

# An Introduction to PAC-Bayesian Theory and its Modern Applications

Pierre Alquier



Statistics and Learning Theory Summer School  
Armenia, July 2026

## Cergy (near Paris)



- Kamelia DAUDEL
- Olga KLOPP
- Marie KRATZ
- Roberto RENO
- Pierre JACOB
- Guillaume CHEVILLON
- Jeroen ROMBOUTS
- Mohamed NDAOUD
- Guillaume LECUE
- Mikołaj KASPRZAK
- Maria ALLAYIOTI
- Vincenzo ESPOSITO VINZI

## Singapore



- Jeremy HENG
- Pierre ALQUIER

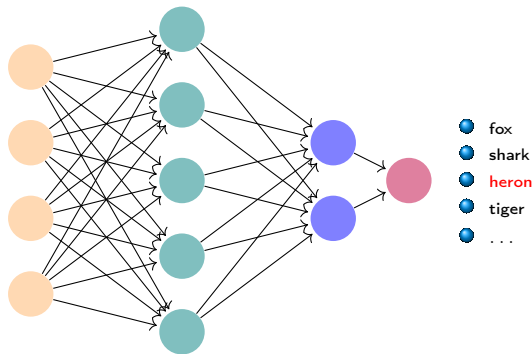
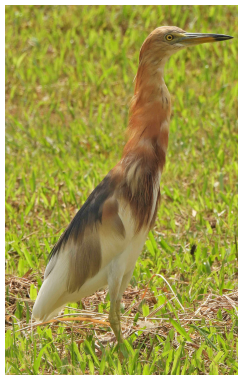
## Rabat



- 1 Introduction : Generalization Bounds and PAC-Bayes
  - Generalization Bounds
  - Minimization of the PAC-Bayes Bound
  - Catoni's PAC-Bayes Bound and its Proof
- 2 Excess Risk Bounds and Mutual Information Bounds
  - From Empirical Bounds to Learning Rates
  - Fast Rates
  - Mutual Information Bounds
- 3 Empirical PAC-Bayes Bounds for Time Series
  - Mathematical Tools
  - PAC-Bayes Bounds for Time Series
  - Towards Empirical Bounds

- 1 Introduction : Generalization Bounds and PAC-Bayes
  - Generalization Bounds
  - Minimization of the PAC-Bayes Bound
  - Catoni's PAC-Bayes Bound and its Proof
- 2 Excess Risk Bounds and Mutual Information Bounds
  - From Empirical Bounds to Learning Rates
  - Fast Rates
  - Mutual Information Bounds
- 3 Empirical PAC-Bayes Bounds for Time Series
  - Mathematical Tools
  - PAC-Bayes Bounds for Time Series
  - Towards Empirical Bounds

- Objects  $x \in \mathcal{X}$ , labels  $y \in \mathcal{Y}$ .
- Predictor : function  $f_\theta : \mathcal{X} \rightarrow \mathcal{Y}$  indexed by  $\theta \in \Theta$ .



- Prediction error measured through loss function  $\ell$  :

$$\ell(y, f_{\theta}(x)).$$

- Risk :

$$R(\theta) := \mathbb{E}_{(X,Y) \sim P} \left[ \ell(Y, f_{\theta}(X)) \right]$$

where  $P$  is the probability distribution of pairs  
object-label we want to learn to classify.

- Objective :

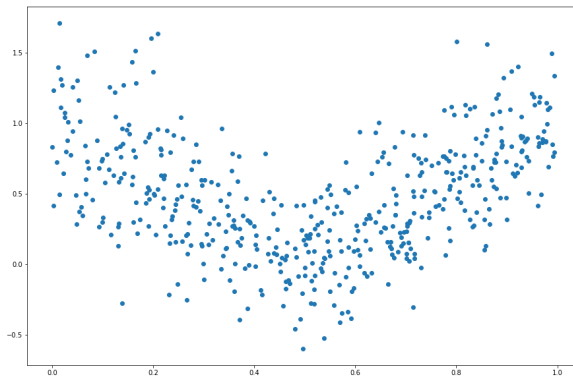
$$R^* = \inf_{\theta \in \Theta} R(\theta).$$

- Data  $(X_1, Y_1), \dots, (X_n, Y_n)$  i.i.d. from  $P$ . Empirical risk :

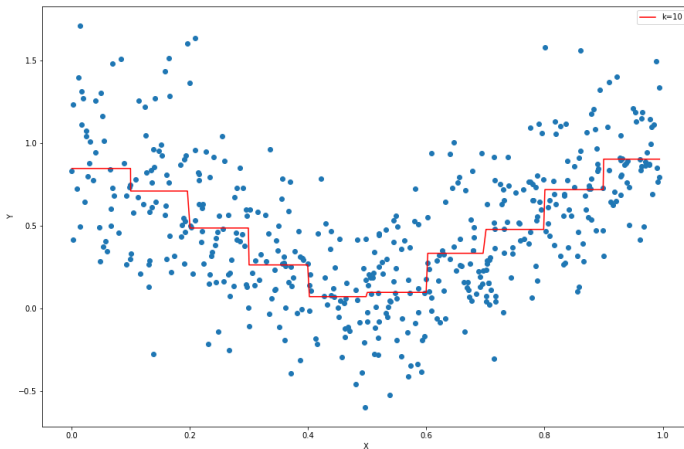
$$R_n(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i, f_{\theta}(X_i)).$$

Toy example :

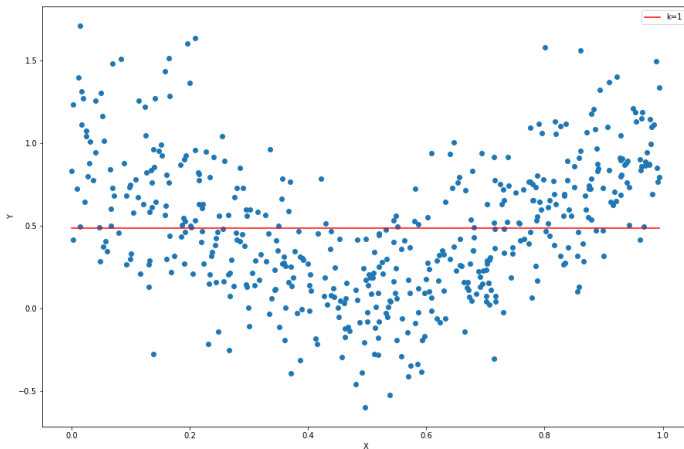
- $X$  uniform on  $[0, 1]$ ,
- $Y = |2X - 1| + \epsilon$ .



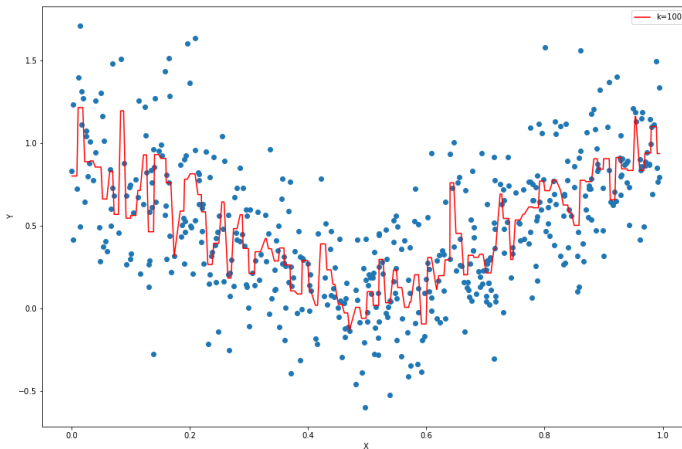
- Prediction by regular histogram with  $k$ -bins.
- $\ell(y, f_\theta(x)) = (y - f_\theta(x))^2$ .

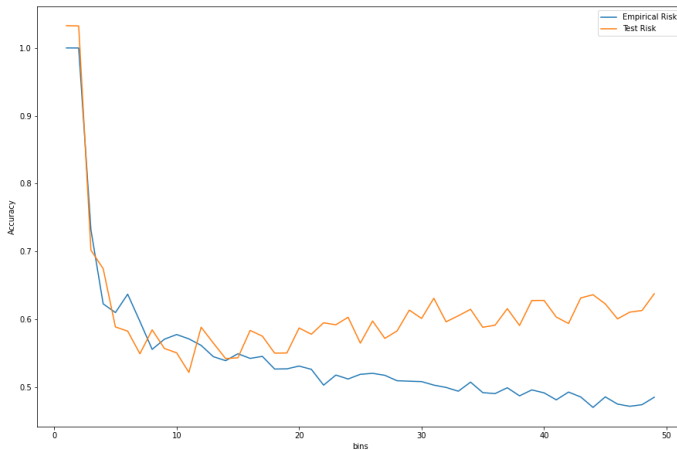


- Prediction by regular histogram with  $k$ -bins.
- $\ell(y, f_\theta(x)) = (y - f_\theta(x))^2$ .



- Prediction by regular histogram with  $k$ -bins.
- $\ell(y, f_\theta(x)) = (y - f_\theta(x))^2$ .





Law of large numbers : for a fixed  $\theta$ ,

$$R_n(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i, f_\theta(X_i)) \xrightarrow[n \rightarrow \infty]{} R(\theta).$$

But  $\hat{\theta} = \hat{\theta}((X_1, Y_1), \dots, (X_n, Y_n)) = \hat{\theta}(\mathcal{S})$  learnt from data.

Can we quantify  $R(\hat{\theta}) - R_n(\hat{\theta})$  when  $\hat{\theta}$  is learnt ?

Various approaches :

- Vapnik-Chervonenkis theory,
- algorithmic stability,
- information bounds : MDL, PAC-Bayes, etc.

## Assumption for **whole WEEK**

Unless specified otherwise,  $0 \leq \ell \leq 1$  and data is i.i.d. from  $P$ .

## Vapnik-Chervonenkis – classification ( $\mathcal{Y} = \{0, 1\}$ )

With probability at least  $1 - \delta$  on the data, for any  $\hat{\theta}$  learnt from the data,

$$R(\hat{\theta}) \leq R_n(\hat{\theta}) + \sqrt{\frac{8d \log\left(\frac{2en}{d}\right) + 8 \log\left(\frac{4}{\delta}\right)}{n}}$$

where  $d$  : the VC-dimension of the set of classifiers  $(f_{\theta}, \theta \in \Theta)$ .

Statistical estimation / ERM etc.

data  $\longrightarrow$  estimator

$$(\mathcal{X} \times \mathcal{Y})^n \longrightarrow \Theta$$

$$\mathcal{S} \longmapsto \hat{\theta} = \hat{\theta}(\mathcal{S})$$

Randomized estimators :

$$(\mathcal{X} \times \mathcal{Y})^n \longrightarrow \mathcal{M}(\Theta) \dashrightarrow \Theta$$

$$\mathcal{S} \longmapsto \hat{\rho} = \hat{\rho}(\mathcal{S}) \dashrightarrow^{\theta \sim \hat{\rho}} \theta$$

## McAllester's PAC-Bayes bound

Fix a prior distribution  $\pi \in \mathcal{M}(\Theta)$ . With probability at least  $1 - \delta$  on the data  $\mathcal{S}$ , for any probability distribution  $\rho$  learnt on the data,

$$\mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \sqrt{\frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{2n}}.$$

$\text{KL}(\rho \parallel \pi)$  = Küllback-Leibler divergence between  $\rho$  and  $\pi$

$KL(\rho||\pi)$  = Küllback-Leibler divergence between  $\rho$  and  $\pi$

- discrete case :

$$KL(\rho||\pi) = \sum_{\theta \in \Theta} \rho(\theta) \log \frac{\rho(\theta)}{\pi(\theta)}$$

and  $KL(\rho||\pi) = \infty$  if for some  $\theta$ ,  $\pi(\theta) = 0$  and  $\rho(\theta) > 0$ .

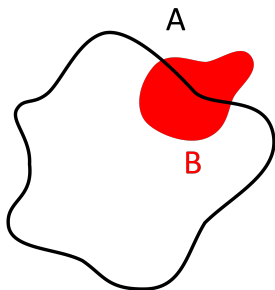
- general case

$$KL(\rho||\pi) = \mathbb{E}_{\theta \sim \rho} \left[ \log \frac{d\rho}{d\pi}(\theta) \right]$$

and  $KL(\rho||\pi) = \infty$  if  $\rho$  has no density  $\frac{d\rho}{d\pi}$  w.r.t.  $\pi$ ...

$$KL(\rho||\pi) \geq 0 \text{ and } KL(\rho||\pi) = 0 \Leftrightarrow \rho = \pi.$$

Intuition on KL :



- $\pi$  uniform on  $A$

$$\pi(\theta) = \frac{1_A(\theta)}{\mathcal{V}(A)}$$

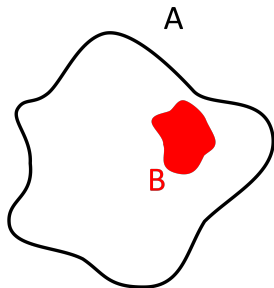
- $\rho$  uniform on  $B$

$$\rho(\theta) = \frac{1_B(\theta)}{\mathcal{V}(B)}$$

$\frac{d\rho}{d\pi}$  not defined here.

$$B \not\subseteq A \Rightarrow \text{KL}(\rho \parallel \pi) = +\infty.$$

Intuition on KL :



- $\pi$  uniform on  $A$

$$\pi(\theta) = \frac{1_A(\theta)}{\mathcal{V}(A)}$$

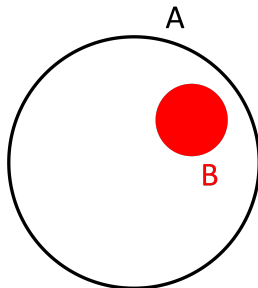
- $\rho$  uniform on  $B$

$$\rho(\theta) = \frac{1_B(\theta)}{\mathcal{V}(B)}$$

$$B \subseteq A \Rightarrow \frac{d\rho}{d\pi}(\theta) = \frac{\mathcal{V}(A)1_B(\theta)}{\mathcal{V}(B)}$$

$$\text{KL}(\rho \parallel \pi) = \mathbb{E}_{\theta \sim \rho} \left[ \log \frac{d\rho}{d\pi}(\theta) \right] = \log \frac{\mathcal{V}(A)}{\mathcal{V}(B)}.$$

Intuition on KL :



$B_d(x, r)$  ball centered on  $x$ , with radius  $r$  in  $\mathbb{R}^d$

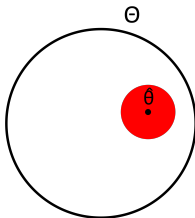
$$\mathcal{V}(B_d(x, r)) = \frac{r^d \pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2} + 1)}$$

- $\pi$  uniform on  $A = B_d(0, C)$
- $\rho$  uniform on  $B = B_d(\theta_0, \epsilon)$

$$\text{KL}(\rho \parallel \pi) = \log \frac{\mathcal{V}(A)}{\mathcal{V}(B)} = d \log \frac{C}{\epsilon}.$$

## McAllester's PAC-Bayes bound

$$\mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \sqrt{\frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{2n}}.$$



- $\pi$  uniform on  $\Theta = B_d(0, C)$
- $\rho = \hat{\rho}$  uniform on  $B_d(\hat{\theta}, \epsilon)$

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \sqrt{\frac{d \log \frac{C}{\epsilon} + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{2n}}.$$

## Minimization of the PAC-Bayes bound.

## McAllester's PAC-Bayes bound

Fix prior  $\pi \in \mathcal{M}(\Theta)$ . With proba. at least  $1 - \delta$ ,  $\forall \rho \in \mathcal{M}(\Theta)$ ,

$$\mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \sqrt{\frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{2n}}.$$

$$\sqrt{\frac{a}{b}} = \inf_{\lambda > 0} \left\{ \frac{a}{\lambda} + \frac{\lambda}{4b} \right\}.$$

$$\begin{aligned} \mathbb{E}_{\theta \sim \rho}[R(\theta)] &\leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \\ &\quad + \inf_{\lambda > 0} \left\{ \frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{\lambda} + \frac{\lambda}{8n} \right\}. \end{aligned}$$

$$\forall \lambda > 0, \mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \\ + \frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{\lambda} + \frac{\lambda}{8n}.$$

### Definition - Gibbs posterior

$$\hat{\pi}_\lambda(d\theta) = \frac{\exp(-\lambda R_n(\theta))}{\mathbb{E}_{\vartheta \sim \pi}[\exp(-\lambda R_n(\vartheta))]} \pi(d\theta).$$

### Proposition

$$\hat{\pi}_\lambda = \arg \min_{\rho \in \mathcal{M}(\Theta)} \left\{ \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \frac{\text{KL}(\rho \parallel \pi)}{\lambda} \right\}.$$

Sampling from  $\hat{\pi}_\lambda$  by Monte-Carlo techniques...

## Lemma (Donsker-Varadhan's variational formula)

Let  $h(\theta)$  be any bounded, measurable function of  $\theta$ . Then

$$F(\rho) := \mathbb{E}_{\theta \sim \rho}[h(\theta)] - \text{KL}(\rho \parallel \pi)$$

reaches its unique maximum at  $\rho^*$  defined by

$$\frac{d\rho^*}{d\pi}(\theta) = \frac{\exp(h(\theta))}{\mathbb{E}_{\vartheta \sim \rho}[h(\vartheta)]}$$

and

$$F(\rho^*) = \log \mathbb{E}_{\theta \sim \pi} \exp[h(\theta)].$$

To prove the previous proposition :  $h(\theta) \leftarrow -\lambda R_n(\theta)$ .

## Proof of the Lemma.

$$\begin{aligned}\text{KL}(\rho \parallel \rho^*) &= \mathbb{E}_{\theta \sim \rho} \left[ \log \frac{d\rho}{d\rho^*}(\theta) \right] \\ &= \mathbb{E}_{\theta \sim \rho} \left[ \log \frac{\mathbb{E}_{\vartheta \sim \pi}[h(\vartheta)]}{\exp(h(\theta))} \frac{d\rho}{d\pi}(\theta) \right] \\ &= \mathbb{E}_{\theta \sim \rho} \left[ \log \frac{d\rho}{d\pi}(\theta) \right] - \mathbb{E}_{\theta \sim \rho}[h(\theta)] + \log \mathbb{E}_{\vartheta \sim \pi}[h(\vartheta)] \\ &\geq 0\end{aligned}$$

with equality if and only if  $\rho = \rho^*$ .

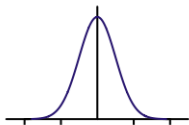
Approximate minimization of the PAC-Bayes bound.

$$\forall \lambda > 0, \mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \frac{\text{KL}(\rho \parallel \pi) + \log\left(\frac{2\sqrt{n}}{\delta}\right)}{\lambda} + \frac{\lambda}{8n}.$$

Alternative approach : optimize  $\rho$  in a smaller set  $\mathcal{F} \subsetneq \mathcal{M}(\Theta)$ .

Definition - variational approximation of Gibbs posterior

$$\tilde{\rho}_\lambda = \arg \min_{\rho \in \mathcal{F}} \left\{ \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] + \frac{\text{KL}(\rho \parallel \pi)}{\lambda} \right\}.$$



Example :  $\rho = \mathcal{N}(\mu, \Sigma)$ , optimize  $(\mu, \Sigma)$ .

Application : generalization bounds for deep learning.

Train a neural network for classification (0-1 loss).

Vapnik-type bound usually lead to something larger than 1, for example :

$$R(\hat{\theta}) \leq 35.4$$

As  $R(\hat{\theta}) = \mathbb{P}(Y \neq f_{\hat{\theta}}(X))$ , the bound brings no information (vacuous).

---

## Computing Nonvacuous Generalization Bounds for Deep (Stochastic) Neural Networks with Many More Parameters than Training Data

---

Gintare Karolina Dziugaite  
Department of Engineering  
University of Cambridge

Daniel M. Roy  
Department of Statistical Sciences  
University of Toronto

### Abstract

One of the defining properties of deep learning is that models are chosen to have many more parameters than available training data. In light of this capacity for overfitting, it is remarkable that simple algorithms like SGD reliably return solutions with low test error. One roadblock to explaining these phenomena in terms of implicit regularization, structural properties of the solution, and/or easiness of the data is that many learning bounds are quantitatively vacuous when applied to networks learned by SGD in this “deep learning” regime. Logically, in order to explain generalization, we need nonvacuous bounds. We return to an idea by Langford and Caruana (2001), who used PAC-Bayes bounds to compute nonvacuous numerical bounds on generalization error for *stochastic* two-layer two-hidden-unit neural networks via a sensitivity analysis. By optimizing the PAC-Bayes bound directly, we are able to extend their approach and obtain nonvacuous generalization bounds for deep stochastic neural network classifiers with millions of parameters trained on only tens of thousands of examples. We connect our findings to recent and old work on flat minima and MDL-based explanations of generalization.

### 1 INTRODUCTION

By optimizing a PAC-Bayes bound, we show that it is possible to compute nonvacuous numerical bounds on the generalization error of deep *stochastic* neural networks with millions of parameters, despite the training data sets being one or more orders of magnitude smaller than the number of parameters. To our knowledge, these are the first explicit and nonvacuous numerical bounds computed

for trained neural networks in the modern deep learning regime where the number of network parameters eclipses the number of training examples.

The bounds we compute are data dependent, incorporating millions of components optimized numerically to identify a large region in weight space with low average empirical error around the solution obtained by stochastic gradient descent (SGD). The data dependence is essential: indeed, the VC dimension of neural networks is typically bounded below by the number of parameters, and so one needs as many training data as parameters before (uniform) PAC bounds are nonvacuous, i.e., before the generalization error falls below 1. To put this in concrete terms, on MNIST, having even 72 hidden units in a fully connected first layer yields vacuous PAC bounds.

Evidently, we are operating far from the worst case: observed generalization cannot be explained in terms the regularizing effect of the size of the neural network alone. This is an old observation, and one that attracted considerable theoretical attention two decades ago: Bartlett [Bar97; Bar98] showed that, in large (sigmoidal) neural networks, when the learned weights are small in magnitude, the fat-shattering dimension is more important than the VC dimension for characterizing generalization. In particular, Bartlett established classification error bounds in terms of the empirical margin and the fat-shattering dimension, and then gave fat-shattering bounds for neural networks in terms of the *magnitudes* of the weights and the depth of the network alone. Improved norm-based bounds were obtained using Rademacher and Gaussian complexity by Bartlett and Mendelson [BM02] and Koltchinskii and Panchenko [KP02].

These norm-based bounds are the foundation of our current understanding of neural network generalization. It is widely accepted that these bounds explain observed generalization, at least “qualitatively” and/or when the weights are explicitly regularized. Indeed, recent work by Neyshabur, Tomioka, and Srebro [NTS14] puts forth

| Experiment      | T-600 | T-1200 | T-300 <sup>2</sup> | T-600 <sup>2</sup> | T-1200 <sup>2</sup> | T-600 <sup>3</sup> | R-600  |
|-----------------|-------|--------|--------------------|--------------------|---------------------|--------------------|--------|
| Train error     | 0.001 | 0.002  | 0.000              | 0.000              | 0.000               | 0.000              | 0.007  |
| Test error      | 0.018 | 0.018  | 0.015              | 0.016              | 0.015               | 0.013              | 0.508  |
| SNN train error | 0.028 | 0.027  | 0.027              | 0.028              | 0.029               | 0.027              | 0.112  |
| SNN test error  | 0.034 | 0.035  | 0.034              | 0.033              | 0.035               | 0.032              | 0.503  |
| PAC-Bayes bound | 0.161 | 0.179  | 0.170              | 0.186              | 0.223               | 0.201              | 1.352  |
| KL divergence   | 5144  | 5977   | 5791               | 6534               | 8558                | 7861               | 201131 |
| # parameters    | 471k  | 943k   | 326k               | 832k               | 2384k               | 1193k              | 472k   |
| VC dimension    | 26m   | 56m    | 26m                | 66m                | 187m                | 121m               | 26m    |

Table 1: Results for experiments on binary class variant of MNIST. SGD is either trained on (T) true labels or (R) random labels. The network architecture is expressed as  $N^L$ , indicating  $L$  hidden layers with  $N$  nodes each. Errors are classification error. The reported VC dimension is the best known upper bound (in millions) for ReLU networks. The SNN error rates are tight upper bounds (see text for details). The PAC-Bayes bounds upper bound the test error with probability 0.965.

## Results taken from :



Dzuigaite, G. K. and Roy, D. M. (2017). Computing Nonvacuous Generalization Bounds for Deep (Stochastic) Neural Networks with Many More Parameters than Training Data. *UAI*.

## Catoni's PAC-Bayes Bound (with Proof!)

## PAC-Bayesian Model Averaging

David A. McAllester  
AT&T Shannon Labs  
180 Park Avenue  
Florham Park, NJ 07932-0971  
dmac@research.att.com

### Abstract

PAC-Bayesian learning methods combine the informative priors of Bayesian methods with distribution-free PAC guarantees. Building on earlier methods for PAC-Bayesian model selection, this paper presents a method for PAC-Bayesian model averaging. The method constructs an optimized weighted mixture of concepts analogous to a Bayesian posterior distribution. Although the main result is stated for bounded loss, a preliminary analysis for unbounded loss is also given.

### 1 INTRODUCTION

A PAC-Bayesian approach to machine learning attempts to combine the advantages of both PAC and Bayesian approaches [12, 8]. The Bayesian approach has the advantage of using arbitrary domain knowledge in the form of a Bayesian prior. The PAC approach has the advantage that one can prove guarantees for generalization error without assuming the truth of the prior. A PAC-Bayesian approach combines the features of the PAC and Bayesian approaches — it bases the bias of the learning algorithm on an arbitrary prior distribution, thus allowing the incorporation of domain knowledge, and yet provides a guarantee on generalization error that is independent of any truth of the prior.

PAC-Bayesian approaches are related to structural risk minimization (SRM) [6]. Here we interpret this broadly as describing any learning algorithm optimizing a tradeoff between the “complexity”, “structure”, or “prior probability” of the concept or model and the “goodness of fit”, “description length”, or “likelihood” of the training data. Under this interpretation of SRM, Bayesian algorithms which select a concept of maximum posterior probability (MAP algorithms) are viewed as a kind of SRM algorithm. Various approaches to SRM

are compared both theoretically and experimentally by Krume et al. in [6]. They give experimental evidence that Bayesian and MDL algorithms tend to over-fit in experimental settings where the Bayesian assumptions fail. A PAC-Bayesian approach uses a prior distribution analogous to that used in MAP or MDL but provides a theoretical guarantee against over-fitting independent of the truth of the prior.

Earlier work on PAC-Bayesian algorithms has focused on model selection — selecting either a single concept or a uniformly weighted set of concepts. Here we consider nonuniform model averaging, i.e., selecting a weighted mixture of the concepts.

Model averaging is empirically important in certain applications. For example, in statistical language modeling for speech recognition one “smooths” a trigram model with a bigram model and smooths the bigram model with a unigram model. This smoothing is essential for minimizing the cross entropy between, say, the model and a test corpus of newspaper sentences. It turns out that smoothing in statistical language modeling is more naturally formulated as model averaging than as model selection. A smoothed language model is very large — it contains a full trigram model, a full bigram model and a full unigram model as parts. If one uses MDL to select the structure of a language model, selecting model parameters with maximum likelihood, the resulting structure is much smaller than that of a smoothed trigram model. Furthermore, the MDL model performs quite badly. However, a smoothed trigram model can be theoretically derived as a compact representation of a Bayesian mixture of an exponential number of (smaller) suffix tree models [10].

Model averaging can also be applied to decision trees. A common method of constructing decision trees is to first build an overly large tree which over-fits the training data and then prune the tree in some way so as to get a smaller tree that does not over-fit the data [11, 5]. An alternative to pruning is to construct a weighted mixture of the subtrees of the original over-fit tree. It is possible to construct a concise representation of a weighting over exponentially many different subtrees [3, 9, 4].

This paper proves a new PAC-Bayesian theorem giving a bound on the generalization error of weighted mixtures. A weighted mixture which gives too much weight to models with low prior probability will over-fit the

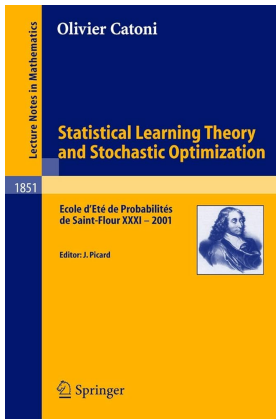
Seminal paper, that contains the bound stated earlier today.

Since then, various bounds published :

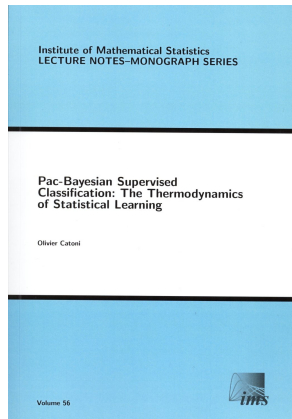
- tighter,
- with less assumptions (i.i.d, bounded loss),
- easier to optimize,
- ...

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are made within the limits of not for profit or commercial sale, and that copies bear this notice and the full citation on the first page. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. COLT '98, Santa Cruz, CA, USA. © 1998 ACM 1-58113-167-6/98/0008...\$5.00

## Classical references :



Connections with information theory  
and MDL.



Very tight bounds, applications to  
Support Vector Machines.

## Catoni's PAC-Bayes bound, 2003

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda}{8n}.$$



Catoni, O. (2003). *A PAC-Bayesian approach to adaptive classification*. Preprint LPMA 840.

- Cons : not optimized in  $\lambda$ .
- Pros : easily generalized to other settings.

## Proof of Catoni's bound : step 1.

### Hoeffding's lemma

Let  $U$  be a random variable with  $a \leq U \leq b$  almost surely and  $\mathbb{E}(U) = 0$ . Then, for any  $s \in \mathbb{R}$ ,

$$\mathbb{E}[\exp(sU)] \leq \exp\left(\frac{s^2(b-a)^2}{8}\right).$$

**Proof of Hoeffding's lemma.** Put  $f(s) = \log \mathbb{E}[\exp(sU)]$ .

By differentiation under the integral sign theorems :

$$\begin{aligned} f'(s) &= \frac{\mathbb{E}[U \exp(sU)]}{\mathbb{E}[\exp(sU)]} = \int_{[a,b]} u \frac{\exp(su)}{\mathbb{E}[\exp(sU)]} P_U(du) = \int_{[a,b]} u P_V(du) \\ &= \mathbb{E}(V) \text{ for some } a \leq V \leq b. \end{aligned}$$

Check :

$$f''(s) = \text{Var}(V) \leq \frac{(b-a)^2}{4}.$$

Finally :

$$\begin{aligned} f(s) &= f(0) + sf'(0) + \int_0^s f''(t)(s-t)dt \\ &\leq 0 + 0 + \int_0^s \frac{(b-a)^2}{4}(s-t)dt = \frac{s^2(b-a)^2}{8}. \end{aligned}$$

## Proof of Catoni's bound : step 2.

Put  $\ell_i(\theta) := \ell(Y_i, f_\theta(X_i))$  for short.

$$\begin{aligned}\mathbb{E}_{\mathcal{S}} \exp \left\{ \lambda \left[ R(\theta) - R_n(\theta) \right] \right\} &= \mathbb{E}_{\mathcal{S}} \exp \left\{ \frac{\lambda}{n} \left[ \sum_{i=1}^n (\mathbb{E}_{\mathcal{S}}[\ell_i(\theta)] - \ell_i(\theta)) \right] \right\} \\ &= \prod_{i=1}^n \mathbb{E}_{\mathcal{S}} \exp \left\{ \frac{\lambda}{n} [\mathbb{E}_{\mathcal{S}}[\ell_i(\theta)] - \ell_i(\theta)] \right\} \\ &\leq \prod_{i=1}^n \exp \left( \frac{\lambda^2}{8n^2} \right) \text{ thanks to Hoeffding.}\end{aligned}$$

Lemma : exponential moment of the risk

$$\forall \lambda \in \mathbb{R}, \mathbb{E}_{\mathcal{S}} \exp \left\{ \lambda \left[ R(\theta) - R_n(\theta) \right] - \frac{\lambda^2}{8n} \right\} \leq 1.$$

## Proof of Catoni's bound : step 3.

Integrate the previous lemma with respect to  $\pi$ , and use Fubini's theorem to get :

$$\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\theta \sim \pi} \exp \left\{ \lambda \left[ R(\theta) - R_n(\theta) \right] - \frac{\lambda^2}{8n} \right\} \leq 1.$$

Remind Donsker-Varadhan's formula :

$$\log \mathbb{E}_{\theta \sim \pi} \exp[h(\theta)] = \sup_{\rho} [\mathbb{E}_{\theta \sim \rho} [h(\theta)] - \text{KL}(\rho \| \pi)] .$$

### Catoni's PAC-Bayes key lemma

$$\forall \lambda \in \mathbb{R},$$

$$\mathbb{E}_{\mathcal{S}} \exp \sup_{\rho} \left\{ \lambda \mathbb{E}_{\theta \sim \rho} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\rho \| \pi) \right\} \leq 1.$$

## Proof of Catoni's bound : final step.

$$\mathbb{E}_{\mathcal{S}} \exp \sup_{\rho} \left\{ \lambda \mathbb{E}_{\theta \sim \rho} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\rho \parallel \pi) \right\} \leq 1.$$

Put  $\delta \in (0, 1)$ , use Markov's inequality :

$$\mathbb{P} \left( U \geq \log \frac{1}{\delta} \right) = \mathbb{P} \left( \exp(U) \geq \frac{1}{\delta} \right) \leq \delta \mathbb{E}[\exp(U)] \leq \delta.$$

$$\mathbb{P} \left( \sup_{\rho} \left\{ \lambda \mathbb{E}_{\theta \sim \rho} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\rho \parallel \pi) \right\} \geq \log \frac{1}{\delta} \right) \leq \delta.$$

... what we wanted to prove

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda}{8n}.$$

- 1 Introduction : Generalization Bounds and PAC-Bayes
  - Generalization Bounds
  - Minimization of the PAC-Bayes Bound
  - Catoni's PAC-Bayes Bound and its Proof
- 2 Excess Risk Bounds and Mutual Information Bounds
  - From Empirical Bounds to Learning Rates
  - Fast Rates
  - Mutual Information Bounds
- 3 Empirical PAC-Bayes Bounds for Time Series
  - Mathematical Tools
  - PAC-Bayes Bounds for Time Series
  - Towards Empirical Bounds

## Refresher :

- Prediction error measured through loss function  $\ell$  :

$$\ell(y, f_{\theta}(x)).$$

- Risk :

$$R(\theta) := \mathbb{E}_{(X, Y) \sim P} \left[ \ell(Y, f_{\theta}(X)) \right].$$

- Objective :

$$R^* = \inf_{\theta \in \Theta} R(\theta).$$

- Data  $(X_1, Y_1), \dots, (X_n, Y_n)$  **i.i.d. from  $P$** . Empirical risk :

$$R_n(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(Y_i, f_{\theta}(X_i)).$$

Statistical estimation / ERM etc.

data  $\longrightarrow$  estimator

$$(\mathcal{X} \times \mathcal{Y})^n \longrightarrow \Theta$$

$$\mathcal{S} \longmapsto \hat{\theta} = \hat{\theta}(\mathcal{S})$$

Randomized estimators :

$$(\mathcal{X} \times \mathcal{Y})^n \longrightarrow \mathcal{M}(\Theta) \dashrightarrow \Theta$$

$$\mathcal{S} \longmapsto \hat{\rho} = \hat{\rho}(\mathcal{S}) \dashrightarrow^{\theta \sim \hat{\rho}} \theta$$

Yesterday, we proved :

### Catoni's PAC-Bayes bound, 2003

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda}{8n}.$$

In the proof, we stated the following :

Reminder : Catoni's PAC-Bayes key lemma

$$\mathbb{E}_S \exp \sup_{\rho} \left\{ \lambda \mathbb{E}_{\theta \sim \rho} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\rho \| \pi) \right\} \leq 1.$$

$$\mathbb{E}_S \exp \left\{ \lambda \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\hat{\rho} \| \pi) \right\} \leq 1.$$

Jensen's inequality gives :

$$\exp \mathbb{E}_S \left\{ \lambda \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} - \text{KL}(\hat{\rho} \| \pi) \right\} \leq 1.$$

## Catoni's PAC-Bayes bound in expectation, 2003

- Fix  $\lambda > 0$ ,  $\pi$  and a randomized estimator  $\hat{\rho}$ .

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta)] \right] \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi)}{\lambda} + \frac{\lambda}{8n} \right].$$



Catoni, O. (2003). *A PAC-Bayesian approach to adaptive classification*. Preprint LPMA 840.



Zhang, T. (2006). Information-theoretic upper and lower bounds for statistical estimation *IEEE Transactions on Information Theory*.

Note : sometimes referred to as “MAC-Bayes” for “Mean Approximately Correct”...

## Reminder – Gibbs posterior

$$\hat{\pi}_\lambda = \arg \min_{\rho \in \mathcal{M}(\Theta)} \left\{ \mathbb{E}_{\theta \sim \rho} [R_n(\theta)] + \frac{\text{KL}(\rho \| \pi)}{\lambda} \right\}$$
$$\hat{\pi}_\lambda(d\theta) = \frac{\exp(-\lambda R_n(\theta))}{\mathbb{E}_{\vartheta \sim \pi} [\exp(-\lambda R_n(\vartheta))]} \pi(d\theta).$$

Consequence of PAC-Bayes bound in expectation :

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R(\theta)] \right] &\leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R_n(\theta)] + \frac{\text{KL}(\hat{\pi}_\lambda \| \pi)}{\lambda} + \frac{\lambda}{8n} \right] \\ &= \mathbb{E}_{\mathcal{S}} \inf_{\rho} \left[ \mathbb{E}_{\rho} [R_n(\theta)] + \frac{\text{KL}(\rho \| \pi)}{\lambda} + \frac{\lambda}{8n} \right] \\ &\leq \inf_{\rho} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\rho} [R_n(\theta)] + \frac{\text{KL}(\rho \| \pi)}{\lambda} + \frac{\lambda}{8n} \right]. \end{aligned}$$

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq \inf_{\rho} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\rho} [R_n(\theta)] + \frac{\text{KL}(\rho \parallel \pi)}{\lambda} + \frac{\lambda}{8n} \right].$$

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\rho} [R_n(\theta)] &= \mathbb{E}_{\rho} [\mathbb{E}_{\mathcal{S}} R_n(\theta)] \text{ (Fubini-Tonelli)} \\ &= \mathbb{E}_{\rho} [R(\theta)]. \end{aligned}$$

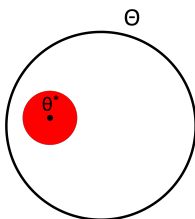
### Catoni's PAC-Bayes excess-risk bound, 2003

- Fix  $\lambda > 0$ ,  $\pi$ , and let  $\hat{\pi}_{\lambda}$  be the Gibbs posterior.

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq \inf_{\rho \in \mathcal{M}(\Theta)} \left[ \mathbb{E}_{\theta \sim \rho} [R(\theta)] + \frac{\text{KL}(\rho \parallel \pi)}{\lambda} + \frac{\lambda}{8n} \right].$$

- In this result, we use  $\hat{\pi}_{\lambda}$  as our randomized estimator.
- But to explicit the right-hand side, we can substitute anything to  $\rho$  in the infimum...

Illustration :



- $\pi$  uniform on  $\Theta = B_d(0, C)$
- $\rho$  uniform on  $B_d(\theta^*, \epsilon)$  where  $R(\theta^*) = R^*$ .

$$\text{KL}(\rho \parallel \pi) = d \log \frac{C}{\epsilon}.$$

Assume  $\theta \mapsto \ell(y, f_\theta(x))$  is  $L$ -Lipschitz around  $\theta^*$ , that is :

$$|\ell(y, f_\theta(x)) - \ell(y, f_{\theta^*}(x))| \leq L \|\theta - \theta^*\|.$$

By taking expectations :

$$R(\theta) - R^* = R(\theta) - R(\theta^*) \leq L \|\theta - \theta^*\|.$$

Thus  $\mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq R^* + L\epsilon$ .

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq \inf_{\epsilon > 0} \left[ R^* + \epsilon + \frac{d \log \frac{C}{\epsilon}}{\lambda} + \frac{\lambda}{8n} \right].$$

The bound is exactly minimized for  $\epsilon = d/\lambda$  :

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq R^* + \frac{d}{\lambda} \left( 1 + \log \frac{C\lambda}{d} \right) + \frac{\lambda}{8n}.$$

In this case, we can calibrate the Gibbs posterior with  $\lambda = \sqrt{n/d}$  which leads to

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq R^* + \mathcal{O} \left( \sqrt{\frac{d}{n}} \log \frac{n}{d} \right).$$

## Reminder : Catoni's PAC-Bayes oracle bound

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] \leq \inf_{\rho \in \mathcal{M}(\Theta)} \left[ \mathbb{E}_{\theta \sim \rho} [R(\theta)] + \frac{\text{KL}(\rho \parallel \pi)}{\lambda} + \frac{\lambda}{8n} \right].$$

Recap on the previous example :

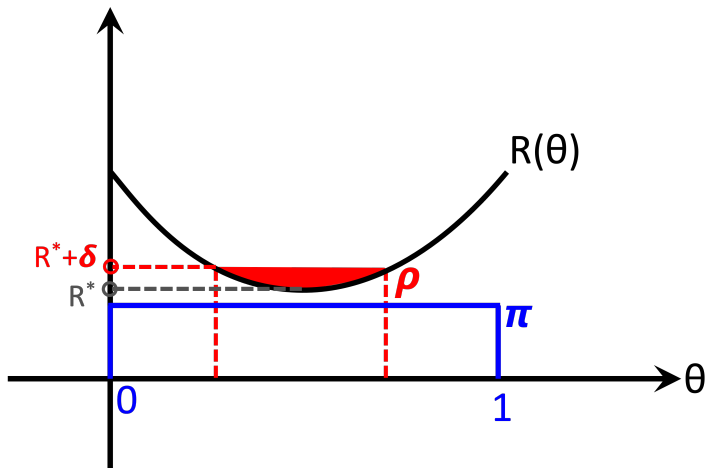
- $\pi$  uniform on  $B_d(0, C)$ ,
- $\rho$  uniform on  $B_d(\theta^*, \epsilon)$ ,
- $L$ -Lipschitz loss function.

$$\Rightarrow \mathbb{E}_{\theta \sim \rho} [R(\theta)] \leq R^* + L\epsilon.$$

More generally, we can consider in the PAC-Bayes oracle bound :

$$\rho = \pi_{\delta} := \text{restriction of } \pi \text{ to } \{\theta : R(\theta) \leq R^* + \delta\}.$$

$\rho = \pi_\delta :=$  restriction of  $\pi$  to  $\{\theta : R(\theta) \leq R^* + \delta\}$ .



$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] \right] &\leq \inf_{\delta > 0} \left[ \mathbb{E}_{\theta \sim \pi_{\delta}} [R(\theta)] + \frac{\text{KL}(\pi_{\delta} \| \pi)}{\lambda} + \frac{\lambda}{8n} \right] \\ &\leq \inf_{\delta > 0} \left[ R^* + \delta + \frac{\log \frac{1}{\pi\{\theta : R(\theta) \leq R^* + \delta\}}}{\lambda} + \frac{\lambda}{8n} \right]. \end{aligned}$$

In the previous example,

$$\log \frac{1}{\pi\{\theta : R(\theta) \leq R^* + \delta\}} \leq d \log \frac{C}{\delta}$$

and we obtained a bound in  $\sqrt{d/n} \log(n/d)$ .

Definition : the **prior mass condition** is satisfied if there are  $C, D > 0$  such that, for any  $\delta > 0$  small enough,

$$\log \frac{1}{\pi\{\theta : R(\theta) \leq R^* + \delta\}} \leq D \log \frac{C}{\delta}.$$

## Theorem - excess risk bound

- Assume the prior mass condition with  $C, D > 0$ .
- Fix  $\lambda = \sqrt{n/D} \log(D/n)$ , and let  $\hat{\pi}_\lambda$  be the Gibbs posterior.

$$\mathbb{E}_S \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R(\theta)] \right] \leq R^* + \mathcal{O} \left( \sqrt{\frac{D}{n}} \log \frac{n}{D} \right).$$

## Fast Rates



Seeger, M. (2002). PAC-Bayesian generalisation error bounds for Gaussian process classification. *Journal of Machine Learning Research*.



Maurer, A. (2004). *A note on the PAC-Bayesian theorem*. Arxiv preprint arXiv :cs/0411099.



Tolstikhin, I. and Seldin, Y. (2013). PAC-Bayes-empirical-Bernstein inequality. *NeurIPS*.

## Tolstikhin and Seldin's PAC-Bayes bound, 2013

With proba. at least  $1 - \delta$ , for any  $\rho$ ,

$$\begin{aligned}\mathbb{E}_{\theta \sim \rho}[R(\theta)] &\leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \\ &\quad + \sqrt{2\mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \frac{\text{KL}(\rho \parallel \pi) + \log \frac{2\sqrt{n}}{\delta}}{n}} \\ &\quad + 2 \frac{\text{KL}(\rho \parallel \pi) + \log \frac{2\sqrt{n}}{\delta}}{n}.\end{aligned}$$

## Tolstikhin and Seldin's PAC-Bayes bound, 2013

With proba. at least  $1 - \delta$ , for any  $\rho$ ,

$$\begin{aligned}\mathbb{E}_{\theta \sim \rho}[R(\theta)] &\leq \mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \\ &\quad + \sqrt{2\mathbb{E}_{\theta \sim \rho}[R_n(\theta)] \frac{\text{KL}(\rho \parallel \pi) + \log \frac{2\sqrt{n}}{\delta}}{n}} \\ &\quad + 2 \frac{\text{KL}(\rho \parallel \pi) + \log \frac{2\sqrt{n}}{\delta}}{n}.\end{aligned}$$

Key observation : if  $\rho$  is such that  $\mathbb{E}_{\theta \sim \rho}[R_n(\theta)] = 0$ , then

$$\mathbb{E}_{\theta \sim \rho}[R(\theta)] \leq 2 \frac{\text{KL}(\rho \parallel \pi) + \log \frac{2\sqrt{n}}{\delta}}{n}.$$

It appears that such fast rates can occur beyond the  $\mathbb{E}_{\theta \sim \rho}[R_n(\theta)] = 0$  case.

**Example 1** : classification,  $\ell(y, f_\theta(x)) = 1_{y \neq f_\theta(x)}$ .

Recall :

- $\eta(x) = \mathbb{P}_S(Y = 1|X = x)$ ,
- the “Bayes classifier”  $f^*(x) = 1_{\eta(x) \geq 1/2}$ .

### Mammen and Tsybakov margin assumption

- $\mathbb{P}_S(|\eta(X) - 1/2| < \tau) = 0$  for some small enough  $\tau > 0$ .
- there is  $\theta^*$  such that  $f^* = f_{\theta^*}$ .



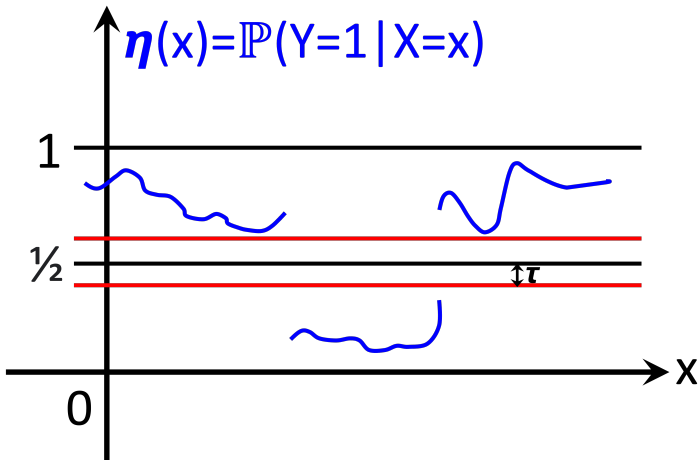
Mammen, E. and Tsybakov, A. B. (1999). Smooth discrimination analysis. *The Annals of Statistics*.



Tsybakov, A. B. (2003). Optimal rates of aggregation. *COLT*.

Under the margin assumption, they prove fast rates in  $\frac{1}{n}$  for various predictors.

$\mathbb{P}_{\mathcal{S}}(|\eta(X) - 1/2| < \tau) = 0$  for some small enough  $\tau > 0$ .



## Example 2 : “strongly convex, Lipschitz loss”.

For short, put  $g(\theta) = g_{x,y}(\theta) = \ell(y, f_\theta(x))$ .

Assume there are  $L, \alpha > 0$  such that, for any  $x, y$ , there is a  $\delta(\cdot, \cdot) = \delta_{x,y}(\cdot, \cdot) \geq 0$  with

$$\forall \theta, \frac{g(\theta) + g(\theta^*)}{2} - g\left(\frac{\theta + \theta^*}{2}\right) \geq \frac{1}{2\alpha} \delta^2(\theta, \theta^*),$$

$$\text{and } g(\theta) - g(\theta^*) \leq L\delta(\theta, \theta^*).$$



Bartlett, P. L., Jordan, M. I. and McAuliffe, J. D. (2003). Convexity, classification, and risk bounds. *Journal of the American Statistical Association*.

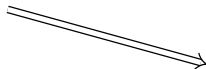
→ fast rates also in this case.

## Definition – Bernstein condition

Bernstein condition is satisfied with constant  $K$  if

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \left[ \left( \ell(Y, f_{\theta}(X)) - \ell(Y, f_{\theta^*}(X)) \right)^2 \right] \\ \leq K \underbrace{\mathbb{E}_{\mathcal{S}} \left[ \ell(Y, f_{\theta}(X)) - \ell(Y, f_{\theta^*}(X)) \right]}_{=R(\theta) - R^*}. \end{aligned}$$

Mammen & Tsybakov margin assumption



Bernstein condition  $\implies$  fast rates



Bartlett et al. convexity condition

## Proof that Example 2 $\Rightarrow$ Bernstein.

Put  $g(\theta) = \ell(Y, f_\theta(X))$ . Then

$$\begin{aligned}\mathbb{E}_S \left[ \left( g(\theta) - g(\theta^*) \right)^2 \right] &\leq L^2 \mathbb{E}_S \left[ \delta^2(\theta, \theta^*) \right]^2 \\ &\leq 2\alpha L^2 \mathbb{E}_S \left[ \frac{g(\theta) + g(\theta^*)}{2} - g\left(\frac{\theta + \theta^*}{2}\right) \right] \\ &= 2\alpha L^2 \left[ \frac{R(\theta) + R(\theta^*)}{2} - R\left(\frac{\theta + \theta^*}{2}\right) \right] \\ &\leq 2\alpha L^2 \left[ \frac{R(\theta) + R(\theta^*)}{2} - R(\theta^*) \right] \\ &= \alpha L^2 [R(\theta) - R(\theta^*)].\end{aligned}$$

Let us re-examine elements of yesterday's proof.

Reminder : lemma on exponential moment of risk

$$\mathbb{E}_{\mathcal{S}} \exp \left\{ -\frac{\lambda}{n} [R(\theta) - R_n(\theta)] \right\} \leq \exp \left( \frac{\lambda^2}{8n} \right).$$

$$\mathbb{E}_{\mathcal{S}} \exp \left\{ -\frac{\lambda}{n} [R(\theta) - R_n(\theta)] - \frac{\lambda^2}{8n} \right\} \leq 1.$$

$$\mathbb{P}_{\mathcal{S}} \left\{ R(\theta) \leq R_n(\theta) + \frac{\lambda}{8n} + \frac{\log \frac{1}{\delta}}{n} \right\} \leq \delta.$$

$\lambda \sim \sqrt{n}$  is the best choice...

## Alternative bound based on Bernstein's inequality

$$\begin{aligned}\mathbb{E}_{\mathcal{S}} \exp \left\{ -\frac{\lambda}{n} [R(\theta) - R(\theta^*) - R_n(\theta) + R_n(\theta^*)] \right\} \\ \leq \exp \left( \frac{\lambda^2 \text{Var}[\ell(Y, f_{\theta}(X)) - \ell(Y, f_{\theta^*}(X))]}{2(n - 2\lambda)} \right).\end{aligned}$$

Constraint :  $\lambda \leq n/4$ , and use Bernstein assumption

$$\begin{aligned}\mathbb{E}_{\mathcal{S}} \exp \left\{ -\frac{\lambda}{n} [R(\theta) - R(\theta^*) - R_n(\theta) + R_n(\theta^*)] \right\} &\leq \exp \left( \frac{\lambda^2 K [R(\theta) - R(\theta^*)]}{n} \right). \\ \Rightarrow \mathbb{P}_{\mathcal{S}} \left\{ \left( 1 - \frac{\lambda K}{n} \right) [R(\theta) - R(\theta^*)] \leq R_n(\theta) - R(\theta^*) + \frac{\log \frac{1}{\delta}}{\lambda} \right\} &\leq \delta.\end{aligned}$$

We can take  $\lambda = cn$  for  $c$  small enough...

## PAC-Bayes oracle inequality under Bernstein condition

- Assume Bernstein condition is satisfied with constant  $K$ .

Put  $\lambda = n / \max(2K, 1)$ ,

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] - R^* \right] \\ \leq 2 \inf_{\rho \in \mathcal{M}(\Theta)} \left[ \mathbb{E}_{\theta \sim \rho} [R(\theta)] - R^* + \frac{\max(2K, 1) \text{KL}(\rho \| \pi)}{n} \right]. \end{aligned}$$

- Assume moreover the prior mass condition with  $C, D > 0$ .

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_{\lambda}} [R(\theta)] - R^* \right] \leq \frac{2 \max(2K, 1) D}{n} \log \left( \frac{e C n}{D} \right).$$

For a full proof, see



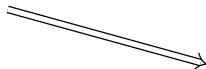
Alquier, P. (2024). *User-Friendly Introduction to PAC-Bayes bounds*. Foundations and Trends in Machine Learning, Now Publishers.

## Definition – Bernstein condition

Bernstein condition is satisfied with constant  $K$  if

$$\mathbb{E}\left[\left(\ell(Y, f_{\theta}(X)) - \ell(Y, f_{\theta^*}(X))\right)^2\right] \leq K \underbrace{\mathbb{E}\left[\ell(Y, f_{\theta}(X)) - \ell(Y, f_{\theta^*}(X))\right]}_{=R(\theta) - R^*}.$$

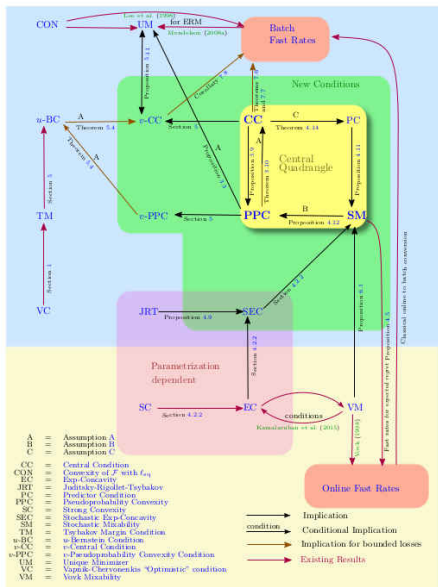
Mammen & Tsybakov margin assumption



Bernstein condition  $\implies$  fast rates



Bartlett et al. convexity condition



Van Erven, T., Grünwald, P., Mehta, N., Reid, M. and Williamson, R. (2015). Fast Rates in Statistical and Online Learning. *JMLR*.

## Mutual Information Bounds

## Reminder – Catoni's PAC-Bayes bound, 2003

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any randomized estimator  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda}{8n}.$$

For  $\lambda$  and  $\pi$  are fixed, this motivated the introduction of the Gibbs posterior  $\hat{\rho} = \hat{\pi}_\lambda$ , that minimizes the r.h.s. Then, we applied the bound in expectation to derive rates of convergence :

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R(\theta)] \right] \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R_n(\theta)] + \frac{\text{KL}(\hat{\pi}_\lambda \parallel \pi)}{\lambda} + \frac{\lambda}{8n} \right].$$

But... why did we keep the same  $\lambda$  and  $\pi$  ?

## PAC-Bayes bound in expectation – v2.0

- Fix  $\Lambda > 0$ ,  $\Pi$  and the randomized estimator  $\hat{\rho}$  (for example  $\hat{\rho} = \hat{\pi}_\lambda$ ).

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta)] \right] \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \Pi)}{\Lambda} + \frac{\Lambda}{8n} \right].$$

Thus,

$$\begin{aligned} & \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta)] \right] \\ & \leq \inf_{\Lambda > 0} \inf_{\Pi \in \mathcal{M}(\Theta)} \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \Pi)}{\Lambda} + \frac{\Lambda}{8n} \right] \\ & = \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] \right] + \inf_{\Lambda > 0} \inf_{\Pi \in \mathcal{M}(\Theta)} \mathbb{E}_{\mathcal{S}} \left[ \frac{\text{KL}(\hat{\rho} \parallel \Pi)}{\Lambda} + \frac{\Lambda}{8n} \right]. \end{aligned}$$

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] \right] + \inf_{\Lambda > 0} \inf_{\Pi \in \mathcal{M}(\Theta)} \mathbb{E}_{\mathcal{S}} \left[ \frac{\text{KL}(\hat{\rho} \| \Pi)}{\Lambda} + \frac{\Lambda}{8n} \right],$$

the infimum is reached, as shown by :



Catoni, O. (2007). *PAC-Bayesian supervised learning : the thermodynamics of statistical learning*.  
 IMS lecture notes – monograph series.

$$\mathbb{E}_{\mathcal{S}} \text{KL}(\hat{\rho} \| \Pi) = \mathbb{E}_{\mathcal{S}} \text{KL}(\hat{\rho} \| \mathbb{E}_{\mathcal{S}} \hat{\rho}) + \underbrace{\text{KL}(\mathbb{E}_{\mathcal{S}} \hat{\rho} \| \Pi)}_{=0 \text{ if } \Pi = \mathbb{E}_{\mathcal{S}} \hat{\rho}}.$$

- $\mathbb{E}_{\mathcal{S}} \hat{\rho} \in \mathcal{M}(\Theta)$  defined by  $[\mathbb{E}_{\mathcal{S}} \hat{\rho}](E) = \mathbb{E}_{\mathcal{S}} [\hat{\rho}(E)]$ .
- the first term in the r.h.s. has a nice interpretation...

Let  $(U, V) \sim P$ . Let  $P_U$  and  $P_V$  denote their marginals. If  $U$  and  $V$  were independent,  $P = P_U \otimes P_V$ .

## Mutual information between two random variables

$$\mathcal{I}(U, V) := \text{KL}(P \| P_U \otimes P_V).$$

Note :  $\mathcal{I}(U, V)$  depends on the distribution  $P$  of  $(U, V)$ , not on  $(U, V)$ . This is confusing... remember that  $\mathbb{E}(U)$  is not a function of  $U$ !

## Proposition

$$\mathcal{I}(U, V) = \mathbb{E}_U \left[ \text{KL}(P_{V|U} \| P_V) \right].$$

Thus,

$$\mathbb{E}_S \text{KL}(\hat{\rho} \| \Pi) = \underbrace{\mathbb{E}_S \text{KL}(\hat{\rho} \| \mathbb{E}_S \hat{\rho})}_{=: \mathcal{I}(\theta, S)} + \underbrace{\text{KL}(\mathbb{E}_S \hat{\rho} \| \Pi)}_{=0 \text{ if } \Pi = \mathbb{E}_S \hat{\rho}}.$$

$$\begin{aligned} & \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{p}} [R(\theta)] \right] \\ & \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{p}} [R_n(\theta)] \right] + \inf_{\Lambda > 0} \inf_{\Pi \in \mathcal{M}(\Theta)} \mathbb{E}_{\mathcal{S}} \left[ \frac{\text{KL}(\hat{p} \parallel \Pi)}{\Lambda} + \frac{\Lambda}{8n} \right] \\ & = \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{p}} [R_n(\theta)] \right] + \inf_{\Lambda > 0} \left[ \frac{\mathcal{I}(\theta, \mathcal{S})}{\Lambda} + \frac{\Lambda}{8n} \right]. \end{aligned}$$

## Mutual information bound

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{p}} [R(\theta)] \right] \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{p}} [R_n(\theta)] \right] + \sqrt{\frac{\mathcal{I}(\theta, \mathcal{S})}{2n}}.$$

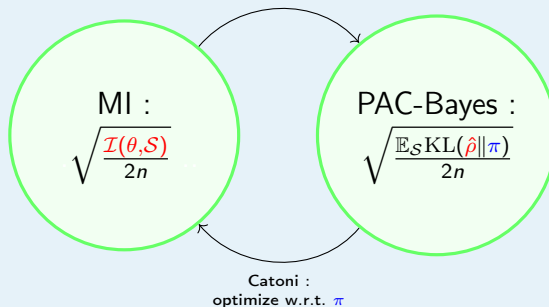


Russo, D. and Zou, J. (2019). How much does your data exploration overfit? controlling bias via information usage. *IEEE Transactions on Information Theory*.

## Mutual information bound

$$\mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R(\theta)] \right] \leq \mathbb{E}_{\mathcal{S}} \left[ \mathbb{E}_{\theta \sim \hat{\rho}} [R_n(\theta)] \right] + \sqrt{\frac{\mathcal{I}(\theta, \mathcal{S})}{2n}}.$$

$$\mathcal{I}(\theta, \mathcal{S}) = \mathbb{E}_{\mathcal{S}} \text{KL}(\hat{\rho} \| \mathbb{E}_{\mathcal{S}} \hat{\rho}) \leq \mathbb{E}_{\mathcal{S}} \text{KL}(\hat{\rho} \| \pi)$$



Example of improvement thanks to MI bounds :

### Theorem - excess risk bound based on PAC-Bayes

- Assume the prior mass condition with  $C, D > 0$ .
- Fix  $\lambda = \sqrt{n/D} \log(D/n)$ , and let  $\hat{\pi}_\lambda$  be the Gibbs posterior.

$$\mathbb{E}_S \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R(\theta)] \right] \leq R^* + \mathcal{O} \left( \sqrt{\frac{D}{n}} \log \frac{n}{D} \right).$$

### Informal statement - excess risk bound based on MI

$$\mathbb{E}_S \left[ \mathbb{E}_{\theta \sim \hat{\pi}_\lambda} [R(\theta)] \right] \leq R^* + \mathcal{O} \left( \sqrt{\frac{D}{n}} \right).$$

For more details, see :



Alquier, P. (2024). *User-Friendly Introduction to PAC-Bayes bounds*. Foundations and Trends in Machine Learning, Now Publishers.

- 1 Introduction : Generalization Bounds and PAC-Bayes
  - Generalization Bounds
  - Minimization of the PAC-Bayes Bound
  - Catoni's PAC-Bayes Bound and its Proof
- 2 Excess Risk Bounds and Mutual Information Bounds
  - From Empirical Bounds to Learning Rates
  - Fast Rates
  - Mutual Information Bounds
- 3 Empirical PAC-Bayes Bounds for Time Series
  - Mathematical Tools
  - PAC-Bayes Bounds for Time Series
  - Towards Empirical Bounds

From now on, we no longer assume that  $(X_1, Y_1), \dots, (X_n, Y_n)$  are independent.

- We still assume that they are indentially distributed. Actually, we will work under the assumption that  $(X_t, Y_t)_{t \in \mathbb{Z}}$  is a stationary time series.
- **Special case** :  $(X_t)_{t \in \mathbb{Z}}$  is a time series, and  $Y_t = X_{t+1}$ .

Some mathematical tools.

Let  $\mathcal{A} \subset \mathcal{E}$  and  $\mathcal{B} \subset \mathcal{E}$  be  $\sigma$ -fields. If  $\mathcal{A}$  and  $\mathcal{B}$  are independent, then for  $A \in \mathcal{A}$  and  $B \in \mathcal{B}$ , we have

$$\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B).$$

### Definition – $\alpha$ -coefficient

$$\alpha(\mathcal{A}, \mathcal{B}) = \sup_{A \in \mathcal{A}, B \in \mathcal{B}} \left| \mathbb{P}(A \cap B) - \mathbb{P}(A)\mathbb{P}(B) \right|.$$

### Definition – $\varphi$ -coefficient

$$\varphi(\mathcal{A}, \mathcal{B}) = \sup_{A \in \mathcal{A}, B \in \mathcal{B}} \left| \mathbb{P}(B|A) - \mathbb{P}(B) \right|.$$

### Definition – $\beta$ -coefficient

$$\beta(\mathcal{A}, \mathcal{B}) = \mathbb{E} \left[ \sup_{B \in \mathcal{B}} \left| \mathbb{P}(B|\mathcal{A}) - \mathbb{P}(B) \right| \right].$$

## Proposition

- $0 \leq \alpha(\mathcal{A}, \mathcal{B}) \leq 1/4$ .
- $0 \leq \varphi(\mathcal{A}, \mathcal{B}) \leq 1$ .
- $2\alpha(\mathcal{A}, \mathcal{B}) \leq \beta(\mathcal{A}, \mathcal{B}) \leq \varphi(\mathcal{A}, \mathcal{B})$ .

Application : we can measure the dependence between the future and the past of a time series !

## $\alpha$ -mixing coefficients

Let  $U = (U_t)_{t \in \mathbb{N}}$  be a sequence of random variables. We put, for any  $k \in \mathbb{N}$ ,

$$\alpha_U(k) = \sup_{t \in \mathbb{N}} \alpha \left( \sigma(U_0, U_1, \dots, U_t), \sigma(U_{t+k}, U_{t+k+1}, \dots) \right).$$

Similar definitions for  $\beta_U(k)$  and  $\varphi_U(k)$ .

## Example : Markov chains.

- state space  $\mathcal{U}$ ,
- initial distribution :  $U_0 \sim \mu(\cdot)$ ,
- transition kernel :  $U_{t+1}|U_t = u \sim P(u, \cdot)$ .

Notations :

$$\mu P(\cdot) = \int \mu(\mathrm{d}u) P(u, \cdot),$$

$$P^2(u, \cdot) = \int P(u, \mathrm{d}v) P(v, \cdot), \text{ } P^k \text{ defined recursively.}$$

A proba. measure  $\mu$  is invariant iff :  $\mu P = \mu$ .

Consequence of these notations :

- $U_t \sim \mu P^t(\cdot)$ .
- if  $\mu$  is invariant :  $U_t \sim \mu P^t(\cdot) = \mu$ , the chain is stationary.

**Example continued.** If  $(U_t)_{t \in \mathbb{N}}$  is a Markov chain :

$$\begin{aligned}\varphi_U(k) &= \sup_{t \in \mathbb{N}} \varphi \left( \sigma(U_0, U_1, \dots, U_t), \sigma(U_{t+k}, U_{t+k+1}, \dots) \right) \\ &= \sup_{t \in \mathbb{N}} \varphi \left( \sigma(U_t), \sigma(U_{t+k}) \right).\end{aligned}$$

If  $\mu$  is invariant (stationary chain) :

$$\varphi_U(k) = \varphi \left( \sigma(U_0), \sigma(U_k) \right).$$

### Theorem – Davydov 1973

$$\begin{aligned}\varphi_U(k) &= \sup_{u \in \mathcal{U}} \|P^k(u, \cdot) - \mu(\cdot)\|_{\text{TV}}, \\ \beta_U(k) &= \sum_{u \in \mathcal{U}} \mu(u) \|P^k(u, \cdot) - \mu(\cdot)\|_{\text{TV}}.\end{aligned}$$

$$\varphi_U(k) = \sup_{u \in \mathcal{U}} \|P^k(u, \cdot) - \pi(\cdot)\|_{\text{TV}}, \quad \beta_U(k) = \sum_{u \in \mathcal{U}} \mu(u) \|P^k(u, \cdot) - \pi(\cdot)\|_{\text{TV}}.$$

## Definition

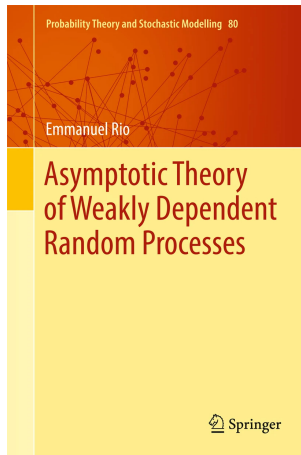
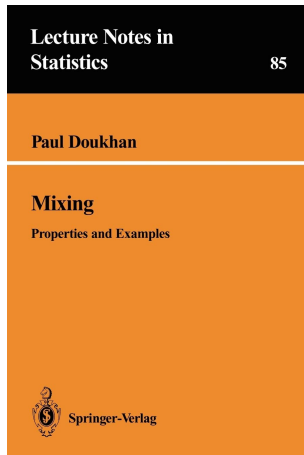
The chain is said to be **uniformly ergodic** if there is  $C > 0$  and  $\rho \in [0, 1)$  such that  $\sup_{u \in \mathcal{U}} \|P^k(u, \cdot) - \mu(\cdot)\|_{\text{TV}} \leq C\rho^k$ .

## Theorem

If  $\mathcal{U}$  is a finite set, and the chain  $(U_t)$  is aperiodic and irreducible, then it is uniformly ergodic.

In the infinite case, we can find examples such that :

- $\varphi_U(k) \leq C\rho^k$ ,
- or  $\beta_U(k) \leq C\rho^k$  but  $\varphi_U(k) = C > 0$ .



Example of the results you can prove with these coefficients :

### Theorem (LLN and CLT under $\beta$ -mixing)

Assume  $(U_t)_{t \in \mathbb{N}}$  is a stationary series with  $\sum_{k=0}^{\infty} \beta_U(k) < \infty$ .  
Let  $g$  be a real-valued function such that  $\mathbb{E}[g^2(U_1)] < \infty$ .

Then :

- $\frac{1}{n} \sum_{t=1}^n g(U_t) \xrightarrow[n \rightarrow \infty]{\text{proba.}} \mathbb{E}[g(U_1)]$ .
- we have  $V(g) = \sum_{t=0}^{\infty} \text{Cov}(g(U_0), g(U_t)) < +\infty$  and

$$\sqrt{n} \left( \frac{1}{n} \sum_{t=1}^n g(U_t) - \mathbb{E}[g(U_1)] \right) \xrightarrow[n \rightarrow \infty]{\text{dist.}} \mathcal{N}(0, V(g)).$$

The (weaker) assumption  $\sum_{k=0}^{\infty} \alpha_U(k) < \infty$  alone implies only LLN.

## Reminder : exponential moment of the risk (i.i.d. case)

$$\forall \lambda \in \mathbb{R}, \mathbb{E}_{\mathcal{S}} \exp \left\{ \lambda \left[ R(\theta) - R_n(\theta) \right] \right\} \leq \exp \left( \frac{\lambda^2}{8n} \right).$$

## Exponential moment of the risk – $\varphi$ -mixing case

Put  $U_t = (X_t, Y_t)$  and assume  $\Sigma := \sum_{k=1}^{\infty} \varphi_U(k) < +\infty$ . For any  $\lambda > 0$ ,

$$\mathbb{E}_{\mathcal{S}} \exp \left\{ \lambda \left[ R(\theta) - R_n(\theta) \right] \right\} \leq \exp \left( \frac{\lambda^2}{8n(1 + 2\Sigma)} \right).$$

Direct application of



Rio, E. (2000). Inégalités de Hoeffding pour les fonctions lipschitziennes de suites dépendantes.  
*Compte Rendus de L'Académie des Sciences.*

## PAC-Bayes Bounds for Time Series

## Theorem – essentially from Alquier and Li (2012)

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\begin{aligned} & \mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \\ & \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda(1 + 2 \sum_{\ell=1}^{\infty} \varphi_{\ell}(U))^2}{8n}. \end{aligned}$$

- Pros : we can also prove a bound in expectation, and derive rate of convergences, etc.
- Cons : in practice, we don't know  $\varphi_{\ell}(U)$  so the result is useless for practical purposes. Similar results proven by other authors suffer the same issue.



Fard, M. M., Pineau, J. and Szepesvári, C. (2012). PAC-Bayesian policy evaluation for reinforcement learning. *UAI*.



Alquier, P. and Wintenberger, O. (2012). Model selection for weakly dependent time series forecasting. *Bernoulli*.



Alquier, P. and Li, X. (2012). Prediction of quantiles by statistical learning and application to GDP forecasting. *15th International Conference on Discovery Science*.



Alquier, P., Li, X. and Wintenberger, O. (2013). Prediction of time series by statistical learning : general losses and fast rates. *Dependence Modeling*.



Alquier, P. and Guedj, B. (2018). Simpler PAC-Bayes bounds for hostile data. *Machine Learning*.



Banerjee, I., Rao, V. A. and Honnappa, H. (2021). PAC-Bayes bounds on variational tempered posteriors for Markov models. *Entropy*.



Haussmann, M., Gerwinn, S., Look, A., Rakitsch, B. and Kandemir, M. (2021). Learning partially known stochastic dynamics with empirical PAC Bayes. *AISTATS*.



Eringis, D., Leth, J.-J., Tan, Z.-H., Wisniewski, R., Esfahani, A. F. and Petreczky, M. (2021). PAC-Bayesian theory for stochastic LTI systems. *60th IEEE Conference on Decision and Control*.



Eringis, D., Leth, J.-J., Tan, Z.-H., Wisniewski, R. and Petreczky, M. (2024). PAC-Bayesian error bound, via Rényi divergence, for a class of linear time-invariant state-space models. *ICML*.



Boroujeni, M.G., Galimberti, C.L., Krause, A. and Ferrari-Trecate, G. (2025). *PAC-Bayesian Optimal Control with Stability and Generalization Guarantees*. ArXiv preprint arXiv :2512.02858.

## Towards Empirical Bounds (in the Markov case)



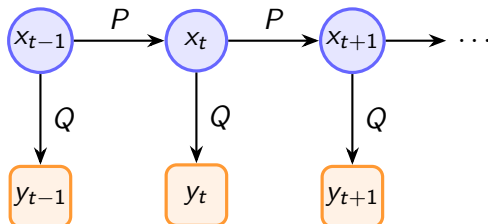
Karagulyan, V. and Alquier, P. (2026). Empirical PAC-Bayes bounds for Markov chains. *AISTATS*.



From now, we assume :

- $(X_t)_t$  is a stationary Markov chain with kernel  $P$ ,
- $Y_t = F(X_t, \varepsilon_t)$  for some i.i.d. sequence  $(\varepsilon_t)$ .

This implies :  $Y_t | X_t = x$  has the distribution of  $F(x, \varepsilon_t)$ . We can let  $Q(x, \cdot)$  denote this distribution.



Note that this makes  $U_t = (X_t, Y_t)$  a Markov chain with kernel

$$P_U((x, y), d(x', y')) = P(x, dx')Q(x', dy').$$

## Definition : reversal kernel

Assume  $P$  has a unique invariant  $\mu$ , we put :

$$P^*(u, dv) = \frac{P(v, du)\mu(dv)}{\mu(du)}.$$

- interpretation : in the stationary regime, the distribution of  $(X_t, X_{t+1})$  is

$$\mu(x_t)P(x_t, dx_{t+1}) = P^*(x_{t+1}, dx_t)\mu(x_{t+1}).$$

- note that  $P^*P$  is also a Markov kernel.

## Definition

For a Markov kernel  $M$ , 1 is an eigenvalue of  $M$ . Let  $m_1$  denotes its multiplicity. Define :

$$\gamma(M) := \begin{cases} 0 & \text{if } m_1 > 1, \\ 1 - \sup \{ \lambda \in \text{sp}(M), \lambda < 1 \} & \text{otherwise.} \end{cases}$$

## Definition – pseudo spectral-gap

$$\gamma_{\text{ps}} := \max_{k \geq 1} \frac{\gamma((P^*)^k P^k)}{k}.$$



Paulin, D. (2015). Concentration inequalities for Markov chains by Marton couplings and spectral methods. *Electronic Journal of Probability*.

## Theorem

Fix  $0 < \lambda < \frac{n}{20}$  and a prior  $\pi$ . With probability at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda \gamma_{\text{ps}}} + \frac{8\lambda}{n}.$$

## Corollary

Assume we have an estimator  $\hat{\gamma}_{\text{ps}}$  of  $\gamma_{\text{ps}}$  such that

$$\mathbb{P} \left( \left| \frac{\hat{\gamma}_{\text{ps}}}{\gamma_{\text{ps}}} - 1 \right| > \epsilon \right) \leq \alpha(n, \epsilon, P) =: \alpha.$$

With probability at least  $1 - \delta - \alpha$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + (1 + \epsilon) \frac{\text{KL}(\hat{\rho} \| \pi) + \log \frac{1}{\delta}}{\lambda \hat{\gamma}_{\text{ps}}} + \frac{8\lambda}{n}.$$

## Example : finite state space $\{1, \dots, d\}$



Wolfer, G. and Kontorovich, A. (2019). Estimating the mixing time of ergodic markov chains. *COLT*.



Wolfer, G. and Kontorovich, A. (2024). Improved estimation of relaxation time in nonreversible Markov chains. *The Annals of Applied Probability*.

$$N^k(i, j) = \sum_{\ell=1}^{\lfloor \frac{n-1}{k} \rfloor} \mathbf{1} [X_{1+(\ell-1)k} = i, X_{1+\ell k} = j],$$

$$N^k(i) = \sum_{\ell=1}^{\lfloor \frac{n-1}{k} \rfloor} \mathbf{1} [X_{1+(\ell-1)k} = i], \text{ and } \hat{L}^k(i, j) = \frac{N^k(i, j)}{\sqrt{N^k(i)N^k(j)}}.$$

Note :  $\hat{L}^k$  estimates  $L^k = \text{diag}(\mu)^{\frac{1}{2}} P^k \text{diag}(\mu)^{-\frac{1}{2}}$ .

Fix  $K > 0$ , put  $\hat{\gamma}_{\text{ps}, K} = \max \left\{ \gamma((\hat{L}^k)^T \hat{L}^k), 1 \leq k \leq K \right\}$ .

## Proposition (Wolfer and Kontorovich)

For  $K = \lceil \frac{2}{\epsilon} \rceil$ , we have

$$\mathbb{P} \left( \left| \frac{\hat{\gamma}_{\text{ps}, K}}{\gamma_{\text{ps}}} - 1 \right| > \epsilon \right) \leq \alpha(n, \epsilon, P)$$

where

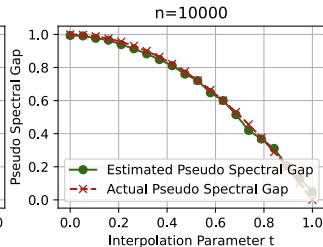
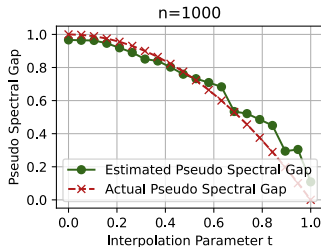
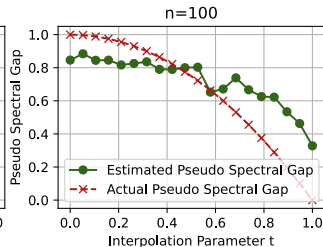
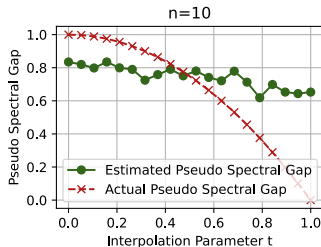
$$\alpha(n, \epsilon, P) = \frac{C_1(P)}{\epsilon} \exp(-n \gamma_{\text{ps}}^2 \epsilon^2 C_2(P)).$$

## Corollary

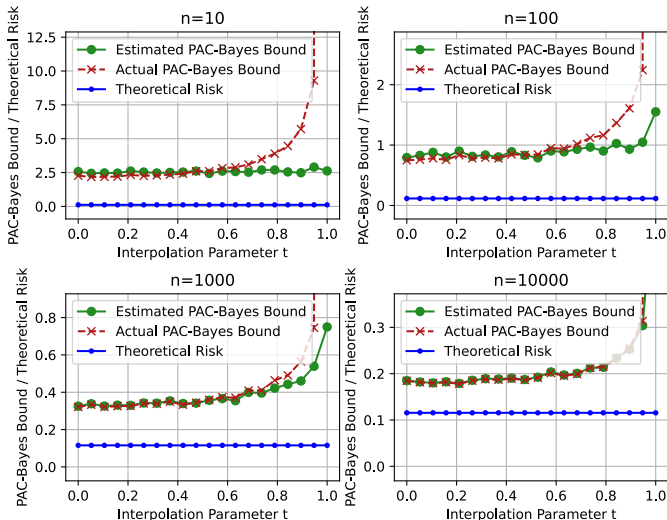
With probability at least  $1 - \delta - \alpha$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + (1 + \epsilon) \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda \hat{\gamma}_{\text{ps}, K}} + \frac{8\lambda}{n}.$$

dim=20



**dim=20**



- Infinite case : not possible to estimate  $\gamma_{\text{ps}}$  *in general*.
- However, under additional assumptions, this can be feasible.

### Proposition

Assume there is  $a \in [-1, 1]$  such that

$$X_{t+1} = aX_t + \varepsilon_t, \text{ with } \varepsilon_t \sim \mathcal{N}(0, 1).$$

Then

$$\gamma_{\text{ps}} = 1 - a^2.$$

## Reminder

Fix  $\lambda > 0$  and  $\pi$ . With proba. at least  $1 - \delta$  on  $\mathcal{S}$ , for any  $\hat{\rho}$ ,

$$\begin{aligned} & \mathbb{E}_{\theta \sim \hat{\rho}}[R(\theta)] \\ & \leq \mathbb{E}_{\theta \sim \hat{\rho}}[R_n(\theta)] + \frac{\text{KL}(\hat{\rho} \parallel \pi) + \log \frac{1}{\delta}}{\lambda} + \frac{\lambda(1 + 2 \sum_{\ell=1}^{\infty} \varphi_{\ell}(U))^2}{8n}. \end{aligned}$$

- Going beyond the Markov case will require to estimate the full sequence of  $\varphi$ -coefficients.
- This is an open problem !
- However, estimation of  $\alpha$  and  $\beta$ -mixing coefficients was done recently :



Khaleghi, A. and Lugosi, G. (2023). Inferring the mixing properties of a stationary ergodic process from a single sample-path. *IEEE Transactions on Information Theory*.

# մերսի