<http://corochann.com/>

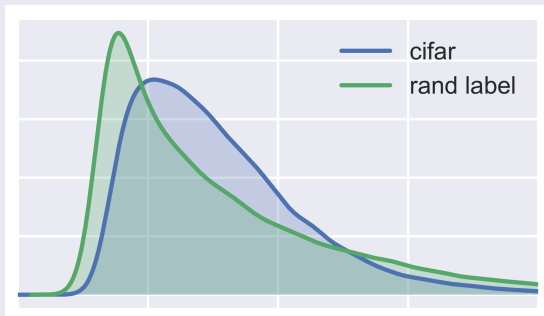
Explaining Generalization Failures

Stochastic Gradient Training Error on CIFAR10



Explaining Generalization Failures

Training margins on CIFAR10 with true and random labels



- How does this match the large margin explanation?
- Need to account for the *scale* of the neural network functions.
- What is the appropriate notion of the size of these functions?

Generalization in Deep Networks

Spectrally-normalized margin bounds for neural networks.
B., Dylan J. Foster, Matus Telgarsky, 2017.
arXiv:1706.08498



Matus Telgarsky
UIUC



Dylan Foster
Cornell

Generalization in Deep Networks

New results for generalization in deep ReLU networks

- Measuring the size of functions computed by a network of ReLUs. (c.f. sigmoid networks: the output y of a layer has $\|y\|_\infty \leq 1$, so $\|w\|_1 \leq B$ keeps the scale under control.)
- Large multiclass versus binary classification.

Definitions

- Consider operator norms: For a matrix A_i ,

$$\|A_i\|_* := \sup_{\|x\| \leq 1} \|A_i x\|.$$

- Multiclass margin function for $f : \mathcal{X} \rightarrow \mathbb{R}^m$, $y \in \{1, \dots, m\}$:

$$M(f(x), y) = f(x)_y - \max_{i \neq y} f(x)_i.$$

Generalization in Deep Networks

Theorem

With high probability, every f_A with $R_A \leq r$ satisfies

$$\Pr(M(f_A(X), Y) \leq 0) \leq \frac{1}{n} \sum_{i=1}^n 1[M(f_A(X_i), Y_i) \leq \gamma] + \tilde{O}\left(\frac{rL}{\gamma\sqrt{n}}\right).$$

Definitions

Network with L layers, parameters $A_1, \dots, A_L$:

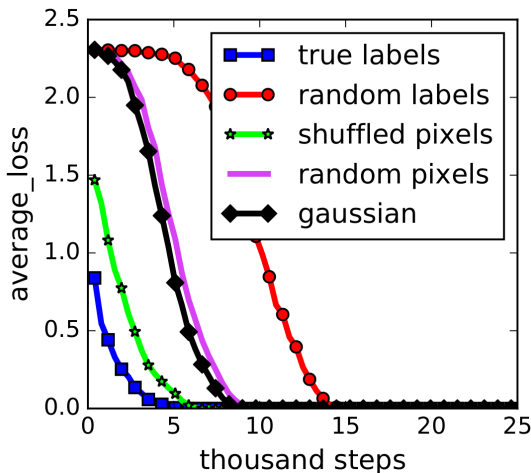
$$f_A(x) := \sigma(A_L \sigma_{L-1}(A_{L-1} \cdots \sigma_1(A_1 x) \cdots)).$$

Scale of f_A : $R_A := \prod_{i=1}^L \|A_i\|_* \sqrt{\sum_{i=1}^L \frac{\|A_i\|_F}{\|A_i\|_*}}$.

(Assume σ_i is 1-Lipschitz, inputs normalized.)

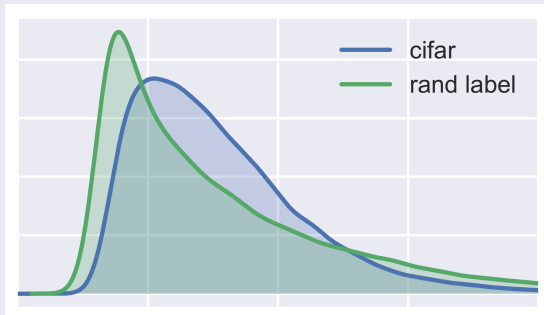
Explaining Generalization Failures

Stochastic Gradient Training Error on CIFAR10



Explaining Generalization Failures

Training margins on CIFAR10 with true and random labels

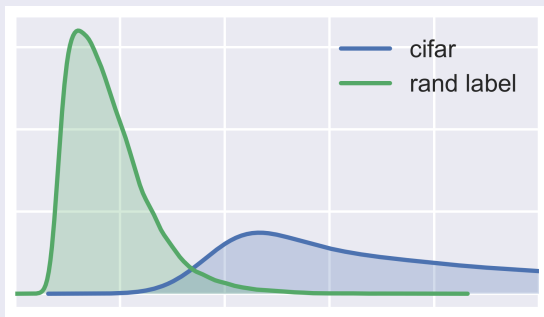


- How does this match the large margin explanation?

Explaining Generalization Failures

If we rescale the margins by R_A (the scale parameter):

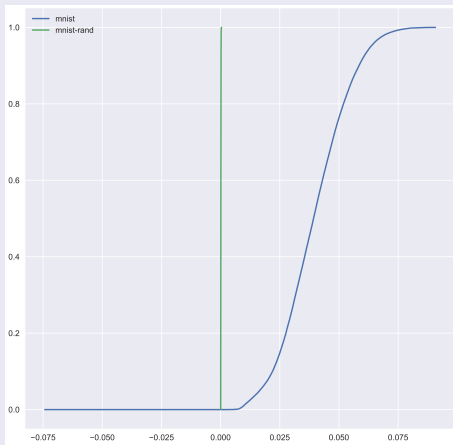
Rescaled margins on CIFAR10



Explaining Generalization Failures

If we rescale the margins by R_A (the scale parameter):

Rescaled margins on MNIST



Generalization in Deep Networks

Theorem

With high probability, every f_A with $R_A \leq r$ satisfies

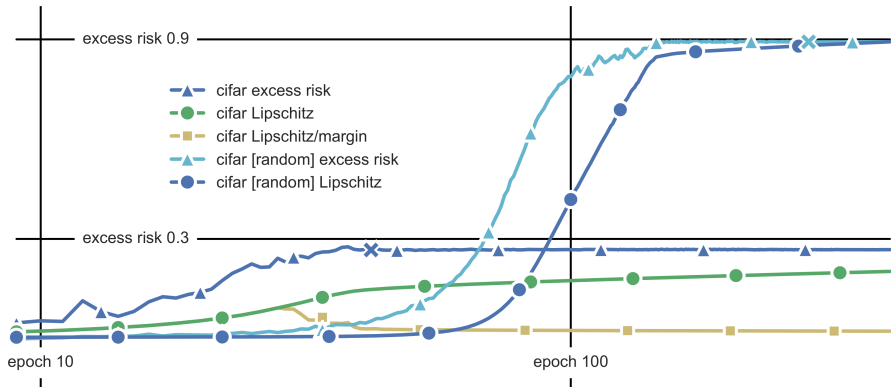
$$\Pr(M(f_A(X), Y) \leq 0) \leq \frac{1}{n} \sum_{i=1}^n 1[M(f_A(X_i), Y_i) \leq \gamma] + \tilde{O}\left(\frac{rL}{\gamma\sqrt{n}}\right).$$

Network with L layers, parameters $A_1, \dots, A_L$:

$$f_A(x) := \sigma(A_L \sigma_{L-1}(A_{L-1} \cdots \sigma_1(A_1 x) \cdots)).$$

Scale of f_A : $R_A := \prod_{i=1}^L \|A_i\|_* \sqrt{\sum_{i=1}^L \frac{\|A_i\|_F}{\|A_i\|_*}}$.

Explaining Generalization Failures



Generalization in Neural Networks

- With appropriate normalization, the margins analysis is qualitatively consistent with the generalization performance.
- Lower bounds?
- Regularization: explicit control of operator norms?
- Role of depth?
- Interplay with optimization?
- Residual networks: Close to identity.