

Sequential Learning, Sequential Optimization (= online machine learning)



Lecture notes in English

vs. course on the blackboard delivered in French

How to reach me:

Gilles Stoltz

CNRS senior researcher at the maths department of Université Paris-Sud

gilles.stoltz@math.u-psud.fr

Lecture notes posted at: <http://stoltz.perso.math.cnrs.fr>

Tentative schedule:

Thursdays 2pm - 4.30pm (or maybe 4.45pm...)

October 3

Lesson #1

October 10

Lesson #2

! October 17

No class!

October 24

Lesson #3

! October 31

No class!

November 7

Lesson #4

November 14

Lesson #5

! November 21

Mid-term exam

(to be decided)

→ November 28

Lesson #6 or No class

December 5

Lesson #7 or Lesson #6

! December 12

No class!

December 19

Lesson #8 or Lesson #7

January 9

Final exam in all cases.

Exams :

Not open book

but { 5 two-sided A4 size "cheat sheets" mid-term exam
10 allowed for the final exam

→ Write your own summary of the course
(= summarize the proof techniques)

Outline / Content :

1. Adversarial online learning (aka: prediction of individual sequences)
2. Stochastic bandits
3. Adversarial bandits (if time permits)

Path followed for each (sub)topic:

- Definition of the setting and of an objective (eg: regret)
- Statement + analysis of an algorithm (preferably computationally efficient)
 - ↳ eg: upper bound on its regret
- Lower bounds on the objective, holding for any algorithm
- No simulation, no application to real data (I could do it but don't want to)
 - ↳ 100% Theoretical with course (definition - Review - proof)

Sequential setting #1 (meta-statistical)

* At each round $t=1, 2, \dots, T$:

- experts output forecasts f_{jt} , $j \in \{1, \dots, N\}$
- statistician aggregates the forecasts as $\hat{y}_t = \sum_j p_{jt} f_{jt}$
where $P_t = (p_{1t} \dots p_{Nt})$ is a convex weight vector
- true observation y_t is revealed

* Assessment of performance: loss function $\ell(\hat{y}_t, y_t)$

↳ cumulative loss $\sum_{t=1}^T \ell(\hat{y}_t, y_t)$

No stochastic model → relative-performance criterion = perform almost as well as the best constant convex combination

$$\text{ie, control } R_T = \sum_{t=1}^T \ell(\hat{y}_t, y_t) - \inf_{q \in \mathcal{X}} \sum_{t=1}^T \ell\left(\sum_{j=1}^N q_j f_{jt}, y_t\right)$$

$\uparrow$ the simplex of all convex weight vectors $\uparrow$ independent of t

Ex: square loss $\ell(\hat{y}_t, y_t) = (\hat{y}_t - y_t)^2$

absolute loss $\ell(\hat{y}_t, y_t) = |\hat{y}_t - y_t|$

absolute percentage of error $\ell(\hat{y}_t, y_t) = |\hat{y}_t - y_t| / |y_t|$

↳ all these loss functions are convex in $\hat{y}_t$, a fact that we will need later on.

* R_T is called the regret

We have some bias/variance decomposition: $\sum_{t=1}^T \ell(\hat{y}_t, y_t) = \inf_q \sum_{t=1}^T \ell\left(\sum_j q_j f_{jt}, y_t\right) + R_T$

We will focus on the regret in these lectures

$\uparrow$
cumulative loss

$\uparrow$
approximation error (= how good the experts are)

$\uparrow$
the regret corresponds to a sequential estimation error

Sequential setting #2

(sequential optimization)

* At each round $t = 1, 2, \dots, T$:

- Mathematician picks $p_t \in \mathcal{X}$
- true loss function l_t is revealed

* Assessment of performance

$$R_T = \sum_{t=1}^T l_t(p_t) - \inf_{q \in \mathcal{X}} \sum_{t=1}^T l_t(q)$$

* Link: $l_t: p \mapsto l(\sum_j p_j f_{jt}, y_t)$

The only (small) difference is that the f_{jt} are only available, in this version, AFTER p_t is selected.

* Special case of interest to start with:

LINEAR loss functions

$$l_t(p) = \sum_{j=1}^N p_j l_{jt} \quad \text{where} \quad l_{jt} \stackrel{\text{not}}{=} l_t(\delta_j)$$

l_t is given by a vector $(l_{t1}, \dots, l_{tN})$,
choosing l_t is choosing such a vector.

↑
Dirac mass
on j

Setting #2/
Linear losses

→ Exponentially weighted average predictor [EWA]

* Algorithm:

$$p_1 = (1/N, \dots, 1/N)$$

formula
also ok
for $t=1$

$$t \geq 2, \quad p_{jt} = \frac{\exp(-\eta \sum_{s=1}^{t-1} l_{js})}{\sum_{k=1}^N \exp(-\eta \sum_{s=1}^{t-1} l_{ks})}$$

where $\eta > 0$ is a learning rate.

Note: Not all weight on the best experts (it wouldn't work!) but most of the weight.

* Bound:

For all choices of linear loss functions $(l_{1t}, \dots, l_{Nt}) \in [m, M]^N$

$$\begin{aligned} R_T &= \sum_{t=1}^T \sum_{j=1}^N p_{jt} l_{jt} - \min_{q \in \mathcal{X}} \sum_{t=1}^T \sum_{j=1}^N q_j l_{jt} \\ &= \sum_{t=1}^T \sum_{j=1}^N p_{jt} l_{jt} - \min_{k=1, \dots, N} \sum_{t=1}^T l_{kt} \leq \frac{\ln N}{\eta} + \eta \frac{(M-m)^2}{8} T \end{aligned}$$

PROOF:

based on Hoeffding's lemma:

- for random variables $X \in [m, M]$,

$$\forall \eta \in \mathbb{R}, \quad \ln \mathbb{E}[e^{\eta X}] \leq \eta \mathbb{E}[X] + \frac{\eta^2}{8} (M-m)^2$$

- thus for all convex weight vectors $q \in \mathcal{X}$, for all $l_{jt} \in [m, M]$,

$$\forall \eta \in \mathbb{R}, \quad \ln \sum_j q_j e^{\eta l_{jt}} \leq -\eta \sum_j q_j l_{jt} + \frac{\eta^2}{8} (M-m)^2$$

In particular,

$$\sum_j p_{jt} l_{jt} \leq -\frac{1}{\eta} \ln \sum_j p_{jt} e^{\eta l_{jt}} + \frac{\eta}{8} (M-m)^2$$

$$\begin{aligned} &= \ln \frac{\sum_j \exp(-\eta \sum_{s=1}^{t-1} l_{js})}{\sum_k \exp(-\eta \sum_{s=1}^{t-1} l_{ks})} \\ &\quad \uparrow \\ &\text{by definition} \\ &\text{of the algorithm} \end{aligned}$$

Summing over t (telescoping sum in the right-hand side):

$$\sum_{t=1}^T \sum_{j=1}^N p_{jt} l_{jt} \leq -\frac{1}{\eta} \ln \frac{\sum_{j=1}^N \exp(-\eta \sum_{t=1}^T l_{jt})}{N} + \frac{\eta}{8} (M-m)^2 T$$

Conclude using $\ln \sum_j e^{-\eta \sum_{t=1}^T l_{jt}} \geq \ln \max_j e^{-\eta \sum_{t=1}^T l_{jt}} = \max_j \ln e^{-\eta \sum_{t=1}^T l_{jt}} = \max_j \{-\eta \sum_{t=1}^T l_{jt}\}$

and (LINEARITY): $\min_{j=1..N} \sum_{t=1}^T l_{jt} = \min_{q \in \mathcal{X}} \sum_{t=1}^T \sum_k q_k l_{kt}$ $\perp$

Reminder :

PROOF OF Hoeffding's Lemma:

$\Psi(\eta) = \ln \mathbb{E}[e^{\eta X}]$ defined for all $\eta \in \mathbb{R}$

$$\Psi'(\eta) = \frac{\mathbb{E}[X e^{\eta X}]}{\mathbb{E}[e^{\eta X}]}$$

(differentiability: if X bounded thus $\eta \mapsto X e^{\eta X}$ locally dominated by an integrable r.v. independent of η)

(Ψ twice differentiable: similar reasons)

$$\Psi''(\eta) = \frac{\mathbb{E}[X^2 e^{\eta X}] \mathbb{E}[e^{\eta X}] - (\mathbb{E}[X e^{\eta X}])^2}{(\mathbb{E}[e^{\eta X}])^2} = \text{Var}_{\mathbb{Q}}(X)$$

under the probability $\mathbb{Q}$ defined by

$$\frac{d\mathbb{Q}}{d\mathbb{P}}(w) = \frac{e^{\eta X(w)}}{\mathbb{E}[e^{\eta X}]}$$

$$X \in [m, M] : \text{Var}_{\mathbb{Q}}(X) = \inf_{\mu \in \mathbb{R}} \mathbb{E}_{\mathbb{Q}}[(X-\mu)^2] \leq \mathbb{E}_{\mathbb{Q}}[(X - \frac{M+m}{2})^2] \leq \frac{(M-m)^2}{4}$$

Taylor: Ψ is actually C^2

$$\exists x \text{ s.t. } \Psi(\eta) = \overbrace{\Psi(0)}^{=0} + \eta \Psi'(0) + \frac{\eta^2}{2} \overbrace{\Psi''(x)}$$

$$\text{i.e. } \ln \mathbb{E}[e^{\eta X}] \leq \eta \mathbb{E}[X] + \frac{\eta^2}{8} (M-m)^2 \quad \perp$$

* How good is the obtained bound?

$$\min_{\eta > 0} \frac{\ln N}{\eta} + \frac{\eta}{8} (M-m)^2 T = (M-m) \sqrt{\frac{T \ln N}{2}} \quad \text{for } \eta^* = \frac{1}{(M-m)} \sqrt{\frac{8 \ln N}{T}}$$

↑
nice because $o(T)$ and weak dependence in N

↑
not nice when T, m, M unknown

Why bother? (exercise) #1

Why not put all mass on the best j in EWA?

Why only compare the performance to the best global expert?

How good is EWA?

1) Consider the strategy

$$p_1 = (1/N \dots 1/N)$$

$$p_t \leftrightarrow \begin{cases} \text{equal mass on } k \text{ s.t. } \sum_{s=t-1}^T \ell_{ks} = \min_j \sum_{s=t-1}^T \ell_{js} \\ \text{zero mass on } k \text{ s.t. } \sum_{s=t-1}^T \ell_{ks} > \max_j \sum_{s=t-1}^T \ell_{js} \end{cases}$$

called: "follows the leader(s)" (FIL)

Shows that it fails: there exist sequences $(\ell_{1t}, \dots, \ell_{Nt}) \in [0, 1]^N$ such that $\exists \delta > 0, \forall T, R_T = \sum_{t=1}^T \sum_{j=1}^N p_{jt} \ell_{jt} - \min_{k=1 \dots N} \sum_{t=1}^T \ell_{kt} \geq \delta T$.

EWA appears as a smoothed version of FIL!

2) Can we be more ambitious? People often hope to get close to in statistics when they resort to aggregation.

Shows that no strategy can be such that for all sequences $(\ell_{1t}, \dots, \ell_{Nt}) \in [0, 1]^N$, $R_T = \sum_{t,j} p_{jt} \ell_{jt} - \sum_t \min_k \ell_{kt} = o(T)$.

3) Show that for EWA, the regret is never negative:

$\forall (\ell_{1t}, \dots, \ell_{Nt}) \in [m, M]^N, R_T = \sum_{j,t} p_{jt} \ell_{jt} - \min_{k=1 \dots N} \sum_{t=1}^T \ell_{kt} \geq 0$

Application of the EWA forecaster / Sion's lemma.

Statement: Let X, Y two convex sets, $f: X \times Y \rightarrow [0, M]$
 a function s.t. $\forall x \in X, f(x, \cdot)$ is concave
 $\forall y \in Y, f(\cdot, y)$ is convex
 then (under additional regularity assumptions):

$$\inf_{x \in X} \sup_{y \in Y} f(x, y) = \sup_{y \in Y} \inf_{x \in X} f(x, y).$$

Proof:

1) $\geq$ always holds: $\forall x, \forall y, f(x, y) \geq \inf_{x' \in X} f(x', y)$

taking $\sup_{y \in Y}$ in both sides: $\forall x, \sup_{y \in Y} f(x, y) \geq \sup_{y \in Y} \inf_{x' \in X} f(x', y)$

Get $\inf \sup f \geq \sup \inf f$ by taking the $\inf_{x \in X}$ (the right-hand side is a constant independent of x).

2) A (fictitious) statistician and a (fictitious) opponent play as follows:

First, the statistician sets $N \geq 2$ and $x^{(1)} \dots x^{(N)}$ in X , as well as $T \geq 1$

Then, at each round, they simultaneously pick

$$x_t = \sum_{j=1}^N p_{jt} x^{(j)} \in X \quad \text{and} \quad y_t \in Y$$

How?

$$p_{jt} = \frac{\exp(-\eta \sum_{s=1}^{t-1} f(x^{(j)}, y_s))}{\sum_{k=1}^N \exp(-\eta \sum_{s=1}^{t-1} f(x^{(k)}, y_s))}$$

with $\eta = \frac{1}{M} \sqrt{\frac{\ln N}{T}}$

and (since p_{jt} only depends on the past, the opponent can compute it and pick:)

$$y_t \text{ s.t. (by definition of sup)} \quad f(x_t, y_t) \geq \sup_{y \in Y} f(x_t, y) - \frac{1}{T}$$

By definition of the exponentially weighted average strategy with a well-chosen learning rate $\eta \geq 0$

$$\sum_{t=1}^T \sum_{j=1}^N p_{jt} \underbrace{f(x^{(j)}, y_t)}_{\substack{\text{corresponds} \\ \text{to } y_t \in [0, M]}} - \min_{k=1, \dots, N} \sum_{t=1}^T f(x^{(k)}, y_t) \leq \underbrace{\frac{\ln N}{\eta} + \eta \frac{M^2}{8} T}_{= M \sqrt{\frac{T}{2} \ln N}} \quad \left. \begin{array}{l} \text{given} \\ \text{choice} \\ \text{for } \eta \end{array} \right\}$$

From the convexity of $f(\cdot, y_t)$,
we finally get:

$$\sum_{t=1}^T f\left(\underbrace{\sum_j p_{jt} x^{(j)}}_{= \bar{x}_t \text{ def.}}, y_t\right) - \min_{k=1, \dots, N} \sum_{t=1}^T f(x^{(k)}, y_t) \leq M \sqrt{\frac{T}{2} \ln N} \quad (*)$$

3) Now, $\inf_x \sup_y f(x, y) \leq \sup_y f(\bar{x}, y)$ where $\bar{x} = \frac{1}{T} \sum_{t=1}^T x_t$

$$\leq \sup_y \frac{1}{T} \sum_{t=1}^T f(x_t, y) \quad \left. \begin{array}{l} f(\cdot, y) \\ \text{convex} \\ \text{by } y \end{array} \right\}$$

$$\stackrel{\substack{\sup \sum \\ \leq \sum \sup}}{\leq} \frac{1}{T} \sum_{t=1}^T \sup_y f(x_t, y) \leq \frac{1}{T} \sum_{t=1}^T \sup_y f(x_t, y_t) \stackrel{\substack{\text{def of} \\ y_t}}{=} \frac{1}{T} \sum_{t=1}^T f(x_t, y_t) + \frac{1}{T}$$

then, by (*):

$$\begin{aligned} \frac{1}{T} \sum_{t=1}^T f(x_t, y_t) &\leq \underbrace{M \sqrt{\frac{\ln N}{2T}}}_{= O(1/\sqrt{T})} + \min_{k=1, \dots, N} \frac{1}{T} \sum_{t=1}^T f(x^{(k)}, y_t) \quad \left. \begin{array}{l} \text{by } (*) \end{array} \right\} \\ &\leq O(1/\sqrt{T}) + \min_k f(x^{(k)}, \bar{y}) \quad \left. \begin{array}{l} \text{by concavity} \\ \text{of } f(x^{(k)}, \cdot), \\ \text{where } \bar{y} = \frac{1}{T} \sum_{t=1}^T y_t \end{array} \right\} \\ &\leq \sup_{j \in [N]} \min_k f(x^{(k)}, y) \end{aligned}$$

4) In sections (2)-(3) T , N and $x^{(1)} \dots x^{(N)}$ were fixed, but
we can play with them! We proved

$$\inf_x \sup_y f(x, y) \leq \frac{1}{T} + O(1/\sqrt{T}) + \sup_y \min_k f(x^{(k)}, y)$$

Letting $T \rightarrow +\infty$:

$$\inf_x \sup_y f(x, y) \leq \sup_y \min_k f(x^{(k)}, y)$$

This holds for all N and all $x^{(1)} \dots x^{(N)}$ in X :

$$\inf_x \sup_y f(x, y) \leq \inf_{N \geq 1} \inf_{\{x^{(1)} \dots x^{(N)}\} \subset X} \sup_y \min_{k=1 \dots N} f(x^{(k)}, y)$$

is clearly $\sup_y \inf_x f(x, y)$ but maybe $\leq \sup_y \inf_x f(x, y)$ so that we have an equality?
 We will now state and use regularity / topological assumptions.

- Assume
- X, Y are metric spaces, with distances d_x and d_y
 - $f: X \times Y \rightarrow [a, M]$ uniformly continuous:
 $\forall \varepsilon > 0, \exists \delta > 0 \mid d_x(x, x') + d_y(y, y') \leq \delta \Rightarrow |f(x, y) - f(x', y')| \leq \varepsilon$
 - X compact: $\forall \delta > 0, \exists N$ and $x^{(1)} \dots x^{(N)}$ s.t.
 $X \subset \bigcup_{j=1}^N B(x^{(j)}, \delta)$

Given $\varepsilon > 0$ and the associated $\delta > 0$:

$$\forall x, y, \quad f(x, y) \geq \min_{j=1 \dots N} f(x^{(j)}, y) - \varepsilon \quad \text{by uniform continuity}$$

Taking $\inf_x$ then $\sup_y$:

$$\sup_y \inf_x f(x, y) \geq \sup_y \min_j f(x^{(j)}, y) - \varepsilon$$

$$\geq \inf_N \inf_{\{x^{(j)}\}} \sup_y \min_j f(x^{(j)}, y)$$

and we let $\varepsilon \downarrow 0$

EXERCISE #2

Bonus points at the exam

(Perhaps you can find even better - weaker - assumptions? If so, let me know!
 ↑ while still having a small and easy-to-read proof...)

Strongly convex loss functions $\rightarrow$ Same predictor, faster rates

(Still Setting #2, but ℓ_t only assumed to be convex, not linear)

Exp-concavity: The chosen loss functions ℓ_t are η -exp-concave if $\forall t, e^{-\eta \ell_t}$ is concave on X

Remarks:

- $x \mapsto x^\alpha$ is concave and increasing for $0 < \alpha \leq 1$; thus η -exp-concave functions are also η' -exp-concave for all $0 < \eta' \leq \eta$

- $-\ln$ is convex and decreasing: exp-concave functions are in particular convex functions

Examples:

- Investment in the stock market, without transaction costs

1-exp-concave $\left\{ \begin{array}{l} \ell_t(p) = -\ln\left(\sum_{j=1}^n p_j x_{jt}\right) \\ \text{where } x_{jt} \text{ denotes the multiplicative evolution of stock } j \\ \text{The } -\ln \text{ comes because we consider a cumulative loss (not a reward)} \end{array} \right.$

- The square loss $\ell_t(p) = \left(\sum_j p_j f_{jt} - y_t\right)^2$

If $f_{jt}, y_t \in [0, B]$ $\forall j, t$, then the ℓ_t are $1/2B^2$ -exp concave; it suffices to prove that

$\forall y \in [0, B], \psi_y: x \in [0, B] \mapsto e^{-\eta(x-y)^2}$ is concave for $\eta = 1/2B^2$

But $\psi_y'(x) = -2\eta(x-y)e^{-\eta(x-y)^2}$

$\psi_y''(x) = (4\eta^2(x-y)^2 - 2\eta)e^{-\eta(x-y)^2}$ is ≤ 0

as soon as $2\eta(x-y)^2 \leq 1$, which is guaranteed

by $\eta \leq 1/2B^2$.

Same predictor:

with $\ell_{jt} \stackrel{\text{def}}{=} \ell_t(\delta_j)$; ie $p_{jt} \propto e^{-\eta \sum_{s=1}^{t-1} \ell_{js}}$

Theorem: If ℓ_t is η -exp. concave, then the exponentially weighted average predictor used with this η is such that

$$\sum_{t=1}^T \ell_t(p_t) - \min_{k=1 \dots N} \sum_{t=1}^T \ell_t(\delta_k) \leq \frac{\ln N}{\eta}$$

Ex: For all sequences of bounded forecasts and observations

$$\forall t \quad f_t \in [0, B] \quad \text{and} \quad \ell_t, y_t \in [0, B],$$

$$\frac{1}{T} \sum_{t=1}^T \left(\sum_j p_{jt} f_{jt} - y_t \right)^2 \leq \frac{2B^2 \ln N}{T} + \min_{k=1 \dots N} \frac{1}{T} \sum_{t=1}^T (f_{kt} - y_t)^2$$

In particular, taking $\sqrt{\cdot}$ and using $\sqrt{a+b} \leq \sqrt{a} + \sqrt{b}$:

$$\forall \text{ sequences, } \text{RMSE of meta-predictor on stages } 1, \dots, T \leq \text{RMSE of best predictor} + B \sqrt{\frac{2 \ln N}{T}}$$

↳ Hence the naive "prediction of individual sequences" or "robust aggregation of forecasts"

PROOF: $\ell_t(p_t) = -\frac{1}{\eta} \ln \underbrace{e^{-\eta \ell_t(p_t)}}_{\geq \sum_j p_{jt} e^{-\eta \ell_{jt}}} \leq -\frac{1}{\eta} \ln \sum_j p_{jt} e^{-\eta \ell_{jt}}$

Same telescoping argument:

$$\sum_{t=1}^T \ell_t(p_t) \leq -\frac{1}{\eta} \ln \frac{\sum_{j=1}^N e^{-\eta \sum_{t=1}^T \ell_{jt}}}{N}$$

and same way of concluding as before. $\lrcorner$

What about the more challenging comparison to convex combinations?

$$\inf_{p \in \mathcal{X}} \sum_{t=1}^T \ell_t(p) \quad ?$$

Exp-concavity & comparison to convex losses

Algorithm:
(continuous EWA prediction)

$$p_t = \frac{\int_{\mathcal{X}} p e^{-\eta \sum_{s=1}^{t-1} \ell_s(p)} d\mu(p)}{\int_{\mathcal{X}} e^{-\eta \sum_{s=1}^{t-1} \ell_s(p)} d\mu(p)} \stackrel{\text{not.}}{=} \int_{\mathcal{X}} p d\mu_t(p)$$

where μ is the uniform (Lebesgue) measure on $\mathcal{X}$.

Remarks:

- $p_1 = (1/N \dots 1/N)$
- how to better grasp μ ? $\mathcal{X}$ is included in an affine hyperspace, in which it has a positive Lebesgue measure $\rightarrow \mu$ is a conditional measure of that
- one can show that if $X_1 \dots X_{N-1}$ iid $\sim \mathcal{U}_{[0,1]}$ then $0 \leq X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(N-1)} \leq 1$ (order statistics) define N segments $(X_{(1)}, X_{(2)} - X_{(1)}, \dots, 1 - X_{(N-1)})$ whose joint law is μ .
- how to compute p_t ?
 - $\hookrightarrow$ grid arguments (with S intervals) have a complexity of $O(S^{N+1})$
 - $\hookrightarrow$ Monte-Carlo methods? Simulate $P_1 \dots P_m$ iid $\sim \mu$ and use $\mathbb{E}_g, \int_{\mathcal{X}} g(p) d\mu(p) \approx \frac{1}{m} \sum_{j=1}^m g(P_j)$
 - $\hookrightarrow$ MCMC methods, see Vempala 1996.

in any case,
a nightmare
in practice

(I had to
implement it once,

see Stoltz & Lugosi 2005)

Theorem: For all sequence l_t of η -exp-concave functions, the continuous EWA predictor satisfies:

$$\sum_{t=1}^T l_t(p_t) - \inf_{p \in \mathcal{X}} \sum_{t=1}^T l_t(p) \leq \frac{1}{\eta} (1 + (N-1) \ln(T\eta))$$

PROOF:

$$\begin{aligned} l_t(p_t) &= -\frac{1}{\eta} \ln e^{-\eta l_t(p_t)} \\ &= -\frac{1}{\eta} \ln e^{-\eta l_t(\int_{\mathcal{X}} p d\mu_t(p))} \\ &\geq \int_{\mathcal{X}} e^{-\eta l_t(p)} d\mu_t(p) \quad \text{by exp-concavity} \end{aligned}$$

thus

$$\begin{aligned} l_t(p_t) &\leq -\frac{1}{\eta} \ln \int_{\mathcal{X}} e^{-\eta l_t(p)} d\mu_t(p) \\ &= -\frac{1}{\eta} \ln \frac{\int_{\mathcal{X}} e^{-\eta \sum_{s=1}^t l_s(p)} d\mu(p)}{\int_{\mathcal{X}} e^{-\eta \sum_{s=1}^{t-1} l_s(p)} d\mu(p)} \end{aligned}$$

Telescoping argument:

$$\sum_{t=1}^T l_t(p_t) \leq \cancel{\frac{\ln 1}{\eta}} - \frac{1}{\eta} \ln \int_{\mathcal{X}} e^{-\eta \sum_{t=1}^T l_t(p)} d\mu(p)$$

$$\left[\text{Fix } \delta > 0, \text{ let } p_\delta^* \text{ be s.t. } \inf_p \sum_{t=1}^T l_t(p) \leq \delta + \sum_{t=1}^T l_t(p_\delta^*) \right]$$

Why is the $\inf$ not a $\min$? l_t is convex thus continuous on the interior of $\mathcal{X}$ but can be discontinuous on its border

$$\Delta_{\delta, \varepsilon}^* = \left\{ (1-\varepsilon) p_\delta^* + \varepsilon r, \quad r \in \mathcal{X} \right\} \quad \text{for } \varepsilon \in]0, 1[$$

→ for such $p = (1-\varepsilon) p_\delta^* + \varepsilon r$, we have by exp-concavity:

$$\forall t, \quad e^{-\eta l_t(p)} \geq (1-\varepsilon) e^{-\eta l_t(p_\delta^*)} + \underbrace{\varepsilon e^{-\eta l_t(r)}}_{\geq 0}$$

→ also, $\mu(\Delta_{\delta, \varepsilon}^*) = \varepsilon^{N-1}$

y. properties of the Lebesgue measure (and ambient dimension $N-1$)

Thus

$$\int_{\mathcal{X}} e^{-\eta \sum_{t=1}^T \ell_t(p)} d\mu(p) \geq (1-\varepsilon)^T \varepsilon^{N-1} e^{-\eta \sum_{t=1}^T \ell_t(p_s^*)}$$

Substituting in the bound:

$$\frac{T}{\sum_{t=1}^T \ell_t(p_t)} \leq -\frac{1}{\eta} \ln((1-\varepsilon)^T \varepsilon^{N-1}) + \frac{T}{\sum_{t=1}^T \ell_t(p_s^*)}$$

$$\leq \delta + \inf_{p \in \mathcal{X}} \frac{T}{\sum_{t=1}^T \ell_t(p)}$$

Let $\delta \downarrow 0$, we have proved:

$$\frac{T}{\sum_{t=1}^T \ell_t(p_t)} - \inf_{p \in \mathcal{X}} \frac{T}{\sum_{t=1}^T \ell_t(p)} \leq \inf_{\varepsilon \in (0,1)} -\frac{1}{\eta} \ln((1-\varepsilon)^T \varepsilon^{N-1})$$

Choosing, e.g., $\varepsilon = 1/T$:

$$\frac{1}{(1-\varepsilon)^T} = \left(1 - \frac{1}{T}\right)^{-T} = \left(\frac{T}{T-1}\right)^T = \left(1 + \frac{1}{T-1}\right)^T$$

$$= \exp\left(T \ln\left(1 + \frac{1}{T-1}\right)\right) \leq e$$

$\leq 1/T$

hence the stated bound. $\square$ A better bound:

$$B(\varepsilon) = T \ln(1-\varepsilon) + (N-1) \ln \varepsilon$$

$$B'(\varepsilon) = \frac{-T}{1-\varepsilon} + \frac{N-1}{\varepsilon}$$

$$B''(\varepsilon) < 0$$

thus the ε^* s.t. $B'(\varepsilon^*) = 0$ is the global minimum of B

$$\text{We have } B'(\varepsilon^*) = 0 \Leftrightarrow \frac{\varepsilon^*}{N-1} = \frac{1}{T} - \frac{\varepsilon^*}{T} \Leftrightarrow \varepsilon^* = \frac{1/T}{1/T + 1/(N-1)} = \frac{N-1}{T+N-1}$$

We thus get the sharper final bound

$$\frac{1}{\eta} \left(T \ln\left(\frac{T+N-1}{T}\right) + (N-1) \ln\left(\frac{T}{T+N-1}\right) \right) = \frac{T+N-1}{\eta} H\left(\frac{N-1}{T+N-1}\right)$$

$$= \frac{1}{\eta} \left((N-1) \ln\left(1 + \frac{N-1}{T}\right) + (N-1) \ln\left(1 + \frac{T}{N-1}\right) \right)$$

$\leq \frac{(N-1)^2}{T}$

where $H(x) = -x \ln x - (1-x) \ln(1-x)$ is the binary entropy

* However * both bounds are $\sim \frac{N-1}{\eta} \ln T$ as $T \rightarrow +\infty$. $\hookrightarrow$ And the second bound seems less readable to me...

Third case :Convex Loss Functions

$$\ell_t: X \rightarrow [m, M]$$

convex

1. Comparison to the best individual expert

Regret $R_T = \sum_{t=1}^T \ell_t(p_t) - \min_{i=1, \dots, N} \sum_{t=1}^T \ell_t(S_i)$ where S_i has mass at i

is bounded by convexity by

$$\hat{R}_T \leq \underbrace{\sum_{t=1}^T \sum_{j=1}^N p_{jt} \ell_{jt}}_{\text{linear upper bound, which we know how to control with EWA}} - \min_{i=1, \dots, N} \sum_{t=1}^T \ell_{it} \quad \text{where we denoted } \ell_{jt} = \ell_t(S_j)$$

2. Comparison to the best convex vector

Exercise #3: An inefficient algorithm to do so: the continuous EWA predictor

↳ We'll see a more efficient way (algorithm with a better computational complexity) + a better regret bound

1) Show that the continuous EWA strategy is such that:
Against all sequences $\ell: X \rightarrow [m, M]$ of convex functions,

$$\sum_{t=1}^T \ell_t(p_t) - \inf_{q \in X} \sum_{t=1}^T \ell_t(q) \leq \frac{\eta}{8} (M-m)^2 T + \inf_{q \in X} \frac{1}{\eta} \ln \left(e^{N-1} - \eta \varepsilon T (M-m) \right)$$

2) When T, M and m are known, which η is a good choice to minimize the obtained upper bound, and can you provide an easy-to-read final upper bound?