



AALBORG
UNIVERSITET

PAC-Chernoff Bounds: Understanding Generalization in the Interpolation Regime

Andrés R. Masegosa Luis A. Ortega



Universidad Autónoma
de Madrid

Motivations

- Why current machine learning techniques find overparameterized **interpolators with strong generalization** performance is an **open question**.
- Generalization bounds that solely depend on the training data are **provably vacuous** for overparameterized model classes; unable to explain generalization.
- Explaining the generalization of overparametrized interpolators require **new tools**.

Contributions

- A **perfectly tight distribution-dependent PAC-Chernoff bound** for **interpolators**, even in over-parameterized models.
- A **theoretical framework** that explains why some interpolators generalize well, while others do not, based on a novel **characterization of smoothness**.
- We explain why regularization, data augmentation, invariant architectures, and over-parameterization, produce smoother interpolators with superior generalization.

The Rate Function

Chernoff Theorem. For any fixed $\theta \in \Theta$ and $a > 0$, it satisfies

$$\mathbb{P}_{D \sim \nu^n} \left(L(\theta) - \hat{L}(D, \theta) \geq a \right) \leq e^{-n\mathcal{I}(a)}.$$

with

$$\mathcal{I}(a) = \sup_{\lambda > 0} \lambda a - J_{\theta}(\lambda) \quad \text{and} \quad J_{\theta}(\lambda) = \ln \mathbb{E}_{\nu} \left[e^{\lambda(L(\theta) - \ell(\mathbf{y}, \mathbf{x}, \theta))} \right],$$

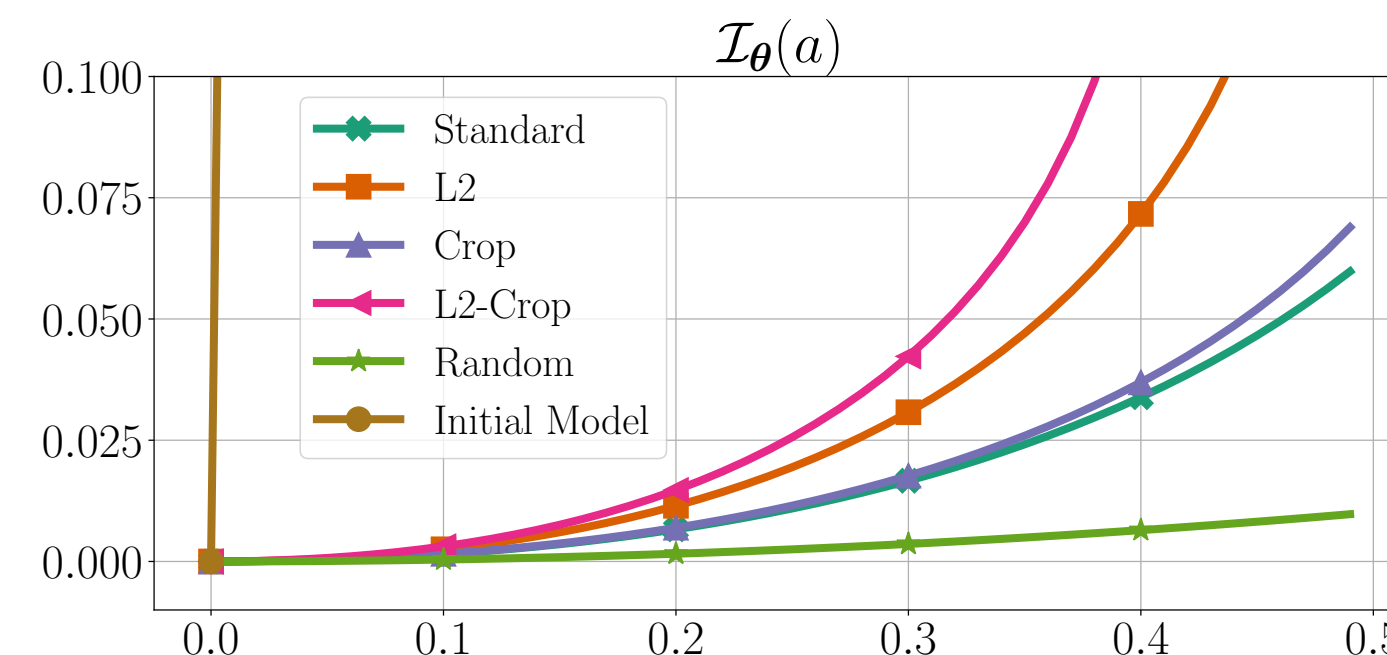


Figure 1: Rate Function of Inception on Cifar10.

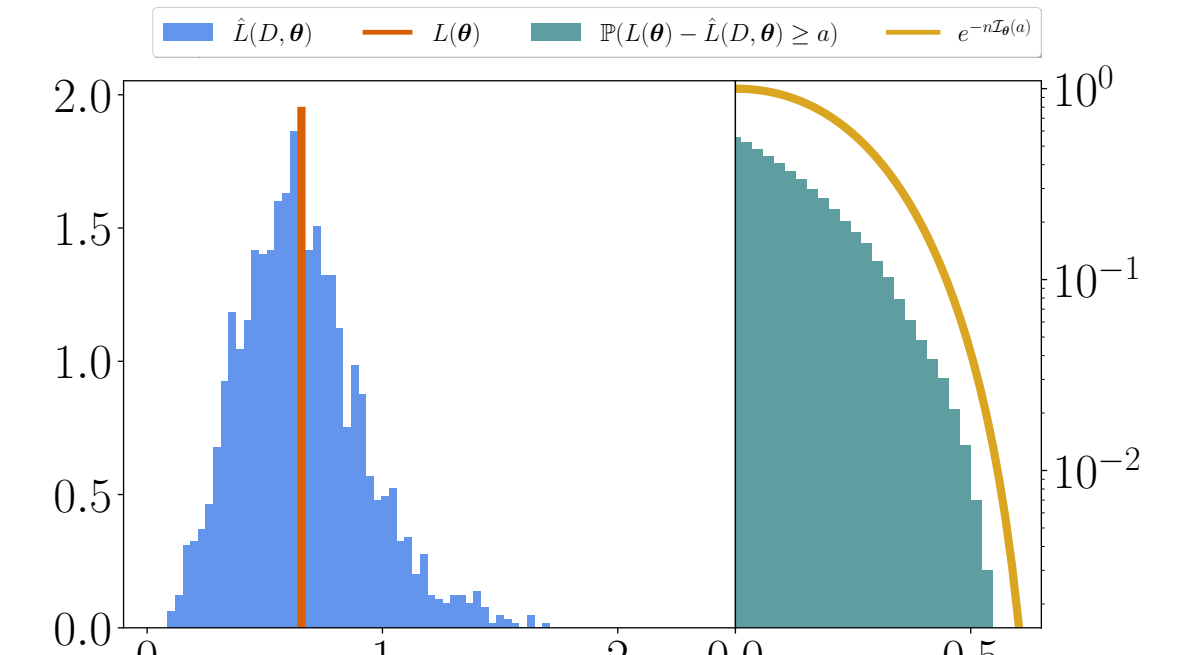


Figure 2: $\hat{L}(D, \theta)$ with $D \sim \nu^{50}$

PAC-Chernoff Bound

Theorem 4.1. With h.p., for all $\theta \in \Theta$, simultaneously,

$$L(\theta) \leq \hat{L}(D, \theta) + \mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right).$$

with

$$\mathcal{I}_{\theta}^{-1}(s) = \inf_{\lambda > 0} \frac{J_{\theta}(\lambda) + s}{\lambda} \quad \forall s \geq 0.$$

Where p is the number of **parameters** of the model class.

Tightness

Proposition. With h.p., for all $\theta \in \Theta$, simultaneously,

$$L(\theta) \leq \hat{L}(D, \theta) + \mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right) \leq L(\theta) + \hat{L}(D, \theta).$$

The bound is **perfectly tight** for interpolators.

Smoother Interpolators Generalize Better

Smoothness: A model $\theta \in \Theta$ is β -smoother than a model $\theta' \in \Theta'$ if

$$\forall a \in (0, \beta] \quad \mathcal{I}(a) \geq \mathcal{I}'(a).$$

Theorem 4.5. For any $\epsilon \geq 0$, with h.p., for all θ, θ' , simultaneously,

if $\hat{L}(D, \theta) \leq \epsilon$ and θ is β -smoother than θ' , then, $L(\theta) \leq L(\theta') + \epsilon$.

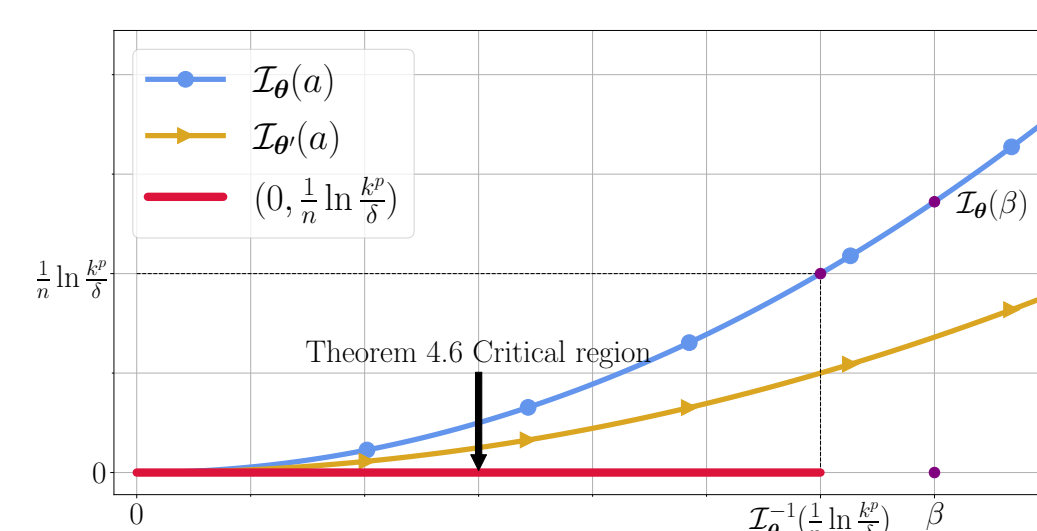


Figure 3: $D_0 \sim \nu_0^{50}$ estimated with CIFAR-10's test set; $D_1 \sim \nu_1^{50}$ adds random translations of 3 pixels; $D_2 \sim \nu_2^{50}$ also adds rotations up to 20° .

$$\exists \beta > 0 \text{ s.t. } \theta \text{ is } \beta\text{-smoother than } \theta' \\ \Updownarrow \\ \nabla_{\nu}(\ell(\mathbf{y}, \mathbf{x}, \theta)) \leq \nabla_{\nu}(\ell(\mathbf{y}, \mathbf{x}, \theta'))$$

Optimal Regularization

The inverse rate is an **optimal regularizer**.

Theorem 5.1 For any $\epsilon > 0$, with h.p. $1 - \delta$ over $D \sim \nu^n$

$$|L(\theta_{\epsilon}^*) - L(\theta_{\epsilon}^{\times})| \leq \epsilon.$$

Where the two models are **interpolators** defined as:

$$\theta_{\epsilon}^{\times} = \arg \min_{\theta: \hat{L}(D, \theta) \leq \epsilon} \hat{L}(D, \theta) + \underbrace{\mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right)}_{\text{Regularizer}},$$

and

$$\theta_{\epsilon}^* = \arg \min_{\theta: \hat{L}(D, \theta) \leq \epsilon} L(\theta).$$

Understanding Existing Regularizers

Many common regularization techniques are **approximations** to the **optimal regularizer**:

Distance from initialization and ℓ_2 -norm:

$$\mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right) \leq \sqrt{2Ma} \|\theta - \theta_0\|_2,$$

Exponential Family and Large Data Sets:

$$\left| \mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right) - \sqrt{2 \frac{1}{n} \ln \frac{k^p}{\delta}} \sqrt{\theta^T \text{Cov}_{\nu}(s(\mathbf{y}, \mathbf{x})) \theta} \right| \leq \epsilon,$$

Input-gradient norm:

$$\mathcal{I}_{\theta}^{-1} \left(\frac{1}{n} \ln \frac{k^p}{\delta} \right) \leq \sqrt{\frac{H}{n} \ln \frac{k^p}{\delta}} \sqrt{\mathbb{E}_{\nu} \left[\|\nabla_{\mathbf{x}} \ell(\mathbf{y}, \mathbf{x}, \theta)\|_2^2 \right]}.$$

Transformed Input Data

- Input data** in many machine learning problems undergo **transformations**, often due to the measuring process, such as sensor noise or image distortions.
- Transformed input-data** makes the expected loss $L(\theta)$ **higher** and the distribution of $\hat{L}(D, \theta)$ with $D \sim \nu^n$ **less concentrated**.

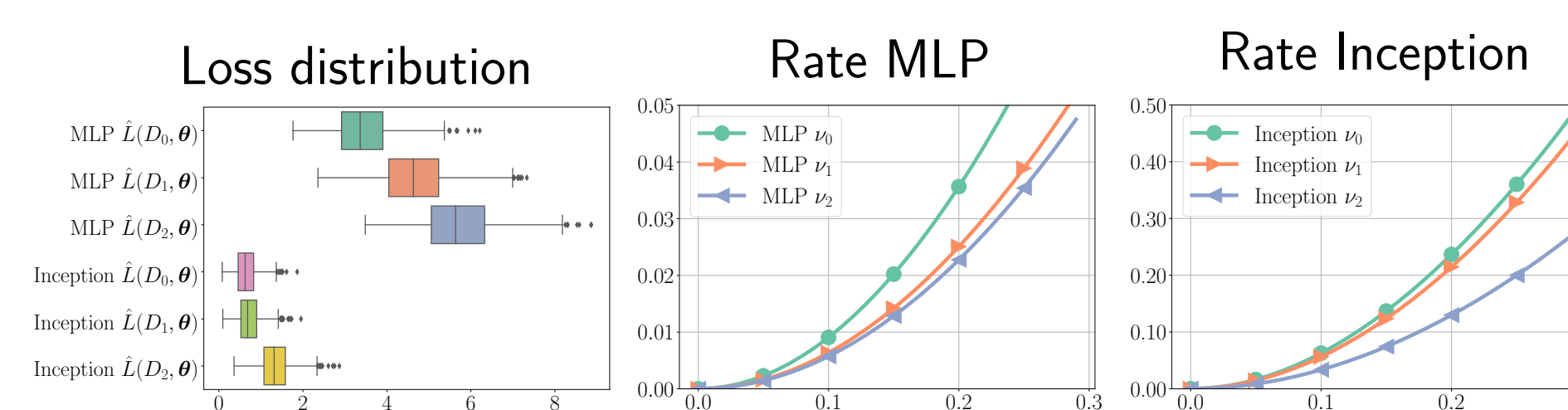


Figure 3: $D_0 \sim \nu_0^{50}$ estimated with CIFAR-10's test set; $D_1 \sim \nu_1^{50}$ adds random translations of 3 pixels; $D_2 \sim \nu_2^{50}$ also adds rotations up to 20° .

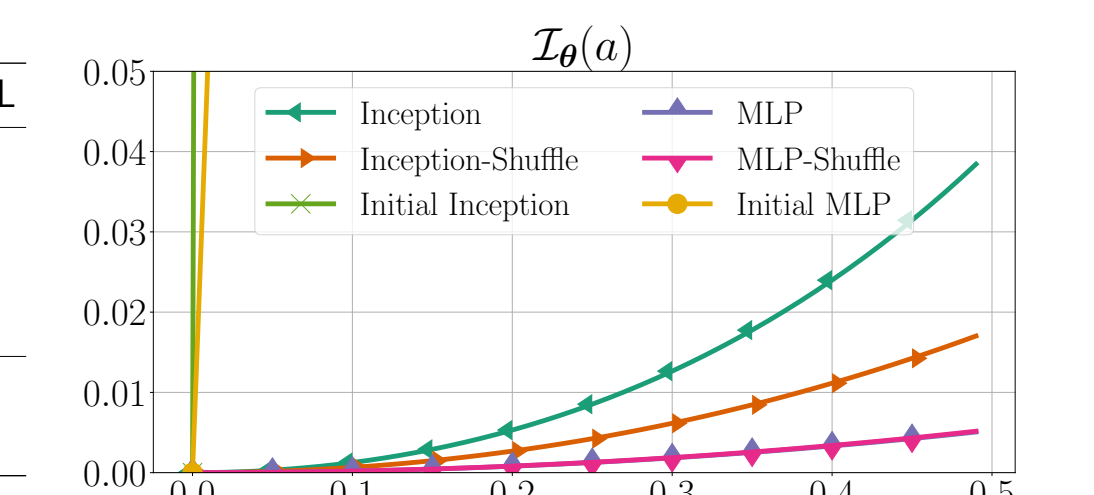
Invariant Architectures

Proposition 6.4 If a model $\theta \in \Theta$ is invariant to transformed-inputs ν_{t+1} ,

$$L^{\nu_{t+1}}(\theta) = L^{\nu_t}(\theta) \quad \text{and} \quad \mathcal{I}_{\theta}^{\nu_{t+1}}(a) = \mathcal{I}_{\theta}^{\nu_t}(a) \quad \forall a > 0.$$

- The $\hat{L}(D, \theta)$ of invariant architectures is **more concentrated** under transformed inputs.
- PAC-Chernoff bounds explain why interpolating with **invariant architecture** leads to better generalization performance.

Model	Train Acc.	Test Acc.	Test NLL
Inception	100.0%	74.08%	1.00
Inception-Shuffle	100.0%	42.46%	2.45
MLP	99.99%	51.69%	3.29
MLP-Shuffle	99.99%	51.12%	3.29
Initial Inception	10.00%	10.00%	2.30
Initial MLP	10.00%	9.96%	2.30



Over-parameterization

The distribution-dependent PAC-Chernoff Bound can be used to obtain **bounds over the number of parameters** of **interpolators**:

Theorem 7.1. For any $\epsilon \in (0, L^*)$ and any $\delta \in (0, 1)$, with high probability $1 - \delta$ over $D \sim \nu^n$, for all $\theta \in \Theta$, simultaneously,

$$\text{if } \hat{L}(D, \theta) \leq \epsilon \text{ then } p \geq \frac{n\mathcal{I}_{\theta}(L^* - \epsilon) + \ln \delta}{\ln k}.$$

This result can be extended to create bounds over **Lipschitz constants**:

$$\text{if } \hat{L}(D, \theta) \leq \epsilon \text{ then } \text{Lip}(\theta) \geq \sqrt{\frac{nd}{2c(p \ln k - \ln \delta)}} (L^* - \epsilon),$$

and over **parameter norms**:

$$\text{if } \hat{L}(D, \theta) \leq \epsilon \text{ then } \|\theta - \theta_0\|_2 \geq \sqrt{\frac{n}{8M(p \ln k - \ln \delta)}} (L^* - \epsilon).$$

Conclusions and Limitations

- Traditional bounds relying solely on training data are **unable** to explain **generalization of over-parameterized interpolators**.
- Distribution-dependent **PAC-Chernoff bounds** are a promising tool able to **explain a wide range of learning techniques**.
- Smoother interpolators** generalize better.
- Connected to a **wide range of regularization methods**.
- Explain why **invariant architectures and data-augmentation** works under transformed input-data.
- Over-parameterization is a **necessary condition** for smooth interpolation.
- Limitation:** Assumption of a finite model class. It can be addressed by using PAC-Bayes Chernoff bounds ([Casado et al. 2024](#)).

Preprint



Code

