

LECTURE #17

CS 170

Spring 2021



Last time:

zero-sum games

row player and column player can be associated to dual LPs

value of the game is the optimum value of these LPs

(an example of duality)

Today:

experts problem

multiplicative weight updates

THEORY OF COMPUTING, Volume 8 (2012), pp. 121–164
www.theoryofcomputing.org

RESEARCH SURVEY

The Multiplicative Weights Update Method: A Meta-Algorithm and Applications

Sanjeev Arora*

Elad Hazan

Satyen Kale

Received: July 22, 2008; revised: July 2, 2011; published: May 1, 2012.

Abstract: Algorithms in varied fields use the idea of maintaining a distribution over a certain set and use the *multiplicative update rule* to iteratively change these weights. Their analyses are usually very similar and rely on an exponential potential function.

In this survey we present a simple meta-algorithm that unifies many of these disparate algorithms and derives them as simple instantiations of the meta-algorithm. We feel that since this meta-algorithm and its analysis are so simple, and its applications so broad, it should be a standard part of algorithms courses, like “divide and conquer.”

n experts (people whose advice you take or not)

T days (may or may not be known a priori)

Each day $t \in [T]$, you choose an expert $i(t) \in [n]$ (according to some rule), and incur a loss $f_{i(t)}^t \in [0, 1]$. (Losses are bounded.)

↑ you chose to "follow" the prediction of expert $i(t)$ and incurred a corresponding loss

expert	day 1		day 2		day 3		...	day T	
	am	pm	am	pm	am	pm		am	pm
1		f_1^1		f_1^2		f_1^3	...		f_1^T
2		f_2^1		f_2^2		f_2^3	...		f_2^T
3		f_3^1		f_3^2		f_3^3	...		f_3^T
⋮		⋮		⋮		⋮	...		⋮
n		f_n^1		f_n^2		f_n^3	...		f_n^T
(possibly probabilistic) choices: $i(1)$ $i(2)$ $i(3)$... $i(T)$									

a day's losses are adversarially set at beginning of day but revealed after our choice

→ total loss = $\sum_{t=1}^T f_{i(t)}^t$

(it is a random variable if the choices are probabilistic)

Ex: select stocks based on opinions of n other people

GOAL: minimize (expected) total loss

Observation: all experts could be idiots so we cannot expect to design a selection strategy that always achieves small total loss

We refine the goal as follows:

minimize the (expected) regret $R = \mathbb{E} \left[\left(\sum_{t=1}^T f_{i(t)}^t \right) \right] - \left(\min_{i \in [n]} \sum_{t=1}^T f_i^t \right)$
(across all adversarial losses)

That is, we minimize the total loss wrt best expert in hindsight.

no need to consider a distribution here bc there is a best expert

• Why not minimize $\mathbb{E} \left[\left(\sum_{t=1}^T f_{i(t)}^t \right) \right] - \left(\sum_{t=1}^T \min_{i \in [n]} f_i^t \right)$?

A: We cannot hope for a selection strategy that competes with each day's best expert.

For each t , adversarially set $\begin{cases} \text{loss } 0 & \text{for expert with least probability on day } t \\ \text{loss } 1 & \text{for all other experts} \end{cases}$

This gives $\mathbb{E} \left[\left(\sum_{t=1}^T f_{i(t)}^t \right) \right] = \sum_{t=1}^T \sum_{i=1}^n p_i^t f_i^t \geq \sum_{t=1}^T \left(1 - \frac{1}{n} \right) = \frac{n-1}{n} \cdot T$

$\sum_{t=1}^T \min_{i \in [n]} f_i^t = \sum_{t=1}^T 0 = 0$

largest value for the smallest probability

difference is $\frac{n-1}{n} \cdot T$ (it's large)

Q: how to pick expert each day?

temporary simplification: binary losses $f_i^t \in \{0, 1\}$
right wrong

Try #1: always pick expert 1 ($i(t)=1 \forall t$, regardless of losses)

the losses could be

	1	2	...	T
1	1	1	...	1
2	0	0	...	0
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
n	0	0	...	0

which leads to $R \geq T$

Try #2: choose majority opinion (this is well-defined for binary predictions)

the losses could be

	1	2	...	T
1	0	0	...	0
2	1	1	...	1
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
n	1	1	...	1

which leads to $R \geq T$

Try #3: choose expert at random (intuition is to use randomness to defeat adversarial losses)

the (expected regret) is

$$[F_i := \sum_{t=1}^T f_i^t]$$

$$\sum_{t=1}^T \left(\sum_{i=1}^n \frac{1}{n} \cdot f_i^t \right) - \left(\min_{i \in [n]} \sum_{t=1}^T f_i^t \right) = \underbrace{\sum_{i=1}^n \frac{1}{n} \cdot F_i}_{\text{average}} - \underbrace{\min_{i \in [n]} F_i}_{\text{minimum}} \leq \frac{n-1}{n} \cdot T$$

The upper bound is tight (eg. $F_1=0, F_2=T, \dots, F_n=T$ as for the bad weights in Try #2)

Try #4: choose best expert so far (intuition is to take the past into account)
 [& break ties lexicographically]

↑ While losses can be adversarially chosen so that every day's best expert so far does poorly, this makes experts overall worse, reducing regret.

But can still arrange losses so that

$$\left. \begin{array}{l} \text{total loss} = T \\ \text{best expert's loss} = T/n \end{array} \right\} R = \frac{n-1}{n} \cdot T \text{ (still large)}$$

Here is the example with $n=3$:

expert	day 1			day 2			day 3			day 4			...
	am	pm	tot	am	pm	tot	am	pm	tot	am	pm	tot	
1		1	1		0	1		0	1		1	2	
2		0	0		1	1		0	1		0	1	
3		0	0		0	0		1	1		0	1	
choices:	1			2			3			1			...

The chosen expert has loss 1 each day $\Rightarrow$ total loss is T .

Each expert has loss 1 once every n days, and loss 0 all other days $\Rightarrow$ best expert's loss is T/n .

Note: the example can be tweaked to avoid abusing the tie-breaking rule by relying on fractional losses

Try #5: choose expert according to a weighted majority (this is well-defined for binary predictions)

Fix a parameter ϵ .

- initialization: set weights $w_1^0, w_2^0, \dots, w_n^0$ to 1

- expert choice at time t :

$$E_A^t = \{i \mid \text{expert } i \text{ predicts A on am of day } t\}$$

$$E_B^t = \{i \mid \text{expert } i \text{ predicts B on am of day } t\}$$

if $\sum_{i \in E_A^t} w_i^t \geq \sum_{i \in E_B^t} w_i^t$ then predict A; else predict B

- update: $w_i^{t+1} := w_i^t \cdot (1 - \epsilon)^{f_i^t}$

Theorem: $\forall \epsilon > 0$, $WM(\epsilon)$ achieves the following guarantee

$$\left(\sum_{i=1}^T f_{i(t)}^t \right) \leq 2(1+\epsilon) \left(\min_{i \in [n]} \sum_{t=1}^T f_i^t \right) + \frac{2 \ln(n)}{\epsilon}$$

The theorem gives a multiplicative guarantee.

But we can do even better!

MULTIPLICATIVE WEIGHT UPDATES (uses past losses and randomness)

⊗ in some references the update is

$$w_i^{t+1} := w_i^t \cdot (1 - \epsilon) f_i^t$$

for which a similar analysis applies

Fix a parameter ϵ .

- initialization: set weights $w_1^0, w_2^0, \dots, w_n^0$ to 1
- expert choice at time t : choose $i \in [n]$ w.p. $p_i^t := \frac{w_i^t}{\sum_{j=1}^n w_j^t}$
- update: $w_i^{t+1} := w_i^t \cdot (1 - \epsilon f_i^t)$ ⊗

Theorem: $\forall \epsilon \in (0, \frac{1}{2}]$ MWU(ϵ) achieves the following (expected) regret:

$$\left(\sum_{t=1}^T \langle p^t, f^t \rangle \right) - \left(\min_{i \in [n]} \sum_{t=1}^T f_i^t \right) \leq \epsilon \cdot T + \frac{\ln(n)}{\epsilon}$$

if losses are in $[a, b]$
then the bound becomes
 $(b-a) \cdot \left(\epsilon T + \frac{\ln(n)}{\epsilon} \right)$
by re-scaling

- The guarantee is better than an approx factor of ≈ 2 bc loss of best expert could grow with T .
- The regret per day tends to $\epsilon + o(1)$ as $T \rightarrow \infty$.
- If we know T in advance then we can set $\epsilon = \sqrt{\ln(n)/T}$ so that $R \leq 2\sqrt{T \ln(n)}$, and the regret per day is $2\sqrt{\frac{\ln(n)}{T}}$, which tends to 0 very quickly.

Proof is based on the potential function

$$\Phi^t := \sum_{i=1}^n w_i^t$$

① claim: $\Phi^T \leq n \cdot e^{-\epsilon \sum_{t=1}^T \langle p^t, f^t \rangle}$

$$\begin{aligned} \Phi^{t+1} &= \sum_{i=1}^n w_i^{t+1} = \sum_{i=1}^n w_i^t \cdot (1 - \epsilon f_i^t) \\ &= \sum_{i=1}^n (p_i^t \Phi^t) \cdot (1 - \epsilon f_i^t) = \Phi^t \sum_{i=1}^n p_i^t \cdot (1 - \epsilon f_i^t) \\ &= \Phi^t \cdot (\sum_{i=1}^n p_i^t - \epsilon \sum_{i=1}^n p_i^t f_i^t) \\ &= \Phi^t \cdot (1 - \epsilon \langle p^t, f^t \rangle) \\ &\leq \Phi^t e^{-\epsilon \langle p^t, f^t \rangle} \end{aligned}$$

if there is an expected loss
then the potential drops accordingly

Hence $\Phi^T \leq \Phi^0 \prod_{t=1}^T e^{-\epsilon \langle p^t, f^t \rangle} = n \cdot e^{-\epsilon \sum_{t=1}^T \langle p^t, f^t \rangle}$

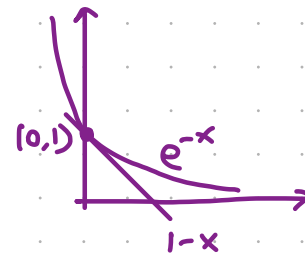
② claim: $\forall i \in [n] \quad \Phi^T \geq e^{-\epsilon \sum_{t=1}^T f_i^t - \epsilon^2 \sum_{t=1}^T (f_i^t)^2}$

if there is a good expert
then Φ^T cannot be too small

$$\Phi^T \geq w_i^T = \prod_{t=1}^T (1 - \epsilon f_i^t) \geq \prod_{t=1}^T e^{-\epsilon f_i^t - \epsilon^2 (f_i^t)^2} = e^{-\epsilon \sum_{t=1}^T f_i^t - \epsilon^2 \sum_{t=1}^T (f_i^t)^2}$$

Two inequalities:

① $1 - x \leq e^{-x}$



② $1 - x \geq e^{-x - x^2}$ for $x \leq \frac{1}{2}$

By Taylor expansion:

$$\begin{aligned} \ln(1-x) &= -x - \frac{x^2}{2} - \frac{x^3}{3} \dots \\ &\geq -x - x^2 \text{ for } x \leq \frac{1}{2} \end{aligned}$$

Proof is based on the potential function

$$\Phi^t := \sum_{i=1}^n w_i^t$$

① claim: $\Phi^T \leq n \cdot e^{-\epsilon \sum_{t=1}^T \langle p^t, f^t \rangle}$

② claim: $\forall i \in [n], \Phi^T \geq e^{-\epsilon \sum_{t=1}^T f_i^t - \epsilon^2 \sum_{t=1}^T (f_i^t)^2}$

Combining ① and ② we get that $\forall i \in [n]$:

$$n \cdot e^{-\epsilon \sum_{t=1}^T \langle p^t, f^t \rangle} \geq \Phi^T \geq e^{-\epsilon \sum_{t=1}^T f_i^t - \epsilon^2 \sum_{t=1}^T (f_i^t)^2}$$

$$\ln(n) - \epsilon \sum_{t=1}^T \langle p^t, f^t \rangle \geq -\epsilon \sum_{t=1}^T f_i^t - \epsilon^2 \sum_{t=1}^T (f_i^t)^2$$

$$\ln(n) + \epsilon^2 \sum_{t=1}^T (f_i^t)^2 \geq \epsilon \left(\sum_{t=1}^T \langle p^t, f^t \rangle - \sum_{t=1}^T f_i^t \right)$$

take $\ln(\cdot)$

shuffle terms

As the inequality holds $\forall i \in [n]$ we deduce that:

$$\ln(n) + \epsilon^2 \sum_{t=1}^T (f_i^t)^2 \geq \epsilon \left(\sum_{t=1}^T \langle p^t, f^t \rangle - \min_{i \in [n]} \sum_{t=1}^T f_i^t \right) = \epsilon \cdot R$$

As losses are bounded in $[0,1]$ we have $\sum_{t=1}^T (f_i^t)^2 \leq T$. Hence

$$\ln(n) + \epsilon^2 \cdot T \geq \epsilon R \Rightarrow R \leq \epsilon \cdot T + \frac{\ln(n)}{\epsilon}$$

analysis holds even if losses $f^t = (f_i^t)_{i \in [n]}$ are adversarially chosen based on prior losses $f^1, \dots, f^{t-1}$, prior choices $i(1), \dots, i(t-1)$, and selection probabilities (all we need is

that f^t is indep of randomness to choose $i(t)$)