

COMS 4771 Lecture 26

1. Online learning

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy
- ▶ You have access to N “experts”: `weather.com`, `wunderground.com`, ..., `groundhog`
- ▶ Each expert makes it's own prediction, sunny or rainy

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy
- ▶ You have access to N “experts”: `weather.com`, `wunderground.com`, ..., `groundhog`
- ▶ Each expert makes it's own prediction, sunny or rainy
- ▶ Based on their prediction, you make your own prediction

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy
- ▶ You have access to N “experts”: `weather.com`, `wunderground.com`, ..., `groundhog`
- ▶ Each expert makes it's own prediction, sunny or rainy
- ▶ Based on their prediction, you make your own prediction
- ▶ **Question:** without knowing which expert is the best, can you predict almost as well as the best expert?

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy
- ▶ You have access to N “experts”: `weather.com`, `wunderground.com`, ..., `groundhog`
- ▶ Each expert makes it's own prediction, sunny or rainy
- ▶ Based on their prediction, you make your own prediction
- ▶ **Question:** without knowing which expert is the best, can you predict almost as well as the best expert?
- ▶ More precisely: can you ensure that

$$\#(\text{mistakes you make}) \leq \#(\text{mistakes best expert makes}) + o(T)?$$

A SEQUENTIAL PREDICTION TASK

- ▶ Over a sequence of T days, you want to predict whether it will be sunny or rainy
- ▶ You have access to N “experts”: `weather.com`, `wunderground.com`, ..., `groundhog`
- ▶ Each expert makes it's own prediction, sunny or rainy
- ▶ Based on their prediction, you make your own prediction
- ▶ **Question:** without knowing which expert is the best, can you predict almost as well as the best expert?
- ▶ More precisely: can you ensure that

$$\#(\text{mistakes you make}) \leq \#(\text{mistakes best expert makes}) + o(T)?$$

- ▶ $\Rightarrow$ your error rate $\rightarrow$ best expert's error rate as $T \rightarrow \infty$.

PREDICTION SETUP

- ▶ Feature space: $\mathcal{X}$
- ▶ Label space: $\mathcal{Y}$

PREDICTION SETUP

- ▶ Feature space: $\mathcal{X}$
- ▶ Label space: $\mathcal{Y}$
- ▶ Loss function: $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$. E.g.
 - ▶ Zero-one loss in classification: $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$
 - ▶ Square loss in regression: $\ell(y, \hat{y}) = (\hat{y} - y)^2$
 - ▶ Absolute loss in regression: $\ell(y, \hat{y}) = |\hat{y} - y|$
- ▶ Given a labeled example $(x, y) \in \mathcal{X} \times \mathcal{Y}$, error of predicting $\hat{y}$ is measured by $\ell(y, \hat{y})$

PREDICTION SETUP

- ▶ Feature space: $\mathcal{X}$
- ▶ Label space: $\mathcal{Y}$
- ▶ Loss function: $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$. E.g.
 - ▶ Zero-one loss in classification: $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$
 - ▶ Square loss in regression: $\ell(y, \hat{y}) = (\hat{y} - y)^2$
 - ▶ Absolute loss in regression: $\ell(y, \hat{y}) = |\hat{y} - y|$
- ▶ Given a labeled example $(x, y) \in \mathcal{X} \times \mathcal{Y}$, error of predicting $\hat{y}$ is measured by $\ell(y, \hat{y})$
- ▶ Class $\mathcal{F}$ of functions $f : \mathcal{X} \rightarrow \mathcal{Y}$ (e.g. decision trees, nearest-neighbor classifiers, neural networks, linear predictor for regression, etc.)
- ▶ Informally: *Learning* over $\mathcal{F}$ means making predictions as well as the best predictor in $\mathcal{F}$.

ONLINE LEARNING FRAMEWORK

Online learning algorithm

For round $t = 1, 2, \dots, T$:

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.
4. Suffer loss $\ell(y_t, \hat{y}_t)$ (e.g., $\ell(y_t, \hat{y}_t) = \mathbb{1}\{\hat{y}_t \neq y_t\}$).

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.
4. Suffer loss $\ell(y_t, \hat{y}_t)$ (e.g., $\ell(y_t, \hat{y}_t) = \mathbb{1}\{\hat{y}_t \neq y_t\}$).

Underlying (non-probabilistic) assumption

The function class $\mathcal{F}$ models the data well:

i.e. best-in-hindsight (average) loss $\min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t))$ is small.

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.
4. Suffer loss $\ell(y_t, \hat{y}_t)$ (e.g., $\ell(y_t, \hat{y}_t) = \mathbb{1}\{\hat{y}_t \neq y_t\}$).

Underlying (non-probabilistic) assumption

The function class $\mathcal{F}$ models the data well:

i.e. best-in-hindsight (average) loss $\min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t))$ is small.

All “data” (e.g., x_t, y_t) could be controlled by an adversary!

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.
4. Suffer loss $\ell(y_t, \hat{y}_t)$ (e.g., $\ell(y_t, \hat{y}_t) = \mathbb{1}\{\hat{y}_t \neq y_t\}$).

Underlying (non-probabilistic) assumption

The function class $\mathcal{F}$ models the data well:

i.e. best-in-hindsight (average) loss $\min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t))$ is small.

All “data” (e.g., x_t, y_t) could be controlled by an adversary!

Goal: perform (almost) as well as the best-in-hindsight $f \in \mathcal{F}$.

ONLINE LEARNING FRAMEWORK

Online learning algorithm

For round $t = 1, 2, \dots, T$:

1. Observe data point $x_t \in \mathcal{X}$.
2. Choose prediction $\hat{y}_t \in \mathcal{Y}$.
3. Observe label $y_t \in \mathcal{Y}$.
4. Suffer loss $\ell(y_t, \hat{y}_t)$ (e.g., $\ell(y_t, \hat{y}_t) = \mathbb{1}\{\hat{y}_t \neq y_t\}$).

Underlying (non-probabilistic) assumption

The function class $\mathcal{F}$ models the data well:

i.e. best-in-hindsight (average) loss $\min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t))$ is small.

All “data” (e.g., x_t, y_t) could be controlled by an adversary!

Goal: perform (almost) as well as the best-in-hindsight $f \in \mathcal{F}$.

$$\begin{array}{lcl} \text{(Average) Regret:} & \frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) & - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t)) \\ & \text{your avg. loss} & \text{avg. loss of best } f \in \mathcal{F} \end{array}$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t									
$\hat{y}_t$									
y_t									

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7								
$\hat{y}_t$									
y_t									

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7								
$\hat{y}_t$	0								
y_t									

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7								
$\hat{y}_t$	0								
y_t	1								

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8							
$\hat{y}_t$	0								
y_t	1								

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8							
$\hat{y}_t$	0	1							
y_t	1								

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8							
$\hat{y}_t$	0	1							
y_t	1	1							

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1						
$\hat{y}_t$	0	1							
y_t	1	1							

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1						
$\hat{y}_t$	0	1	1						
y_t	1	1							

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1						
$\hat{y}_t$	0	1	1						
y_t	1	1	0						

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1	9	2	6	3	1	4
$\hat{y}_t$	0	1	1	1	1	1	1	0	1
y_t	1	1	0	1	0	1	1	0	0

(Average) **Regret:**
$$\frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t))$$

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1	9	2	6	3	1	4
$\hat{y}_t$	0	1	1	1	1	1	1	0	1
y_t	1	1	0	1	0	1	1	0	0

$$\begin{aligned} \text{(Average) Regret:} \quad & \frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t)) \\ & = \quad \frac{4}{9} \quad - \quad \frac{1}{9} \end{aligned}$$

$(f(x) = \mathbb{1}\{x > 2\})$ makes just one mistake.)

EXAMPLE

$\mathcal{X} = \mathbb{N}$, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) := \mathbb{1}\{\hat{y} \neq y\}$, $\mathcal{F}$ = threshold functions.

t	1	2	3	4	5	6	7	8	9
x_t	7	8	1	9	2	6	3	1	4
$\hat{y}_t$	0	1	1	1	1	1	1	0	1
y_t	1	1	0	1	0	1	1	0	0

$$\begin{aligned} \text{(Average) Regret:} \quad & \frac{1}{9} \sum_{t=1}^9 \ell(y_t, \hat{y}_t) \quad - \quad \min_{f \in \mathcal{F}} \frac{1}{9} \sum_{t=1}^9 \ell(y_t, f(x_t)) \\ & = \quad \frac{4}{9} \quad - \quad \frac{1}{9} \end{aligned}$$

$(f(x) = \mathbb{1}\{x > 2\})$ makes just one mistake.)

Goal: Regret $\rightarrow 0$ as $T \rightarrow \infty$ (i.e., “vanishing regret” or “no regret”).

SOME EASY CASES

SOME EASY CASES

Simplifying assumptions

- ▶ $N := |\mathcal{F}|$ is finite, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$ (zero-one loss).
- ▶ *Realizability*: assume there is a $f \in \mathcal{F}$ with $y_t = f(x_t)$ for all t .

SOME EASY CASES

Simplifying assumptions

- ▶ $N := |\mathcal{F}|$ is finite, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$ (zero-one loss).
- ▶ *Realizability*: assume there is a $f \in \mathcal{F}$ with $y_t = f(x_t)$ for all t .

So **Regret** is simply Average Loss (i.e., average number of mistakes):

$$\frac{1}{T} \sum_{t=1}^T \mathbb{1}\{\hat{y}_t \neq y_t\}.$$

SOME EASY CASES

Simplifying assumptions

- ▶ $N := |\mathcal{F}|$ is finite, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$ (zero-one loss).
- ▶ *Realizability*: assume there is a $f \in \mathcal{F}$ with $y_t = f(x_t)$ for all t .

So **Regret** is simply Average Loss (i.e., average number of mistakes):

$$\frac{1}{T} \sum_{t=1}^T \mathbb{1}\{\hat{y}_t \neq y_t\}.$$

Trivial algorithm: follow-the-leader

In round t , pick any $f_t \in \mathcal{F}$ that is *perfect-so-far*, and choose $\hat{y}_t := f_t(x_t)$.

SOME EASY CASES

Simplifying assumptions

- ▶ $N := |\mathcal{F}|$ is finite, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$ (zero-one loss).
- ▶ *Realizability*: assume there is a $f \in \mathcal{F}$ with $y_t = f(x_t)$ for all t .

So **Regret** is simply Average Loss (i.e., average number of mistakes):

$$\frac{1}{T} \sum_{t=1}^T \mathbb{1}\{\hat{y}_t \neq y_t\}.$$

Trivial algorithm: follow-the-leader

In round t , pick any $f_t \in \mathcal{F}$ that is *perfect-so-far*, and choose $\hat{y}_t := f_t(x_t)$.

What is the average loss after T rounds?

SOME EASY CASES

Simplifying assumptions

- ▶ $N := |\mathcal{F}|$ is finite, $\mathcal{Y} = \{0, 1\}$, $\ell(y, \hat{y}) = \mathbb{1}\{\hat{y} \neq y\}$ (zero-one loss).
- ▶ *Realizability*: assume there is a $f \in \mathcal{F}$ with $y_t = f(x_t)$ for all t .

So **Regret** is simply Average Loss (i.e., average number of mistakes):

$$\frac{1}{T} \sum_{t=1}^T \mathbb{1}\{\hat{y}_t \neq y_t\}.$$

Trivial algorithm: follow-the-leader

In round t , pick any $f_t \in \mathcal{F}$ that is *perfect-so-far*, and choose $\hat{y}_t := f_t(x_t)$.

What is the average loss after T rounds?

Answer: every mistake eliminates one function. So at most $N - 1$ mistakes, and average loss is at most $\frac{N-1}{T}$.

A BETTER ALGORITHM

“Halving” algorithm

In round t ,

$$\hat{y}_t := \arg \max_{y \in \{0,1\}} |\{\text{perfect-so-far } f : f(x_t) = y\}|$$

(i.e., go with majority's prediction).

A BETTER ALGORITHM

“Halving” algorithm

In round t ,

$$\hat{y}_t := \arg \max_{y \in \{0,1\}} |\{\text{perfect-so-far } f : f(x_t) = y\}|$$

(i.e., go with majority's prediction).

What is the average loss after T rounds?

A BETTER ALGORITHM

“Halving” algorithm

In round t ,

$$\hat{y}_t := \arg \max_{y \in \{0,1\}} |\{\text{perfect-so-far } f : f(x_t) = y\}|$$

(i.e., go with majority's prediction).

What is the average loss after T rounds?

Answer: every mistake eliminates half the remaining functions. So, at most $\log_2(N)$ mistakes, and average loss is at most $\frac{\log_2(N)}{T}$.

A BETTER ALGORITHM

“Halving” algorithm

In round t ,

$$\hat{y}_t := \arg \max_{y \in \{0,1\}} |\{\text{perfect-so-far } f : f(x_t) = y\}|$$

(i.e., go with majority's prediction).

What is the average loss after T rounds?

Answer: every mistake eliminates half the remaining functions. So, at most $\log_2(N)$ mistakes, and average loss is at most $\frac{\log_2(N)}{T}$.

What if we drop realizability assumption?

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$									
y_t									
$L_t(f_0)$									
$L_t(f_1)$									
leader									

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1								
y_t									
$L_t(f_0)$									
$L_t(f_1)$									
leader									

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1								
y_t	$2/3$								
$L_t(f_0)$									
$L_t(f_1)$									
leader									

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1								
y_t	$2/3$								
$L_t(f_0)$	$2/3$								
$L_t(f_1)$	$1/3$								
leader	f_1								

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1							
y_t	$2/3$								
$L_t(f_0)$	$2/3$								
$L_t(f_1)$	$1/3$								
leader	f_1								

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1							
y_t	$2/3$	0							
$L_t(f_0)$	$2/3$								
$L_t(f_1)$	$1/3$								
leader	f_1								

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1							
y_t	$2/3$	0							
$L_t(f_0)$	$2/3$	$2/3$							
$L_t(f_1)$	$1/3$	$4/3$							
leader	f_1	f_0							

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0						
y_t	$2/3$	0							
$L_t(f_0)$	$2/3$	$2/3$							
$L_t(f_1)$	$1/3$	$4/3$							
leader	f_1	f_0							

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0						
y_t	$2/3$	0	1						
$L_t(f_0)$	$2/3$	$2/3$							
$L_t(f_1)$	$1/3$	$4/3$							
leader	f_1	f_0							

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0						
y_t	$2/3$	0	1						
$L_t(f_0)$	$2/3$	$2/3$	$5/3$						
$L_t(f_1)$	$1/3$	$4/3$	$4/3$						
leader	f_1	f_0	f_1						

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1					
y_t	$2/3$	0	1						
$L_t(f_0)$	$2/3$	$2/3$	$5/3$						
$L_t(f_1)$	$1/3$	$4/3$	$4/3$						
leader	f_1	f_0	f_1						

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1					
y_t	$2/3$	0	1	0					
$L_t(f_0)$	$2/3$	$2/3$	$5/3$						
$L_t(f_1)$	$1/3$	$4/3$	$4/3$						
leader	f_1	f_0	f_1						

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1					
y_t	$2/3$	0	1	0					
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$					
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$					
leader	f_1	f_0	f_1	f_0					

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0				
y_t	$2/3$	0	1	0					
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$					
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$					
leader	f_1	f_0	f_1	f_0					

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0				
y_t	$2/3$	0	1	0	1				
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$					
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$					
leader	f_1	f_0	f_1	f_0					

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0				
y_t	$2/3$	0	1	0	1				
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$				
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$				
leader	f_1	f_0	f_1	f_0	f_1				

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0	1	0	1	0
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0	1	0	1	0
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

Regret of follow-the-leader is constant (i.e. not going to 0).

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0	1	0	1	0
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

Regret of follow-the-leader is constant (i.e. not going to 0).

Can any learner do much better?

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0	1	0	1	0
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

Regret of follow-the-leader is constant (i.e. not going to 0).

Can any learner do much better?

Cover's impossibility theorem: No *deterministic* learner can ensure regret $< \frac{1}{2}$

EXAMPLE: FAILURE OF FOLLOW-THE-LEADER

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$\hat{y}_t$	1	1	0	1	0	1	0	1	0
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Follow-the-leader: go with the function $f \in \mathcal{F}$ that is best-so-far.

Regret of follow-the-leader is constant (i.e. not going to 0).

Can any learner do much better?

Cover's impossibility theorem: No *deterministic* learner can ensure regret $< \frac{1}{2}$

However:

If the learner can be randomized, then can make regret $\rightarrow 0$ as $T \rightarrow \infty$.

USING RANDOMNESS

Assume $N := |\mathcal{F}|$ is finite, and $\ell(y, \hat{y}) \in [0, 1]$.

Let $L_t(f) := \sum_{\tau=1}^t \ell(y_\tau, f(x_\tau))$ be cumulative loss of f after t rounds.

USING RANDOMNESS

Assume $N := |\mathcal{F}|$ is finite, and $\ell(y, \hat{y}) \in [0, 1]$.

Let $L_t(f) := \sum_{\tau=1}^t \ell(y_\tau, f(x_\tau))$ be cumulative loss of f after t rounds.

Randomized Weighted Majority (Littlestone and Warmuth, 1994)

In round t , randomly pick f_t from distribution w_t on $\mathcal{F}$ given by

$$w_t(f) \propto \exp(-\eta \cdot L_{t-1}(f));$$

predict $\hat{y}_t := f_t(x_t)$.

USING RANDOMNESS

Assume $N := |\mathcal{F}|$ is finite, and $\ell(y, \hat{y}) \in [0, 1]$.

Let $L_t(f) := \sum_{\tau=1}^t \ell(y_\tau, f(x_\tau))$ be cumulative loss of f after t rounds.

Randomized Weighted Majority (Littlestone and Warmuth, 1994)

In round t , randomly pick f_t from distribution w_t on $\mathcal{F}$ given by

$$w_t(f) \propto \exp(-\eta \cdot L_{t-1}(f));$$

predict $\hat{y}_t := f_t(x_t)$.

Expected regret after T rounds:

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t)) \right] \leq \frac{1}{\eta} \cdot \frac{\ln(N)}{T} + \eta \cdot \frac{1}{2}.$$

USING RANDOMNESS

Assume $N := |\mathcal{F}|$ is finite, and $\ell(y, \hat{y}) \in [0, 1]$.

Let $L_t(f) := \sum_{\tau=1}^t \ell(y_\tau, f(x_\tau))$ be cumulative loss of f after t rounds.

Randomized Weighted Majority (Littlestone and Warmuth, 1994)

In round t , randomly pick f_t from distribution w_t on $\mathcal{F}$ given by

$$w_t(f) \propto \exp(-\eta \cdot L_{t-1}(f));$$

predict $\hat{y}_t := f_t(x_t)$.

Expected regret after T rounds:

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t)) \right] \leq \frac{1}{\eta} \cdot \frac{\ln(N)}{T} + \eta \cdot \frac{1}{2}.$$

Optimize bound with $\eta := \sqrt{\frac{2 \ln(N)}{T}}$ to get regret bound $\sqrt{\frac{2 \ln(N)}{T}}$.

USING RANDOMNESS

Assume $N := |\mathcal{F}|$ is finite, and $\ell(y, \hat{y}) \in [0, 1]$.

Let $L_t(f) := \sum_{\tau=1}^t \ell(y_\tau, f(x_\tau))$ be cumulative loss of f after t rounds.

Randomized Weighted Majority (Littlestone and Warmuth, 1994)

In round t , randomly pick f_t from distribution w_t on $\mathcal{F}$ given by

$$w_t(f) \propto \exp(-\eta \cdot L_{t-1}(f));$$

predict $\hat{y}_t := f_t(x_t)$.

Expected regret after T rounds:

$$\mathbb{E} \left[\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(x_t)) \right] \leq \frac{1}{\eta} \cdot \frac{\ln(N)}{T} + \eta \cdot \frac{1}{2}.$$

Optimize bound with $\eta := \sqrt{\frac{2 \ln(N)}{T}}$ to get regret bound $\sqrt{\frac{2 \ln(N)}{T}}$.

Could also use $\eta_t := \sqrt{\frac{2 \ln(N)}{t}}$ in round t .

EXAMPLE: RWM

$\mathcal{Y} = [0, 1]$, $\ell(y, \hat{y}) := |\hat{y} - y|$ (absolute loss),

t	1	2	3	4	5	6	7	8	9
$f_0(x_t)$	0	0	0	0	0	0	0	0	0
$f_1(x_t)$	1	1	1	1	1	1	1	1	1
$w_t(f_1)$	.500	.569	.389	.549	.413	.540	.426	.535	.435
y_t	$2/3$	0	1	0	1	0	1	0	1
$L_t(f_0)$	$2/3$	$2/3$	$5/3$	$5/3$	$8/3$	$8/3$	$11/3$	$11/3$	$14/3$
$L_t(f_1)$	$1/3$	$4/3$	$4/3$	$7/3$	$7/3$	$10/3$	$10/3$	$13/3$	$13/3$
leader	f_1	f_0	f_1	f_0	f_1	f_0	f_1	f_0	f_1

Expected regret of RWM is $O(1/\sqrt{T})$.

WHY RWM WORKS

WHY RWM WORKS

- ▶ Randomization lets us precisely “hedge” over functions in $\mathcal{F}$.

(A generalization of RWM called “Hedge” is due to Freund and Schapire; proposed in same paper as AdaBoost.)

WHY RWM WORKS

- ▶ Randomization lets us precisely “hedge” over functions in $\mathcal{F}$.

(A generalization of RWM called “Hedge” is due to Freund and Schapire; proposed in same paper as AdaBoost.)

- ▶ Minimizer of cumulative loss L_{t-1} can change in every round.
But the distribution given by

$$w_t(f) \propto \exp(-\eta \cdot L_{t-1}(f))$$

(as in RWM) changes relatively slowly when η is small.

APPLICATION TO LEARNING DECISION TREES

Usual approach

Split training data S into two parts, S' and S'' :

- ▶ Use first part S' to **grow tree $\mathcal{T}$ until all leaves are pure.**
- ▶ Use second part S'' to **choose a good pruning of tree $\mathcal{T}$.**

APPLICATION TO LEARNING DECISION TREES

Usual approach

Split training data S into two parts, S' and S'' :

- ▶ Use first part S' to **grow tree $\mathcal{T}$ until all leaves are pure.**
- ▶ Use second part S'' to **choose a good pruning of tree $\mathcal{T}$.**

Alternative approach (Helmbold and Schapire, 1995)

Split training data S into two parts, S' and S'' :

- ▶ Use first part S' to **grow tree $\mathcal{T}$ until all leaves are pure.**
- ▶ Use second part S'' to **learn weighting over all prunings of tree $\mathcal{T}$:**
Let $\mathcal{F}$ = all prunings of $\mathcal{T}$, and run RWM using S'' .

APPLICATION TO LEARNING DECISION TREES

Usual approach

Split training data S into two parts, S' and S'' :

- ▶ Use first part S' to **grow tree $\mathcal{T}$ until all leaves are pure.**
- ▶ Use second part S'' to **choose a good pruning of tree $\mathcal{T}$.**

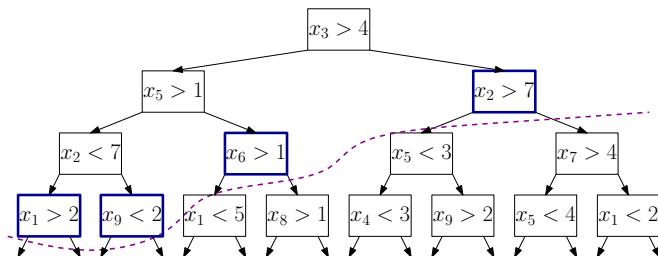
Alternative approach (Helmbold and Schapire, 1995)

Split training data S into two parts, S' and S'' :

- ▶ Use first part S' to **grow tree $\mathcal{T}$ until all leaves are pure.**
- ▶ Use second part S'' to **learn weighting over all prunings of tree $\mathcal{T}$:**
Let $\mathcal{F}$ = all prunings of $\mathcal{T}$, and run RWM using S'' .

Problem: $|\mathcal{F}|$ is exponentially large!

DYNAMIC PROGRAMMING



- ▶ Each $f \in \mathcal{F}$ corresponds to a pruning of the tree $\mathcal{T}$.
- ▶ $L_t(f)$ depends on prediction error statistics at leaves of f ; weight of f is just $\exp(-\eta L_t(f))$.
- ▶ Maintain error statistics for all nodes in $\mathcal{T}$.
- ▶ Can use dynamic programming to combine weights of all $f \in \mathcal{F}$ to make weighted prediction.
- ▶ See (Helmhold & Schapire, 1995) for details.

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes
- ▶ For infinite function classes, bounded VC dimension is *not enough*, unlike the batch case

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes
- ▶ For infinite function classes, bounded VC dimension is *not enough*, unlike the batch case
- ▶ E.g. the class of thresholds in $\mathbb{R}$ has VC-dimension 2, but no online learning algorithm can achieve vanishing regret

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes
- ▶ For infinite function classes, bounded VC dimension is *not enough*, unlike the batch case
- ▶ E.g. the class of thresholds in $\mathbb{R}$ has VC-dimension 2, but no online learning algorithm can achieve vanishing regret
- ▶ Different combinatorial parameter called **Littlestone dimension** characterizes online learnability

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes
- ▶ For infinite function classes, bounded VC dimension is *not enough*, unlike the batch case
- ▶ E.g. the class of thresholds in $\mathbb{R}$ has VC-dimension 2, but no online learning algorithm can achieve vanishing regret
- ▶ Different combinatorial parameter called **Littlestone dimension** characterizes online learnability
- ▶ For $\mathcal{F}$ with bounded Littlestone dimension, generic online learning algorithm: RWM on carefully constructed “experts”

INFINITE FUNCTION CLASSES

- ▶ RWM only works for finite function classes
- ▶ For infinite function classes, bounded VC dimension is *not enough*, unlike the batch case
- ▶ E.g. the class of thresholds in $\mathbb{R}$ has VC-dimension 2, but no online learning algorithm can achieve vanishing regret
- ▶ Different combinatorial parameter called **Littlestone dimension** characterizes online learnability
- ▶ For $\mathcal{F}$ with bounded Littlestone dimension, generic online learning algorithm: RWM on carefully constructed “experts”
- ▶ Caveat: generic online learning algorithm is *inefficient*

ONLINE CONVEX OPTIMIZATION

- ▶ Suppose now that $\mathcal{F}$ is a parametric function class:
 - ▶ $\mathbf{W} \subseteq \mathbb{R}^d$ is a **convex set** of allowable parameters.
 - ▶ For every $\mathbf{w} \in \mathbf{W}$, $f_{\mathbf{w}} : \mathcal{X} \rightarrow \mathcal{Y}$ is predictor parameterized by $\mathbf{w}$
 - ▶ $\mathcal{F} = \{f_{\mathbf{w}} \mid \mathbf{w} \in \mathbf{W}\}$

ONLINE CONVEX OPTIMIZATION

- ▶ Suppose now that $\mathcal{F}$ is a parametric function class:
 - ▶ $\mathbf{W} \subseteq \mathbb{R}^d$ is a **convex set** of allowable parameters.
 - ▶ For every $\mathbf{w} \in \mathbf{W}$, $f_{\mathbf{w}} : \mathcal{X} \rightarrow \mathcal{Y}$ is predictor parameterized by $\mathbf{w}$
 - ▶ $\mathcal{F} = \{f_{\mathbf{w}} \mid \mathbf{w} \in \mathbf{W}\}$
- ▶ Assume, further, that for any example $(x, y) \in \mathcal{X} \times \mathcal{Y}$, the mapping $\mathbf{w} \mapsto \ell(y, f_{\mathbf{w}}(x))$ is **convex** in $\mathbf{w}$.

ONLINE CONVEX OPTIMIZATION

- ▶ Suppose now that $\mathcal{F}$ is a parametric function class:
 - ▶ $\mathbf{W} \subseteq \mathbb{R}^d$ is a **convex set** of allowable parameters.
 - ▶ For every $\mathbf{w} \in \mathbf{W}$, $f_{\mathbf{w}} : \mathcal{X} \rightarrow \mathcal{Y}$ is predictor parameterized by $\mathbf{w}$
 - ▶ $\mathcal{F} = \{f_{\mathbf{w}} \mid \mathbf{w} \in \mathbf{W}\}$
- ▶ Assume, further, that for any example $(\mathbf{x}, y) \in \mathcal{X} \times \mathcal{Y}$, the mapping $\mathbf{w} \mapsto \ell(y, f_{\mathbf{w}}(\mathbf{x}))$ is **convex** in $\mathbf{w}$.
- ▶ Example: linear regression setting
 - ▶ $\mathcal{X} = \mathbb{R}^d$
 - ▶ $\mathcal{Y} = \mathbb{R}$
 - ▶ $\mathbf{W} = \{\mathbf{w} \in \mathbb{R}^d \mid \|\mathbf{w}\|_2 \leq R\}$, a convex set
 - ▶ $f_{\mathbf{w}}(\mathbf{x}) := \langle \mathbf{w}, \mathbf{x} \rangle$
 - ▶ Square loss: $\ell(y, \hat{y}) = (\hat{y} - y)^2$
 - ▶ For any $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathbb{R}$, $\mathbf{w} \mapsto (\langle \mathbf{w}, \mathbf{x} \rangle - y)^2$ is *convex*

ONLINE CONVEX OPTIMIZATION

- ▶ Suppose now that $\mathcal{F}$ is a parametric function class:
 - ▶ $\mathbf{W} \subseteq \mathbb{R}^d$ is a **convex set** of allowable parameters.
 - ▶ For every $\mathbf{w} \in \mathbf{W}$, $f_{\mathbf{w}} : \mathcal{X} \rightarrow \mathcal{Y}$ is predictor parameterized by $\mathbf{w}$
 - ▶ $\mathcal{F} = \{f_{\mathbf{w}} \mid \mathbf{w} \in \mathbf{W}\}$
- ▶ Assume, further, that for any example $(\mathbf{x}, y) \in \mathcal{X} \times \mathcal{Y}$, the mapping $\mathbf{w} \mapsto \ell(y, f_{\mathbf{w}}(\mathbf{x}))$ is **convex** in $\mathbf{w}$.
- ▶ Example: linear regression setting
 - ▶ $\mathcal{X} = \mathbb{R}^d$
 - ▶ $\mathcal{Y} = \mathbb{R}$
 - ▶ $\mathbf{W} = \{\mathbf{w} \in \mathbb{R}^d \mid \|\mathbf{w}\|_2 \leq R\}$, a convex set
 - ▶ $f_{\mathbf{w}}(\mathbf{x}) := \langle \mathbf{w}, \mathbf{x} \rangle$
 - ▶ Square loss: $\ell(y, \hat{y}) = (\hat{y} - y)^2$
 - ▶ For any $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathbb{R}$, $\mathbf{w} \mapsto (\langle \mathbf{w}, \mathbf{x} \rangle - y)^2$ is *convex*
- ▶ This setting is called *online convex optimization*, and admits *efficient* algorithms

ONLINE GRADIENT DESCENT

Algorithm for online convex optimization: an online form of gradient descent!

ONLINE GRADIENT DESCENT

Algorithm for online convex optimization: an online form of gradient descent!

Seen before: online perceptron and stochastic gradient descent

ONLINE GRADIENT DESCENT

Algorithm for online convex optimization: an online form of gradient descent!

Seen before: online perceptron and stochastic gradient descent

Online Gradient Descent (OGD) (Zinkevich, 2003)

- 1: Choose learning rate η .
- 2: Initialize weights $\mathbf{w}_1 \in \mathbf{W}$ arbitrarily.
- 3: **for** $t = 1, 2, \dots, T$ **do**
- 4: Obtain $\mathbf{x}_t \in \mathbb{R}^d$
- 5: Predict $\hat{y}_t = f_{\mathbf{w}_t}(\mathbf{x}_t)$
- 6: Observe y_t , and incur loss $\ell(y_t, \hat{y}_t)$.
- 7: Compute the gradient $\nabla_{\mathbf{w}} \ell(y_t, f_{\mathbf{w}}(\mathbf{x}_t))$ at $\mathbf{w}_t$, call it $\mathbf{g}_t$.
- 8: Update:

$$\mathbf{w}_{t+1} = \Pi_{\mathbf{W}}(\mathbf{w}_t - \eta \mathbf{g}_t)$$

- 9: **end for**

Here, $\Pi_{\mathbf{W}}(\mathbf{w}') := \arg \min_{\mathbf{w} \in \mathbf{W}} \|\mathbf{w} - \mathbf{w}'\|_2^2$. Also known as *projection* on $\mathbf{W}$.

ONLINE GRADIENT DESCENT

Online Gradient Descent (OGD)

- 1: Choose learning rate η .
- 2: Initialize weights $\mathbf{w}_1 \in \mathbf{W}$ arbitrarily.
- 3: **for** $t = 1, 2, \dots, T$ **do**
- 4: Obtain $\mathbf{x}_t \in \mathbb{R}^d$
- 5: Predict $\hat{y}_t = f_{\mathbf{w}_t}(\mathbf{x}_t)$
- 6: Observe y_t , and incur loss $\ell(y_t, \hat{y}_t)$.
- 7: Compute the gradient $\nabla_{\mathbf{w}} \ell(y_t, f_{\mathbf{w}}(\mathbf{x}_t))$ at $\mathbf{w}_t$, call it $\mathbf{g}_t$.
- 8: Update:

$$\mathbf{w}_{t+1} = \Pi_{\mathbf{W}}(\mathbf{w}_t - \eta \mathbf{g}_t)$$

- 9: **end for**

Regret of OGD after T rounds: Suppose for any $\mathbf{w}, \mathbf{w}' \in \mathbf{W}$, we have $\|\mathbf{w} - \mathbf{w}'\|_2 \leq D$ and $\|\mathbf{g}_t\|_2 \leq G$ for all t . Then

$$\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(\mathbf{x}_t)) \leq \frac{1}{\eta} \cdot \frac{D^2}{2T} + \eta \cdot \frac{G^2}{2}.$$

ONLINE GRADIENT DESCENT

Online Gradient Descent (OGD)

- 1: Choose learning rate η .
- 2: Initialize weights $\mathbf{w}_1 \in \mathbf{W}$ arbitrarily.
- 3: **for** $t = 1, 2, \dots, T$ **do**
- 4: Obtain $\mathbf{x}_t \in \mathbb{R}^d$
- 5: Predict $\hat{y}_t = f_{\mathbf{w}_t}(\mathbf{x}_t)$
- 6: Observe y_t , and incur loss $\ell(y_t, \hat{y}_t)$.
- 7: Compute the gradient $\nabla_{\mathbf{w}} \ell(y_t, f_{\mathbf{w}}(\mathbf{x}_t))$ at $\mathbf{w}_t$, call it $\mathbf{g}_t$.
- 8: Update:

$$\mathbf{w}_{t+1} = \Pi_{\mathbf{W}}(\mathbf{w}_t - \eta \mathbf{g}_t)$$

- 9: **end for**

Regret of OGD after T rounds: Suppose for any $\mathbf{w}, \mathbf{w}' \in \mathbf{W}$, we have $\|\mathbf{w} - \mathbf{w}'\|_2 \leq D$ and $\|\mathbf{g}_t\|_2 \leq G$ for all t . Then

$$\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(\mathbf{x}_t)) \leq \frac{1}{\eta} \cdot \frac{D^2}{2T} + \eta \cdot \frac{G^2}{2}.$$

Optimize bound with $\eta := \sqrt{\frac{D}{GT}}$ to get regret bound $\sqrt{\frac{GD}{4T}}$.

ONLINE GRADIENT DESCENT

Online Gradient Descent (OGD)

- 1: Choose learning rate η .
- 2: Initialize weights $\mathbf{w}_1 \in \mathbf{W}$ arbitrarily.
- 3: **for** $t = 1, 2, \dots, T$ **do**
- 4: Obtain $\mathbf{x}_t \in \mathbb{R}^d$
- 5: Predict $\hat{y}_t = f_{\mathbf{w}_t}(\mathbf{x}_t)$
- 6: Observe y_t , and incur loss $\ell(y_t, \hat{y}_t)$.
- 7: Compute the gradient $\nabla_{\mathbf{w}} \ell(y_t, f_{\mathbf{w}}(\mathbf{x}_t))$ at $\mathbf{w}_t$, call it $\mathbf{g}_t$.
- 8: Update:

$$\mathbf{w}_{t+1} = \Pi_{\mathbf{W}}(\mathbf{w}_t - \eta \mathbf{g}_t)$$

- 9: **end for**

Regret of OGD after T rounds: Suppose for any $\mathbf{w}, \mathbf{w}' \in \mathbf{W}$, we have $\|\mathbf{w} - \mathbf{w}'\|_2 \leq D$ and $\|\mathbf{g}_t\|_2 \leq G$ for all t . Then

$$\frac{1}{T} \sum_{t=1}^T \ell(y_t, \hat{y}_t) - \min_{f \in \mathcal{F}} \frac{1}{T} \sum_{t=1}^T \ell(y_t, f(\mathbf{x}_t)) \leq \frac{1}{\eta} \cdot \frac{D^2}{2T} + \eta \cdot \frac{G^2}{2}.$$

Optimize bound with $\eta := \sqrt{\frac{D}{GT}}$ to get regret bound $\sqrt{\frac{GD}{4T}}$.

Could also use $\eta_t := \sqrt{\frac{D}{Gt}}$ in round t .

APPLICATION TO “BIG DATA” LEARNING

- ▶ Online learning algorithms can be easily converted to batch learning algorithms
- ▶ **Online-to-batch conversion:** output a random predictor

APPLICATION TO “BIG DATA” LEARNING

- ▶ Online learning algorithms can be easily converted to batch learning algorithms
- ▶ **Online-to-batch conversion:** output a random predictor
- ▶ Can have high variance, so some form of averaging is useful (if possible)
- ▶ E.g. in online convex optimization settings, output average parameter vector.

APPLICATION TO “BIG DATA” LEARNING

- ▶ Online learning algorithms can be easily converted to batch learning algorithms
- ▶ **Online-to-batch conversion:** output a random predictor
- ▶ Can have high variance, so some form of averaging is useful (if possible)
- ▶ E.g. in online convex optimization settings, output average parameter vector.
- ▶ Utility: can process “big data” by learning example-by-example
- ▶ In fact, SGD analysis directly obtained from online learning techniques.

MORE ON ONLINE LEARNING

Many problems and algorithms in online learning:

- ▶ Partial feedback (“multi-armed bandit”)
- ▶ Bayesian inference
- ▶ Tracking / filtering
- ▶ ...

Many connections between online learning and other subjects:

- ▶ Game theory, calibration, boosting
- ▶ Stochastic optimization
- ▶ Compression, coding theory
- ▶ ...