



Lecture 5. Online Convex Optimization

Advanced Optimization (Fall 2025)

Peng Zhao

zhaop@lamda.nju.edu.cn

Nanjing University

Outline

- Online Optimization
- Online Convex Optimization
- Online-to-Batch Conversion

Part 1. Online Optimization

- Fixed Optimization
- Interactive Optimization
- Online Optimization: Problem and Measure

Optimization

- Originally, we focus on the optimization problem: $\min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x})$

Iterative algorithm (like GD) generates a sequence of iterates $\mathbf{x}_1, \dots, \mathbf{x}_T$, aiming to ensure that $F(\mathbf{x}_T) - F(\mathbf{x}^*)$ is small with respect to T

for simplicity, consider last-iterate convergence

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

- Algorithm outputs a decision $\mathbf{x}_t \in \mathcal{X}$.
- Algorithm has information about $F(\mathbf{x}_t)$ and $\nabla F(\mathbf{x}_t)$.

⇒ Note that the objective function remains *fixed* throughout the entire optimization process as the iterative algorithm progresses.

Interactive Optimization

- However, in many ML applications, the optimization objective *does not stay fixed*.

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

- Algorithm outputs a decision $\mathbf{x}_t \in \mathcal{X}$.
- A new loss function $f_t : \mathcal{X} \rightarrow \mathbb{R}$ is revealed.
- Algorithm suffers the loss $f_t(\mathbf{x}_t)$.

Optimization is no longer performed against a single, fixed target; instead, it involves interactions between an algorithm and *an evolving objective*.

Example 1: Large-Scale ERM

- Consider the model training. Our goal is to minimize ERM

$$F(\mathbf{x}) = \frac{1}{M} \sum_{i=1}^M \ell(\mathbf{x}; z_i)$$

- Big data: facing millions of samples (M is very large).
- Computing full gradient is almost impossible (like due to limited memory of GPU facilities).

Stochastic Optimization as Interaction

Stochastic optimization using *mini-batch* (by SGD/Adam or others)

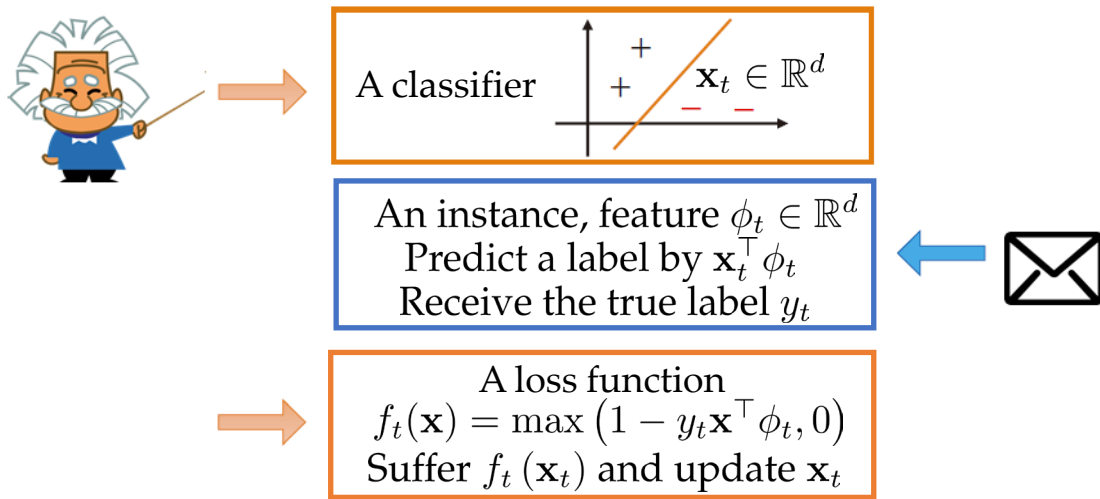
$$f_t(\mathbf{x}) = \frac{1}{|S_t|} \sum_{z_i \in S_t} \ell(\mathbf{x}; z_i)$$

Note that S_t is still sampled from a fixed distribution (over all M samples)

- The optimization objective f_t varies across iterations (due to changing mini-batches).
- This renders the optimization process *interactive* — the algorithm continuously adapts in response to a new batch of samples.

Example 2: Learning with Data Streams

- In applications, data may arrive in a form of stream, which requires the model to make *sequential predictions or even decisions*.
- For example: spam filtering



Sequential Predictions as Interaction

- We can formulate the sequential prediction problem as follows:

At each step t :

- The learner outputs a prediction $\mathbf{x}_t$.
- Then a new sample z_t is revealed.
- The instantaneous loss is $f_t(\mathbf{x}_t) = \ell(\mathbf{x}_t; z_t)$.

Note that the loss function f_t is sampled from (unknown) distribution $\mathcal{D}_\star$ or $\mathcal{D}_t$

- If data are stationary, we can view f_t as sampled from a fixed distribution $\mathcal{D}_\star$.
- Alternatively, if having non-stationarity, distribution $\mathcal{D}_t$ may vary with t .

Online Optimization: distribution-free

- Online optimization: an even general framework without distributional assumption.
- The function f_t may be arbitrary, or even chosen adversarially
- The general interactive procedures of online optimization:

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

- Learner gives a decision $\mathbf{x}_t \in \mathcal{X}$.
- Learner observes $f_t : \mathcal{X} \rightarrow \mathbb{R}$ and suffers the loss $f_t(\mathbf{x}_t)$.

A Game-theoretic Language

- Online (Interactive) optimization as a *repeated game* between
 - **Player**: essentially the learner, or you can think as the “learning model”
 - **Environments**: an abstraction of all factors evaluating the model.

At each round $t = 1, 2, \dots$

- Player first picks a model $\mathbf{x}_t \in \mathcal{X}$.
- Simultaneously environments pick an online function $f_t : \mathcal{X} \rightarrow \mathbb{R}$.
- Player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.

Fixed vs Interactive Optimization

- Fixed Optimization (Closed-world Learning)
 - The training data are all available on hand
 - The objective is fixed with no changes.
- Online Optimization (Open-ended Learning)
 - This may be due to sampling issue, so connected to stochastic optimization
 - Or since data are in the form of stream, so it is crucial for continual learning
 - Even more complicated: decision this round may also influent the environment, decision-theoretic RL/control



Performance Measure

- The best in hindsight

After T rounds, suppose we look back and imagine that we *did* know all $\{f_t\}_{t=1}^T$.

Then the best single decision in hindsight would be:

$$\mathbf{x}^* \in \arg \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T f_t(\mathbf{x})$$

which we obviously could *not* know in advance.

Regret:
$$\text{REG}_T = \sum_{t=1}^T f_t(\mathbf{x}_t) - \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T f_t(\mathbf{x})$$

*benchmark performance with the
offline model (optimal in hindsight)*

*“If I had known all the outcomes ahead of time, I would have
chosen differently — and I **regret** the extra loss I’ve accumulated.”*

Performance Measure

- In online learning, we define *regret* as measure for $\{\mathbf{x}_t\}_{t=1}^T$:

$$\text{REG}_T = \sum_{t=1}^T f_t(\mathbf{x}_t) - \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T f_t(\mathbf{x})$$

benchmark performance with the offline model (optimal in hindsight)

We hope the regret be sub-linear dependence with T

$$\frac{\text{REG}_T}{T} \rightarrow 0 \text{ as } T \rightarrow \infty$$

Hannan Consistency

ALT'16

Hannan Consistency in On-Line Learning
in Case of Unbounded Losses Under Partial
Monitoring^{*,**}

Chamy Allenberg¹, Peter Auer², László Györfi³, and György Ottucsák³

worst-case-dynamic regret $\text{REG}_T(\{\mathbf{x}_t^*\}) = \sum_{t=1}^T f_t(\mathbf{x}_t) - \sum_{t=1}^T f_t(\mathbf{x}_t^*)$ *suffer from overfitting issue*

general dynamic regret $\text{REG}_T(\{\mathbf{u}_t\}) = \sum_{t=1}^T f_t(\mathbf{x}_t) - \sum_{t=1}^T f_t(\mathbf{u}_t)$ *“online ensemble” framework*

Part 2. Online Convex Optimization

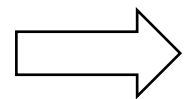
- Online Convex Optimization
- Problem Space
- Online Gradient Descent
- Lower Bound

Is Online Optimization (provably) solvable?

- In general, the online optimization is *too hard to solve*.

At each round $t = 1, 2, \dots$

- Player first picks a model $\mathbf{x}_t \in \mathcal{X}$.
- Simultaneously environments pick an online function $f_t : \mathcal{X} \rightarrow \mathbb{R}$.
- Player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.



A Trackable Case: *Online Convex Optimization*

*requiring feasible domain and online functions to be **convex***

Online Convex Optimization

- Requirements:
 - (1) feasible domain is a convex set
 - (2) online functions are convex

At each round $t = 1, 2, \dots$

- Player first picks a model $\mathbf{x}_t$ from a convex set $\mathcal{X} \subseteq \mathbb{R}^d$.
- Environments pick an online convex function $f_t : \mathcal{X} \rightarrow \mathbb{R}$.
- Player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.

Online Convex Optimization: Hardness

Clearly, curvature information influences:

- Convex vs Strongly convex?
- Lipschitz vs Smooth?

There are other issues due to the *interaction* nature:

- **Feedback:** How much information can the learner access from environments?
- **Environments:** How powerful is the environment?

Given that it can choose the loss function!

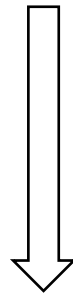
Hardness: Different Feedback

At each round $t = 1, 2, \dots$

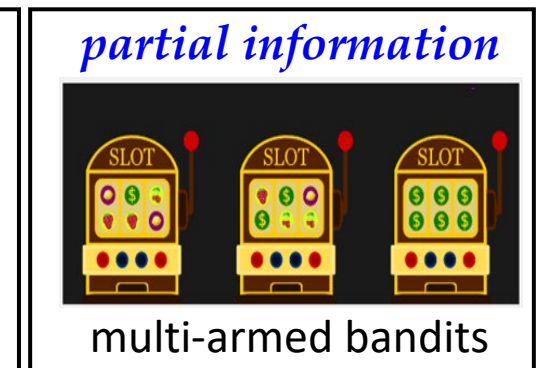
- Player first picks a model $\mathbf{x}_t$ from a convex set $\mathcal{X} \subseteq \mathbb{R}^d$.
- Environments pick an online convex function $f_t : \mathcal{X} \rightarrow \mathbb{R}$.
- Player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.

on the feedback information:

- **full information**: observe entire f_t (or at least gradient $\nabla f_t(\mathbf{x}_t)$)
- **partial information (bandits)**: observe function value $f_t(\mathbf{x}_t)$ only



less information



Hardness: Different Environments

At each round $t = 1, 2, \dots$

- Player first picks a model $\mathbf{x}_t$ from a convex set $\mathcal{X} \subseteq \mathbb{R}^d$.
- **Environments** pick an online convex function $f_t : \mathcal{X} \rightarrow \mathbb{R}$.
- Player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.

on the difficulty of environments:

- stochastic setting

- adversarial setting { oblivious
adaptive
(non-oblivious)

*less restricted
but harder*

oblivious adversary



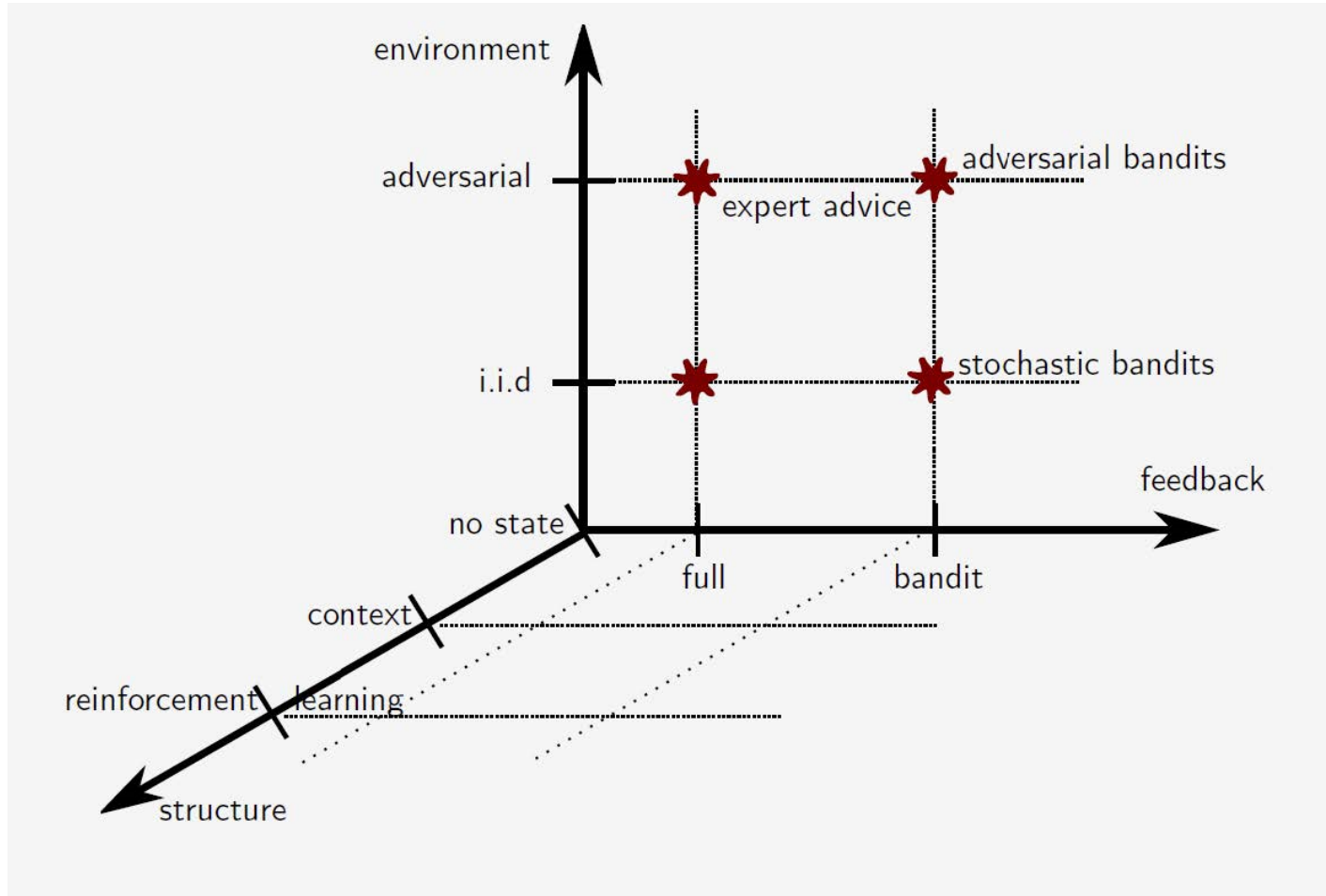
examination

adaptive adversary



interview

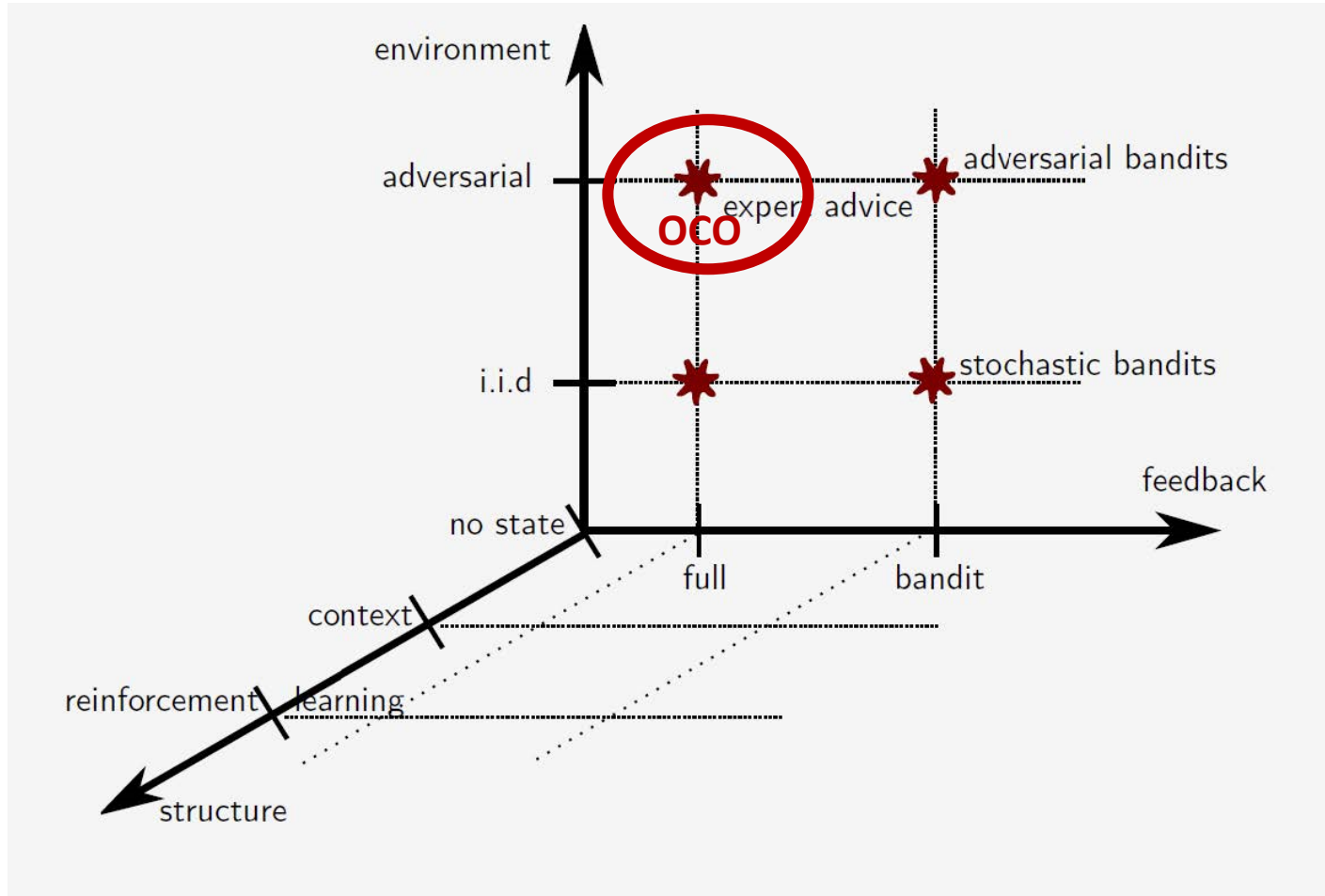
Online Learning: Problem Space



- Full-information setting:
 - Online Convex Optimization
 - Prediction with Expert Advice
 - ...
- Partial-information setting:
 - Multi-Armed Bandits
 - Linear Bandits
 - Parametric Bandits
 - Bandit Convex Optimization
 - ...

Yevgeny Seldin. The Space of Online Learning Problems, ECML-PKDD, Porto, Portugal, 2015.

Online Learning: Problem Space



- Full-information setting:
 - Online Convex Optimization
 - Prediction with Expert Advice
 - ...
- Partial-information setting:
 - Multi-Armed Bandits
 - Linear Bandits
 - Parametric Bandits
 - Bandit Convex Optimization
 - ...

Yevgeny Seldin. The Space of Online Learning Problems, ECML-PKDD, Porto, Portugal, 2015.

OCO: Convex Functions

Definition 2 (Convex Function). A function $f : \mathcal{X} \mapsto \mathbb{R}$ is convex if for any $\mathbf{x}, \mathbf{y} \in \mathcal{X}$, it holds $\forall \alpha \in [0, 1]$, $f((1 - \alpha)\mathbf{x} + \alpha\mathbf{y}) \leq (1 - \alpha)f(\mathbf{x}) + \alpha f(\mathbf{y})$.

Equivalently, if f is differentiable, we have that $\forall \mathbf{x}, \mathbf{y} \in \mathcal{X}$,

$$f(\mathbf{y}) \geq f(\mathbf{x}) + \nabla f(\mathbf{x})^\top (\mathbf{y} - \mathbf{x}).$$

The feasible set $\mathcal{X}$ is closed and convex in Euclidean space, and $f_1, \dots, f_T$ are convex functions.

Regret Minimization in OCO

- We focus on the G -Lipschitz functions

Assumption 1 (Bounded Gradient). The norm of the subgradients is upper bounded by G , i.e., $\|\nabla f_t(\mathbf{x})\| \leq G$ for all $\mathbf{x} \in \mathcal{X}$ and $t \in [T]$.

- The following is *domain boundedness*

Assumption 2 (Bounded Domain). The diameter of the feasible domain $\mathcal{X}$ is upper bounded by D , i.e., $\forall \mathbf{x}, \mathbf{y} \in \mathcal{X}, \|\mathbf{x} - \mathbf{y}\| \leq D$.

Essentially working on Lipschitz (online) optimization with a bounded feasible domain.

On the domain boundedness

Theorem 5.1. Let $\mathcal{V} \subset \mathbb{R}^d$ be any non-empty bounded closed convex subset. Let $D = \sup_{\mathbf{v}, \mathbf{w} \in \mathcal{V}} \|\mathbf{v} - \mathbf{w}\|_2$ be the diameter of $\mathcal{V}$. Let $\mathcal{A}$ be any (possibly randomized) algorithm for OLO on $\mathcal{V}$. Let T be any non-negative integer. Then, there exists a sequence of vectors $\mathbf{g}_1, \dots, \mathbf{g}_T$ with $\|\mathbf{g}_t\|_2 \leq L$ and $\mathbf{u} \in \mathcal{V}$ such that the regret of algorithm $\mathcal{A}$ satisfies

$$\text{Regret}_T(\mathbf{u}) = \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t \rangle - \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{u} \rangle \geq \frac{\sqrt{2}LD\sqrt{T}}{4}.$$

Theorem 5.6. For $T \in \mathbb{N}$ suppose that there is an online convex optimization algorithm that guarantees regret at most ϵ_T against the null competitor on any sequence of T linear and L -Lipschitz losses $\ell_t : \mathbb{R} \rightarrow \mathbb{R}$ for $t = 1, \dots, T$. Let $U > 0$ be such that $1 \leq W\left(\frac{\sqrt{T}UL}{5\epsilon_T}\right) \leq \frac{\sqrt{T}}{2}$, then there exists a sequence of $\mathbf{g}_t$ with $\|\mathbf{g}_t\|_2 \leq L$ and a competitor $\mathbf{u} \in \mathbb{R}^d$ with $\|\mathbf{u}\|_2 = U$, such that

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{u} \rangle \geq R_T(U) := UL\sqrt{T} \left(\sqrt{2W\left(\frac{\sqrt{T}UL}{5\epsilon_T}\right)} - 1 \right) - 2UL + \epsilon_T,$$

where $W : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ is the Lambert function.

Recall Lipschitz (offline) optimization

For G -Lipschitz optimization with $\min_{\mathbf{x}} f(\mathbf{x})$, we use GD to optimize it and obtain $\mathcal{O}\left(\frac{1}{\sqrt{T}}\right)$ optimal rate.

Optimal Result with Known T

Theorem 5. Under the same assumptions with Theorem 1, assume the feasible domain $\mathcal{X}$ is bounded and convex with a diameter $D > 0$, that is, $\|\mathbf{x} - \mathbf{y}\|_2 \leq D$ holds for any $\mathbf{x}, \mathbf{y} \in \mathcal{X}$. Let $\{\mathbf{x}_t\}_{t=1}^T$ be the sequence generated by GD with step size

$$\eta_t = \frac{D}{G\sqrt{T}}.$$

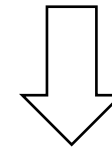
Then

$$f(\bar{\mathbf{x}}_T) - f^* \leq \frac{DG}{\sqrt{T}} = \mathcal{O}\left(\frac{1}{\sqrt{T}}\right),$$

where $\bar{\mathbf{x}}_T \triangleq \arg \min_{\{\mathbf{x}_t\}_{t=1}^T} f(\mathbf{x}_t)$ or $\bar{\mathbf{x}}_T \triangleq \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$.

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \nabla f(\mathbf{x}_t)]$$

Can we extend the GD idea to online optimization?



$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t)]$$

OCO: OGD Algorithm

Online Gradient Descent (OGD)

At each round $t = 1, 2, \dots$

1. the player first picks a model $\mathbf{x}_t \in \mathcal{X}$;
2. and simultaneously environments pick a **convex** online function $f_t : \mathcal{X} \rightarrow \mathbb{R}$;
3. the player suffers loss $f_t(\mathbf{x}_t)$, observes the information of f_t and update the model according to $\mathbf{x}_{t+1} = \Pi_{\mathcal{X}}[\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t)]$.

This belongs to the full-information setting, so player can access the gradient $\nabla f_t(\mathbf{x}_t)$.

Actually, only gradient is required, so it's also called **gradient-feedback OCO model**.

Regret Analysis of OGD

Theorem 3 (Regret bound for OGD). *Under Assumption 1 (G-Lipschitz) and Assumption 2 (D-bounded domain), online gradient descent (OGD) with step sizes $\eta_t = \frac{D}{G\sqrt{t}}$ for $t \in [T]$ guarantees:*

$$\text{REG}_T = \sum_{t=1}^T f_t(\mathbf{x}_t) - \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T f_t(\mathbf{x}) \leq \frac{3}{2}GD\sqrt{T} = \mathcal{O}(\sqrt{T}).$$

The First Gradient Descent Lemma

Lemma 1. Suppose that f_t is proper, closed and convex; the feasible domain $\mathcal{X}$ is nonempty, closed and convex. Let $\{\mathbf{x}_t\}_{t=1}^T$ be the sequence generated by the gradient descent method. Then for any $\mathbf{u} \in \mathcal{X}$ and $t \geq 0$,

$$\|\mathbf{x}_{t+1} - \mathbf{u}\|^2 \leq \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t(f_t(\mathbf{x}_t) - f_t(\mathbf{u})) + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2.$$

Proof:

$$\begin{aligned} \|\mathbf{x}_{t+1} - \mathbf{u}\|^2 &= \|\Pi_{\mathcal{X}}[\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t)] - \mathbf{u}\|^2 \quad (\text{GD}) \\ &\leq \|\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t) - \mathbf{u}\|^2 \quad (\text{Pythagoras Theorem}) \\ &= \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t \langle \nabla f_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2 \\ &\leq \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t(f_t(\mathbf{x}_t) - f_t(\mathbf{u})) + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2 \\ &\quad (\text{convexity: } f_t(\mathbf{x}_t) - f_t(\mathbf{u}) = f_t(\mathbf{x}_t) - f_t(\mathbf{u}) \leq \langle \nabla f_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle) \quad \square \end{aligned}$$

Proof for OGD Regret Bound

Proof: We use the first gradient descent lemma to analyze online gradient descent.

Lemma 1. Suppose that f_t is proper, closed and convex; the feasible domain $\mathcal{X}$ is nonempty, closed and convex. Let $\{\mathbf{x}_t\}_{t=1}^T$ be the sequence generated by the gradient descent method. Then for any $\mathbf{u} \in \mathcal{X}$ and $t \geq 0$,

$$\|\mathbf{x}_{t+1} - \mathbf{u}\|^2 \leq \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t(f_t(\mathbf{x}_t) - f_t(\mathbf{u})) + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2.$$

By Lemma 1 and the gradient boundedness, we have

$$2(f_t(\mathbf{x}_t) - f_t(\mathbf{u})) \leq \frac{\|\mathbf{x}_t - \mathbf{u}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{u}\|^2}{\eta_t} + \eta_t G^2$$

Proof for OGD Regret Bound

Proof: By setting $\eta_t = \frac{D}{G\sqrt{t}}$ (with $\frac{1}{\eta_0} := 0$), summing over T :

$$\begin{aligned} 2 \left(\sum_{t=1}^T f_t(\mathbf{x}_t) - f_t(\mathbf{u}) \right) &\leq \sum_{t=1}^T \frac{\|\mathbf{x}_t - \mathbf{u}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{u}\|^2}{\eta_t} + G^2 \sum_{t=1}^T \eta_t && \text{(GD lemma)} \\ &\leq \sum_{t=1}^T \|\mathbf{x}_t - \mathbf{u}\|^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} \right) + G^2 \sum_{t=1}^T \eta_t && (\|\mathbf{x}_{T+1} - \mathbf{u}\|^2 \geq 0) \\ &\leq D^2 \sum_{t=1}^T \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} \right) + G^2 \sum_{t=1}^T \eta_t \\ &\leq D^2 \frac{1}{\eta_T} + G^2 \sum_{t=1}^T \eta_t && (\eta_t = \frac{D}{G\sqrt{t}} \text{ and } \sum_{t=1}^T \frac{1}{\sqrt{t}} \leq 2\sqrt{T}) \\ &\leq 3DG\sqrt{T}. \end{aligned}$$

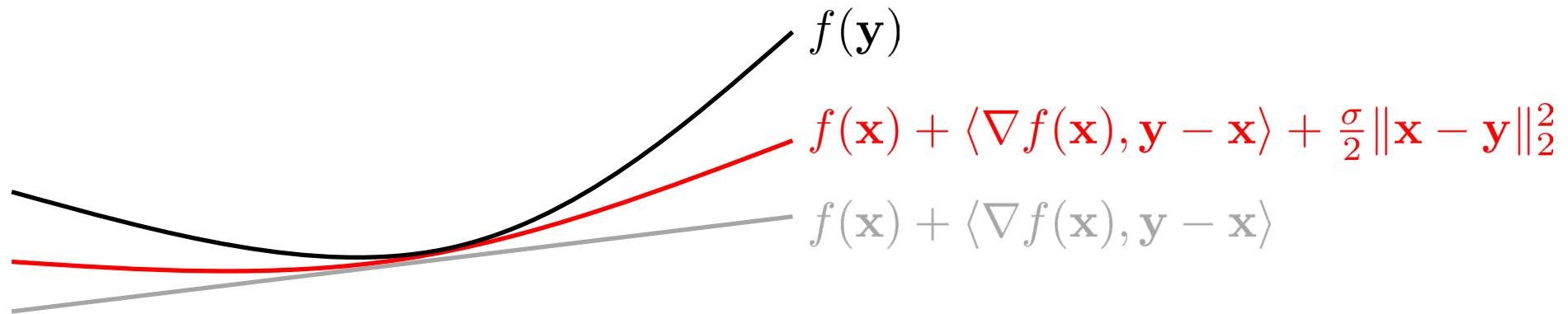
□

OCO: Strongly Convex Functions

Definition 3 (Strong Convexity). A function f is σ -strongly convex if, for any $\mathbf{x}, \mathbf{y} \in \text{dom } f$,

$$f(\mathbf{y}) \geq f(\mathbf{x}) + \nabla f(\mathbf{x})^\top (\mathbf{y} - \mathbf{x}) + \frac{\sigma}{2} \|\mathbf{y} - \mathbf{x}\|^2,$$

or equivalently, $\nabla^2 f(\mathbf{x}) \succeq \alpha I$.



Recall Lipschitz (offline) Optimization

For G -Lipschitz optimization with $\min_{\mathbf{x}} f(\mathbf{x})$, we use GD to optimize it and obtain $\mathcal{O}\left(\frac{1}{T}\right)$ optimal rate.

Strongly Convex and Lipschitz

Theorem 7. Under the same assumptions with Theorem 1, except that f is σ -strongly-convex. Let $\{\mathbf{x}_t\}_{t=1}^T$ be the sequence generated by GD with step size

$$\eta_t = \frac{2}{\sigma(t+1)}.$$

Then (i)

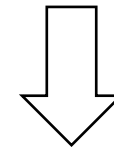
$$f(\bar{\mathbf{x}}_T) - f^* \leq \frac{2G^2}{\sigma(T+1)} = \mathcal{O}\left(\frac{1}{T}\right),$$

where $\bar{\mathbf{x}}_T \triangleq \arg \min_{\{\mathbf{x}_t\}_{t=1}^T} f(\mathbf{x}_t)$ or $\bar{\mathbf{x}}_T \triangleq \sum_{t=1}^T \frac{2t}{T(T+1)} \mathbf{x}_t$.

And (ii)

$$\|\bar{\mathbf{x}}_T - \mathbf{x}^*\| \leq \frac{2G}{\sigma\sqrt{T+1}}.$$

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \nabla f(\mathbf{x}_t)]$$



$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t)]$$

Again, we extend it to online optimization.

OGD for Strongly Convex Functions

Online Gradient Descent (OGD)

At each round $t = 1, 2, \dots$

- (1) the player first picks a model $\mathbf{x}_t \in \mathcal{X}$;
- (2) and simultaneously environments pick a *strongly convex function* $f_t : \mathcal{X} \rightarrow \mathbb{R}$;
- (3) the player suffers loss $f_t(\mathbf{x}_t)$, observes the information of f_t and update the model according to $\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \nabla f_t(\mathbf{x}_t)]$.

OGD for Strongly Convex Functions

Theorem 4 (Regret bound for strongly-convex functions). *Under Assumption 1 (G-Lipschitz), for σ -strongly convex loss functions, online gradient descent with step sizes $\eta_t = \frac{1}{\sigma t}$ achieves the following guarantee*

$$\text{REG}_T \leq \frac{G^2}{2\sigma} (1 + \log T) = \mathcal{O}(\log T).$$

- Strongly convex case compared with convex case: $\mathcal{O}(\log T)$ vs. $\mathcal{O}(\sqrt{T})$
- A caveat is that we now don't need Assumption 2 (bounded domain).

OGD for Strongly Convex Functions

Proof: we start by extending *the first GD lemma* to strongly convex case.

Strongly convex case:

$$\begin{aligned}\|\mathbf{x}_{t+1} - \mathbf{u}\|^2 &\leq \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t \langle \nabla f_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2 \\ &\leq \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t \left(f_t(\mathbf{x}_t) - f_t(\mathbf{u}) + \frac{\sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 \right) + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2 \\ &\quad \text{(strong convexity: } f_t(\mathbf{x}_t) - f_t(\mathbf{u}) + \frac{\sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 \leq \langle \nabla f_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle \text{)} \\ &\leq (1 - \sigma\eta_t) \|\mathbf{x}_t - \mathbf{u}\|^2 - 2\eta_t (f_t(\mathbf{x}_t) - f_t(\mathbf{u})) + \eta_t^2 \|\nabla f_t(\mathbf{x}_t)\|^2 \\ \implies f_t(\mathbf{x}_t) - f_t(\mathbf{u}) &\leq \frac{\eta_t^{-1} - \sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 - \frac{\eta_t^{-1}}{2} \|\mathbf{x}_{t+1} - \mathbf{u}\|^2 + \frac{\eta_t G^2}{2} \quad \text{(rearranging)}\end{aligned}$$

OGD for Strongly Convex Functions

Proof:
$$f_t(\mathbf{x}_t) - f_t(\mathbf{u}) \leq \frac{\eta_t^{-1} - \sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 - \frac{\eta_t^{-1}}{2} \|\mathbf{x}_{t+1} - \mathbf{u}\|^2 + \frac{\eta_t G^2}{2}$$

Summing from $t = 1$ to T , setting $\eta_t = \frac{1}{\sigma t}$ (define $\frac{1}{\eta_0} := 0$):

$$\begin{aligned} 2 \sum_{t=1}^T (f_t(\mathbf{x}_t) - f_t(\mathbf{u})) &\leq \sum_{t=1}^T \|\mathbf{x}_t - \mathbf{u}\|^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - \sigma \right) + G^2 \sum_{t=1}^T \eta_t \quad \left(\frac{1}{\eta_0} := 0 \right) \\ &= 0 + G^2 \sum_{t=1}^T \frac{1}{\sigma t} \quad \left(\frac{1}{\eta_0} \triangleq 0, \|\mathbf{x}_{T+1} - \mathbf{u}\|^2 \geq 0, \frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - \sigma = 0 \right) \\ &\leq \frac{G^2}{\sigma} (1 + \log T). \quad \square \end{aligned}$$

Comparisons

Convex

Property: $f_t(\mathbf{x}) \geq f_t(\mathbf{y}) + \nabla f_t(\mathbf{y})^\top (\mathbf{x} - \mathbf{y})$

$$\text{OGD: } \mathbf{x}_{t+1} = \Pi_{\mathcal{X}} \left[\mathbf{x}_t - \frac{1}{\sqrt{t}} \nabla f_t(\mathbf{x}_t) \right]$$

$$\text{REG}_T \leq \frac{3}{2} G D \sqrt{T}$$

Strongly Convex

Property: $f_t(\mathbf{y}) \geq f_t(\mathbf{x}) + \nabla f_t(\mathbf{x})^\top (\mathbf{y} - \mathbf{x}) + \frac{\sigma}{2} \|\mathbf{y} - \mathbf{x}\|^2$

$$\text{OGD: } \mathbf{x}_{t+1} = \Pi_{\mathcal{X}} \left[\mathbf{x}_t - \frac{1}{\sigma t} \nabla f_t(\mathbf{x}_t) \right]$$

$$\text{REG}_T \leq \frac{G^2}{2\sigma} (1 + \log T)$$

Lower Bounds

- A natural question: whether previous regret can be improved?
- Lower bound argument:

minimax bound: smallest possible worst-case regret of any algorithm

$$\min_{\mathcal{A}} \max_{\ell_1, \dots, \ell_T} \text{REG}_T$$

Theorem 5 (Lower Bound for OCO). *Any algorithm for online convex optimization incurs $\Omega(DG\sqrt{T})$ regret in the worst case. This is true even if the cost functions are generated from a fixed stationary distribution.*

Lower Bounds

Theorem 5 (Lower Bound for OCO). *Any algorithm for online convex optimization incurs $\Omega(DG\sqrt{T})$ regret in the worst case. This is true even if the cost functions are generated from a fixed stationary distribution.*

Proof Sketch.

Construct a *hard* environment:

- Binary classification, loss functions in each iteration are chosen at random
- The hardness comes from the compression of a sequence of T -round random bits

Comparison

	Algo.	Step size	Upper Bound	Lower Bound
Convex	OGD	$\frac{D}{G\sqrt{t}}$	$\mathcal{O}(\sqrt{T})$	$\Omega(\sqrt{T})$
σ -Strongly Convex	OGD	$\frac{1}{\sigma t}$	$\mathcal{O}\left(\frac{\log T}{\sigma}\right)$	$\Omega\left(\frac{\log T}{\sigma}\right)$

Lower bound for strongly convex functions is more non-trivial.

Jacob Abernethy, Peter L. Bartlett, Alexander Rakhlin, and Ambuj Tewari. Optimal strategies and minimax lower bounds for online convex games. In COLT, 2008.

Optimal Strategies and Minimax Lower Bounds for Online Convex Games

Jacob Abernethy*
UC Berkeley
jake@cs.berkeley.edu

Alexander Rakhlin*
UC Berkeley
rakhlin@cs.berkeley.edu

Abstract

A number of learning problems can be cast as an Online Convex Game: on each round, a learner makes a prediction x from a convex set, the environment plays a loss function f_t , and the learner's long-term goal is to minimize regret. Algorithms have been proposed by Zinkevich, when f is assumed to be convex, and Hazan et al., when f is assumed to be strongly convex, that have provably low regret. We consider these two settings and analyze such games from a minimax perspective, proving minimax strategies and lower bounds in each case. These results prove that the existing algorithms are essentially optimal.

1 Introduction

The decision maker's greatest fear is *regret*: knowing, with the benefit of hindsight, that a better alternative existed. Yet, given only hindsight and not the gift of foresight, imperfect decisions can not be avoided. It is thus the decision maker's ultimate goal to suffer as little regret as possible.

In the present paper, we consider the notion of "regret minimization" for a particular class of decision problems. Assume we are given a set X and some set of functions $\mathcal{F}$ on X . On each round $t = 1, \dots, T$, we must choose some x_t from a set X . After we have made this choice, the environment chooses a function $f_t \in \mathcal{F}$. We incur a cost (loss) $f_t(x_t)$, and the game proceeds to the next round. Of course, had we the fortune of perfect foresight and had access to the sum $f_1 + \dots + f_T$, we would know the optimal choice $x^* = \arg \min_{x \in X} \sum_{t=1}^T f_t(x)$. Instead, at time t , we will have only seen $f_1, \dots, f_{t-1}$, and we must make the decision x_t with only historical knowledge. Thus, a natural long-term goal is to minimize the *regret*, which here we define as

$$\sum_{t=1}^T f_t(x_t) - \inf_{x \in X} \sum_{t=1}^T f_t(x).$$

A special case of this setting is when the decision space X is a convex set and $\mathcal{F}$ is some set of convex functions on X . In

Peter L. Bartlett†
UC Berkeley
bartlett@cs.berkeley.edu

Ambuj Tewari
TTI Chicago
ambuj@cs.berkeley.edu

the literature, this framework has been referred to as Online Convex Optimization (OCO), since our goal is to minimize a global function, i.e. $f_1 + f_2 + \dots + f_T$, while this objective is revealed to us but one function at a time. Online Convex Optimization has attracted much interest in recent years [4, 9, 6, 1], as it provides a general analysis for a number of standard online learning problems including, among others, online classification and regression, prediction with expert advice, the portfolio selection problem, and online density estimation.

While instances of OCO have been studied over the past two decades, the general problem was first analyzed by Zinkevich [9], who showed that a very simple and natural algorithm, online gradient descent, elicits a bound on the regret that is on the order of $\sqrt{T}$. Online gradient descent can be described simply by the update $x_{t+1} = x_t - \eta \nabla f_t(x_t)$, where η is some parameter of the algorithm. This regret bound only required that f_t be smooth, convex, and with bounded derivative.

A regret bound of order $\mathcal{O}(\sqrt{T})$ is not surprising: a number of online learning problems give rise to similar bounds. More recently, however, Hazan et al. [4] showed that when $\mathcal{F}$ consists of *curved* functions, i.e. f_t is strongly convex, then we get a bound of the form $\mathcal{O}(\log T)$. It is quite surprising that curvature gives such a great advantage to the player. Curved loss functions, such as square loss or logarithmic loss, are very natural in a number of settings.

Finding algorithms that can guarantee low regret is, however, only half of the story; indeed, it is natural to ask "can we obtain even lower regret?" or "do better algorithms exist?" The goal of the present paper is to address these questions, in some detail, for several classes of such online optimization problems. We answer both in the negative: the algorithms of Zinkevich and Hazan et al. are tight even up to their multiplicative constants.

This is achieved by a game-theoretic analysis: if we pose the above online optimization problem as a game between a Player who chooses x_t and an Adversary who chooses f_t , we may consider the regret achieved when each player is playing optimally. This is typically referred to as the *value* V_T of the game. In general, computing the value of zero-sum games is difficult, as we may have to consider exponentially many, or even uncountably many, strategies of the Player and the Adversary. Ultimately we will show that this value, as well as the optimal strategies of both the player and the adversary,

*Division of Computer Science
†Division of Computer Science, Department of Statistics

Part 3. Online-to-Batch Conversion

- Convex Functions
- Strongly Convex Functions

Offline Optimization

- Consider *offline optimization* $\min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x})$ with stochastic opt method

Computational oracle: only access *noisy* gradient oracle, namely, $\mathbf{g}(\mathbf{x})$, such that

$$\mathbb{E}[\mathbf{g}(\mathbf{x})] = \nabla F(\mathbf{x}), \text{ and } \mathbb{E}[\|\mathbf{g}(\mathbf{x})\|^2] \leq G^2$$

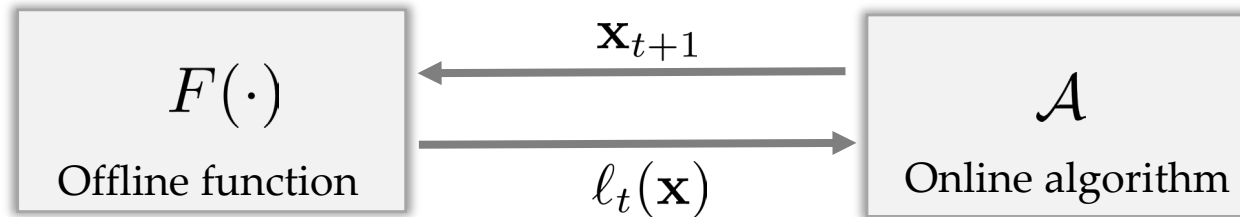
for some $G > 0$.

Example (large-scale opt.). Given dataset $S = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_m, y_m)\}$, ERM optimizes

$$\min_{h \in \mathcal{H}} \sum_{i=1}^m \ell(h(\mathbf{x}_i), y_i) \quad \Rightarrow \quad \begin{array}{l} \text{full gradient computation requires a pass of } \textit{all data} \\ \text{stochastic method only uses a } \textit{mini batch} \text{ at each round} \end{array}$$

Online-to-Batch Conversion

- Reducing *offline optimization* as an *online optimization*.



Algorithm 2 Online-to-Batch Conversion for Convex Functions

Input: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$, online learning algorithm $\mathcal{A}$

- 1: **for** $t = 1, \dots, T$ **do**
 - 2: Obtain noisy gradient $\mathbf{g}(\mathbf{x}_t)$
 - 3: Pass loss function $\ell_t(\mathbf{x}) = \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle$ to online learning algorithm $\mathcal{A}$
 - 4: Receive next point $\mathbf{x}_{t+1}$ from $\mathcal{A}$
 - 5: **end for**
 - 6: **return** $\bar{\mathbf{x}}_T = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$
-

Online-to-Batch Conversion

Theorem 6 (Vanilla O2B Conversion). Assume $F(\mathbf{x})$ is convex, and let $\mathbf{x}^* \in \arg \min_{\mathbf{x}} F(\mathbf{x})$ and $\text{REG}_T^{\mathcal{A}}(\mathbf{u}) \triangleq \sum_{t=1}^T \ell_t(\mathbf{x}_t) - \ell_t(\mathbf{u})$ be the regret of the online learning algorithm $\mathcal{A}$ on the sequence of loss functions $\ell_t(\mathbf{x}) = \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle$. Algorithm 2 achieves the following guarantee:

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \frac{\mathbb{E}[\text{REG}_T^{\mathcal{A}}(\mathbf{x}^*)]}{T}.$$

Proof: Using Jensen's inequality, we have

$$\begin{aligned} \mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{u}) &\stackrel{\text{(Jensen)}}{\leq} \frac{1}{T} \sum_{t=1}^T \mathbb{E}[F(\mathbf{x}_t)] - F(\mathbf{u}) \stackrel{\text{(Convexity)}}{\leq} \frac{1}{T} \sum_{t=1}^T \mathbb{E}[\langle \nabla F(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle] \\ &= \frac{1}{T} \sum_{t=1}^T \mathbb{E}[\langle \mathbf{g}(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle] = \frac{\mathbb{E}[\text{REG}_T^{\mathcal{A}}(\mathbf{u})]}{T}. \quad \square \end{aligned}$$

Stochastic Optimization

- Optimization Goal

$$\min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x})$$

Computational oracle: only access *noisy* gradient oracle, namely, $\mathbf{g}(\mathbf{x})$, such that

$$\mathbb{E}[\mathbf{g}(\mathbf{x})] = \nabla F(\mathbf{x}), \text{ and } \mathbb{E}[\|\mathbf{g}(\mathbf{x})\|^2] \leq G^2.$$

By deploying OGD as $\mathcal{A}$ in Algorithm 2, we obtain a practical algorithm for convex functions, which updates as follows:

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \mathbf{g}(\mathbf{x}_t)]$$

Plugging the $\mathcal{O}(\sqrt{T})$ regret bound of OGD into Theorem 6, we can achieve $\mathcal{O}\left(\frac{1}{\sqrt{T}}\right)$ rate.

Algorithm 2 Online-to-Batch Conversion for Convex Functions

Input: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$, online learning algorithm $\mathcal{A}$

```
1: for  $t = 1, \dots, T$  do
2:   Obtain noisy gradient  $\mathbf{g}(\mathbf{x}_t)$ 
3:   Pass loss function  $\ell_t(\mathbf{x}) = \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle$  to online learning algorithm  $\mathcal{A}$ 
4:   Receive next point  $\mathbf{x}_{t+1}$  from  $\mathcal{A}$ 
5: end for
6: return  $\bar{\mathbf{x}}_T = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$ 
```

Understand SGD from Online Learning

Algorithm 3 Stochastic Gradient Descent

Input: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$

```
1: for  $t = 1, \dots, T$  do  
2:   Obtain noisy gradient  $\mathbf{g}(\mathbf{x}_t)$   
3:   Update the model  $\mathbf{x}_{t+1} = \Pi_{\mathcal{X}}[\mathbf{x}_t - \eta_t \mathbf{g}(\mathbf{x}_t)]$   
4: end for  
5: return  $\bar{\mathbf{x}}_T = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$ 
```

$$\begin{aligned}\mathbf{OGD}: \mathbf{x}_{t+1} &= \Pi_{\mathcal{X}}[\mathbf{x}_t - \eta_t \nabla \ell_t(\mathbf{x}_t)] \\ &= \Pi_{\mathcal{X}}[\mathbf{x}_t - \eta_t \mathbf{g}(\mathbf{x}_t)]\end{aligned}$$

SGD is **equivalent to** deploying OGD in Algorithm 2 over functions $\{\ell_t(\mathbf{x})\}_{t=1}^T$.

Theorem 7 (Convergence of SGD). Assume $F(\mathbf{x})$ is convex, and let $\mathbf{x}^* \in \arg \min_{\mathbf{x}} F(\mathbf{x})$. Algorithm 3 achieves the following guarantee:

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \mathcal{O}\left(\frac{1}{\sqrt{T}}\right).$$

Stochastic Optimization

- Optimization Goal

$$\min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x})$$

Computational oracle: only access *noisy* gradient oracle, namely, $\mathbf{g}(\mathbf{x})$, such that

$$\mathbb{E}[\mathbf{g}(\mathbf{x})] = \nabla F(\mathbf{x}), \text{ and } \mathbb{E}[\|\mathbf{g}(\mathbf{x})\|^2] \leq G^2.$$

By deploying OGD as $\mathcal{A}$ in Algorithm 2, we obtain a practical algorithm for convex functions, which updates as follows:

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}} [\mathbf{x}_t - \eta_t \mathbf{g}(\mathbf{x}_t)]$$

Plugging the $\mathcal{O}(\sqrt{T})$ regret bound of OGD into Theorem 6, we can achieve $\mathcal{O}\left(\frac{1}{\sqrt{T}}\right)$ rate.

⇒ Note that function $\ell_t(\mathbf{x}) \triangleq \mathbf{g}(\mathbf{x}_t)^\top \mathbf{x}$ actually *depends* on the decision $\mathbf{x}_t$, which reveals that OGD regret can hold even against *adaptive adversary*.

Algorithm 2 Online-to-Batch Conversion for Convex Functions

Input: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$, online learning algorithm $\mathcal{A}$

```
1: for  $t = 1, \dots, T$  do  
2:   Obtain noisy gradient  $\mathbf{g}(\mathbf{x}_t)$   
3:   Pass loss function  $\ell_t(\mathbf{x}) = \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle$  to online learning algorithm  $\mathcal{A}$   
4:   Receive next point  $\mathbf{x}_{t+1}$  from  $\mathcal{A}$   
5: end for  
6: return  $\bar{\mathbf{x}}_T = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$ 
```

O2B Conversion for Strongly Convex

- For σ -strongly convex functions, we can revise the O2B conversion and select OGD as $\mathcal{A}$ to achieve an $\mathcal{O}\left(\frac{\log T}{T}\right)$ convergence rate.

Algorithm 4 Online-to-Batch Conversion for Convex Functions

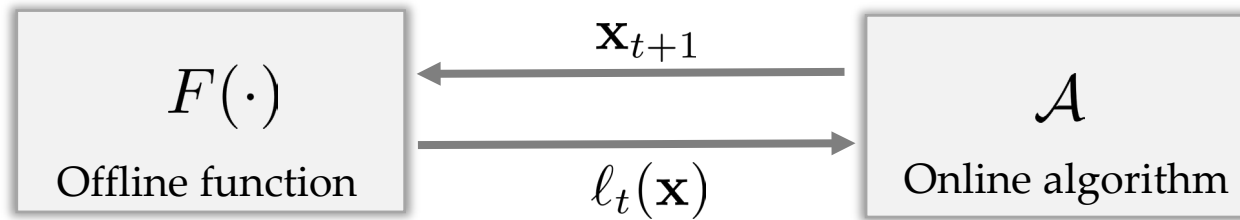
Require: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$, online learning algorithm $\mathcal{A}$

- 1: **for** $t = 1, \dots, T$ **do**
 - 2: Obtain noisy gradient $\mathbf{g}(\mathbf{x}_t)$
 - 3: Pass loss function $\ell_t(\mathbf{x}) = \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle + \frac{\sigma}{2} \|\mathbf{x}_t - \mathbf{x}\|^2$ to online learning algorithm $\mathcal{A}$
 - 4: Receive next point $\mathbf{x}_{t+1}$ from $\mathcal{A}$ *retain strong convexity*
 - 5: **end for**
 - 6: **return** $\bar{\mathbf{x}}_T = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$
-

However, this rate exhibits a gap to the $\Omega(1/T)$ lower bound for stochastic optimization over strongly convex functions.

Weighted O2B Conversion

- To achieve the optimal rate for strongly convex optimization, we introduce the *weighted* Online-to-Batch Conversion.



Algorithm 5 Weighted O2B Conversion for Strongly Convex Functions

Input: noisy gradient oracle $\mathbf{g}(\cdot)$, step sizes $\{\eta_t\}$, online learning algorithm $\mathcal{A}$, and **weights** $\{\alpha_t\}_{t=1}^T$

- 1: **for** $t = 1, \dots, T$ **do**
 - 2: Obtain noisy gradient $\mathbf{g}(\mathbf{x}_t)$
 - 3: Pass loss function $\ell_t(\mathbf{x}) = \langle \alpha_t \mathbf{g}(\mathbf{x}_t), \mathbf{x} \rangle + \frac{\alpha_t \sigma}{2} \|\mathbf{x}_t - \mathbf{x}\|^2$ to online learning algorithm $\mathcal{A}$
 - 4: Receive next point $\mathbf{x}_{t+1}$ from $\mathcal{A}$
 - 5: **end for**
 - 6: **return** $\bar{\mathbf{x}}_T = \frac{1}{\sum_{t=1}^T \alpha_t} \sum_{t=1}^T \alpha_t \mathbf{x}_t$
-

Weighted O2B Conversion

Theorem 8 (Weighted O2B Conversion for Strongly Convex Functions). Assume $F(\mathbf{x})$ is σ -strongly convex, and let $\mathbf{x}^* = \arg \min_{\mathbf{x}} F(\mathbf{x})$ and $\text{REG}_T^{\mathcal{A}}(\mathbf{u}) = \sum_{t=1}^T \ell_t(\mathbf{x}_t) - \ell_t(\mathbf{u})$ be the regret of the online learning algorithm $\mathcal{A}$ on the sequence of loss functions $\ell_t(\mathbf{x}) = \langle \alpha_t \mathbf{g}_t, \mathbf{x} \rangle + \frac{\alpha_t \sigma}{2} \|\mathbf{x}_t - \mathbf{x}\|^2$. Algorithm 4 achieves the following guarantee:

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \frac{\mathbb{E} [\text{REG}_T^{\mathcal{A}}(\mathbf{u})]}{\sum_{t=1}^T \alpha_t}.$$

Proof:

$$\begin{aligned} \mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{u}) &\leq \frac{1}{\sum_{t=1}^T \alpha_t} \sum_{t=1}^T \alpha_t (\mathbb{E}[F(\mathbf{x}_t)] - F(\mathbf{u})) \\ &\leq \frac{1}{\sum_{t=1}^T \alpha_t} \sum_{t=1}^T \mathbb{E} \left[\alpha_t \langle \mathbf{g}(\mathbf{x}_t), \mathbf{x}_t - \mathbf{u} \rangle - \frac{\alpha_t \sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 \right] = \frac{\mathbb{E} [\text{REG}_T^{\mathcal{A}}(\mathbf{u})]}{\sum_{t=1}^T \alpha_t}. \quad \square \end{aligned}$$

O2B for Strongly Convex Functions

Theorem 9 (O2B for Strongly Convex Functions). Assume $F(\mathbf{x})$ is σ -strongly-convex, and let $\mathbf{x}^* = \arg \min_{\mathbf{x}} F(\mathbf{x})$. By setting $\mathcal{A}$ as OGD with learning rate $\eta_t = \frac{1}{\sigma \sum_{i=1}^t \alpha_i}$ and $\alpha_t = t$, Algorithm 4 achieves the following guarantee:

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \mathcal{O}\left(\frac{1}{T}\right).$$

Recall that in Theorem 8, we have

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \frac{\mathbb{E}[\text{REG}_T^{\mathcal{A}}(\mathbf{u})]}{\sum_{t=1}^T \alpha_t}.$$

First, as $\alpha_t = t$, the denominator becomes $\frac{T(T+1)}{2}$.

Next, we focus on $\mathbb{E}[\text{REG}_T^{\mathcal{A}}(\mathbf{u})]$, where $\ell_t(\mathbf{x})$ is defined by $\ell_t(\mathbf{x}) = \langle \alpha_t \mathbf{g}_t, \mathbf{x} \rangle + \frac{\alpha_t \sigma}{2} \|\mathbf{x}_t - \mathbf{x}\|^2$.

O2B for Strongly Convex Functions

$$G_t = \alpha_t G$$

Proof:

$$\mathbb{E} [\text{REG}_T^{\mathcal{A}}] = \mathbb{E} \left[\sum_{t=1}^T (f_t(\mathbf{x}_t) - f_t(\mathbf{u})) \right] \leq \mathbb{E} \left[\sum_{t=1}^T \|\mathbf{x}_t - \mathbf{u}\|^2 \left(\frac{1}{2\eta_t} - \frac{1}{2\eta_{t-1}} - \frac{\sigma}{2} \right) \right] + \mathbb{E} \left[\sum_{t=1}^T \frac{G_t^2 \eta_t}{2} \right]$$

OCO with Strongly Convex Functions

Proof: $f_t(\mathbf{x}_t) - f_t(\mathbf{u}) \leq \frac{\eta_t^{-1} - \sigma}{2} \|\mathbf{x}_t - \mathbf{u}\|^2 - \frac{\eta_t^{-1}}{2} \|\mathbf{x}_{t+1} - \mathbf{u}\|^2 + \frac{\eta_t G^2}{2}$

Summing from $t = 1$ to T , setting $\eta_t = \frac{1}{\sigma t}$ (define $\frac{1}{\eta_0} := 0$):

$$\begin{aligned} 2 \sum_{t=1}^T (f_t(\mathbf{x}_t) - f_t(\mathbf{u})) &\leq \sum_{t=1}^T \|\mathbf{x}_t - \mathbf{u}\|^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - \sigma \right) + G^2 \sum_{t=1}^T \eta_t \quad \left(\frac{1}{\eta_0} := 0 \right) \\ &= 0 + G^2 \sum_{t=1}^T \frac{1}{\sigma t} \quad \left(\frac{1}{\eta_0} \triangleq 0, \|\mathbf{x}_{T+1} - \mathbf{u}\|^2 \geq 0, \frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} - \sigma = 0 \right) \\ &\leq \frac{G^2}{\sigma} (1 + \log T). \quad \square \end{aligned}$$

Choosing $\eta_t = \frac{1}{\sigma \sum_{s=1}^t \alpha_s}$, we have

$$\mathbb{E} [\text{REG}_T^{\mathcal{A}}(\mathbf{x}^*)] \leq \sum_{t=1}^T \frac{G_t^2}{2\sigma \sum_{i=1}^t \alpha_i} = \mathcal{O}(T).$$

Plugging it back to Theorem 8, we obtain

$$\mathbb{E}[F(\bar{\mathbf{x}}_T)] - F(\mathbf{x}^*) \leq \mathcal{O} \left(\frac{1}{T} \right).$$

□

More bits of OGD

- Note that function $\ell_t(\mathbf{x}) \triangleq \mathbf{g}(\mathbf{x}_t)^\top \mathbf{x}$ *depends* on the decision $\mathbf{x}_t$, which actually reveals that OGD regret can hold even against *adaptive adversary*.

At each round $t = 1, 2, \dots$

- (1) the player first picks a model $\mathbf{x}_t \in \mathcal{X}$;
- (2) and *simultaneously* environments pick an online function $f_t : \mathcal{X} \rightarrow \mathbb{R}$;
- (3) the player suffers loss $f_t(\mathbf{x}_t)$, observes some information about f_t and updates the model.

oblivious adversary



examination

adaptive adversary



interview

The “simultaneously” requirement can be sometimes not necessary!

OGD for full-info OCO can handle the case when online functions *depend on $\mathbf{x}_t$* !

History: SGD

Robbins-Monro Method

A STOCHASTIC APPROXIMATION METHOD¹

By HERBERT ROBBINS AND SUTTON MONRO

University of North Carolina

1. Summary. Let $M(x)$ denote the expected value at level x of the response to a certain experiment. $M(x)$ is assumed to be a monotone function of x but is unknown to the experimenter, and it is desired to find the solution $x = \theta$ of the equation $M(x) = \alpha$, where α is a given constant. We give a method for making successive experiments at levels $x_1, x_2, \dots$ in such a way that x_n will tend to θ in probability.

2. Introduction. Let $M(x)$ be a given function and α a given constant such that the equation

$$(1) \quad M(x) = \alpha$$

has a unique root $x = \theta$. There are many methods for determining the value of θ by successive approximation. With any such method we begin by choosing one or more values $x_1, \dots, x_r$ more or less arbitrarily, and then successively obtain new values x_n as certain functions of the previously obtained $x_1, \dots, x_{n-1}$, the values $M(x_1), \dots, M(x_{n-1})$, and possibly those of the derivatives $M'(x_1), \dots, M'(x_{n-1})$, etc. If

$$(2) \quad \lim_{n \rightarrow \infty} x_n = \theta,$$

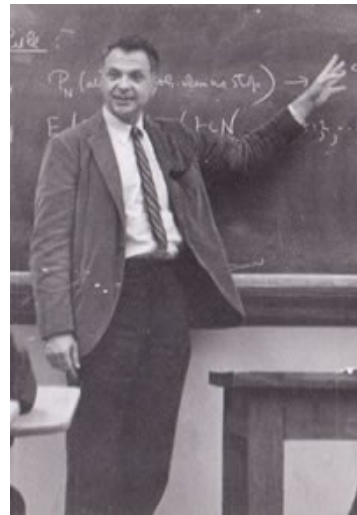
irrespective of the arbitrary initial values $x_1, \dots, x_r$, then the method is effective for the particular function $M(x)$ and value α . The speed of the convergence in (2) and the ease with which the x_n can be computed determine the practical utility of the method.

We consider a stochastic generalization of the above problem in which the nature of the function $M(x)$ is unknown to the experimenter. Instead, we suppose that to each value x corresponds a random variable $Y = Y(x)$ with distribution function $Pr\{Y(x) \leq y\} = H(y | x)$, such that

$$(3) \quad M(x) = \int_{-\infty}^{\infty} y dH(y | x)$$

is the expected value of Y for the given x . Neither the exact nature of $H(y | x)$ nor that of $M(x)$ is known to the experimenter, but it is assumed that equation (1) has a unique root θ , and it is desired to estimate θ by making successive observations on Y at levels $x_1, x_2, \dots$ determined sequentially in accordance with some definite experimental procedure. If (2) holds in probability irrespective of any arbitrary initial values $x_1, \dots, x_r$, we shall, in conformity with usual statistical terminology, call the procedure consistent for the given $H(y | x)$ and value α .

¹ This work was supported in part by the Office of Naval Research.



Herbert Ellis Robbins
(1915 - 2001)

A Stochastic Approximation Method.

Herbert Robbins, Sutton Monro

Ann. Math. Statist. 22(3): 400-407 (September, 1951).

Kiefer-Wolfowitz Method

STOCHASTIC ESTIMATION OF THE MAXIMUM OF A REGRESSION FUNCTION¹

By J. KIEFER AND J. WOLFOWITZ

Cornell University

1. Summary. Let $M(x)$ be a regression function which has a maximum at the unknown point θ . $M(x)$ is itself unknown to the statistician who, however, can take observations at any level x . This paper gives a scheme whereby, starting from an arbitrary point x_1 , one obtains successively $x_2, x_3, \dots$ such that x_n converges to θ in probability as $n \rightarrow \infty$.

2. Introduction. Let $H(y | x)$ be a family of distribution functions which depend on a parameter x , and let

$$(2.1) \quad M(x) = \int_{-\infty}^{\infty} y dH(y | x).$$

We suppose that

$$(2.2) \quad \int_{-\infty}^{\infty} (y - M(x))^2 dH(y | x) \leq S < \infty,$$

and that $M(x)$ is strictly increasing for $x < \theta$, and $M(x)$ is strictly decreasing for $x > \theta$. Let $\{a_n\}$ and $\{c_n\}$ be infinite sequences of positive numbers such that

$$(2.3) \quad c_n \rightarrow 0,$$

$$(2.4) \quad \sum a_n = \infty,$$

$$(2.5) \quad \sum a_n c_n < \infty,$$

$$(2.6) \quad \sum a_n^2 c_n^{-2} < \infty.$$

(For example, $a_n = n^{-1}$, $c_n = n^{-1/3}$.)

We can now describe a recursive scheme as follows. Let z_1 be an arbitrary number. For all positive integral n we have

$$(2.7) \quad z_{n+1} = z_n + a_n \frac{(y_{2n} - y_{2n-1})}{c_n},$$

where y_{2n-1} and y_{2n} are independent chance variables with respective distributions $H(y | z_n - c_n)$ and $H(y | z_n + c_n)$. Under regularity conditions on $M(x)$ which we shall state below we will prove that z_n converges stochastically to θ (as $n \rightarrow \infty$).

The statistical importance of this problem is obvious and need not be discussed. The stimulus for this paper came from the interesting paper by Robbins and Monro [1] (see also Wolfowitz [2]).

¹ Research under contract with the Office of Naval Research. Presented to the American Mathematical Society at New York on April 25, 1952.



Jack Kiefer
(1924 - 1981)



Jacob Wolfowitz
(1910 - 1981)

Stochastic Estimation of the Maximum of a Regression Function

Jack Kiefer, Jacob Wolfowitz

Ann. Math. Statist. 23(3): 462-466 (September, 1952)

History: SGD



Herbert Ellis Robbins
(1915 - 2001)

Statistical Science
1986, Vol. 1, No. 2, 276–284

The Contributions of Herbert Robbins to Mathematical Statistics

Tze Leung Lai and David Siegmund

Herbert Robbins was born on January 12, 1915, in New Castle, Pennsylvania. In 1931 he entered Harvard College at the age of 16. Although his interests until then had been predominantly literary, he found himself increasingly attracted to mathematics under the influence of Marston Morse, who during many long conversations conveyed a vivid sense of the intellectual challenge of creative work in that field (cf. Page, 1984, p. 7). He received the A.B. summa cum laude in 1935, and the Ph.D. in 1938, also from Harvard. His thesis, in the field of combinatorial topology and written under the supervision of Hassler Whitney, was published in 1941 [3]. (Numbers in brackets refer to Robbins' bibliography at the end of this article.)

After graduation, Robbins worked for a year at the Institute for Advanced Study at Princeton as Marston Morse's assistant. He then spent the next three years at New York University as instructor in mathematics. He became nationally known in 1941 as the coauthor,

North Carolina at Chapel Hill. Having read [7] and [10], and greatly impressed by Robbins' mathematical skills, Hotelling offered him the position of associate professor to teach measure theory and probability to the graduate students in the new department. Robbins accepted the position and spent the next six years at Chapel Hill. During this relatively short period Robbins not only studied and developed an increasingly deep interest in statistics, but he also made a number of profound contributions to his new field: complete convergence [12], compound decision theory [25], stochastic approximation [26], and the sequential design of experiments [28], to name a few.

After a Guggenheim Fellowship at the Institute for Advanced Study during 1952–1953, Robbins moved from Chapel Hill to Columbia University as professor and chairman of the Department of Mathematical Statistics. Since 1953, with the exception of the three years 1965–1968 spent at Minnesota, Purdue, Berkeley, and Michigan, he has been at Columbia, where he

History: Two-Player Zero-Sum Games

Theory of repeated games



James Hannan
(1922–2010)



David Blackwell
(1919–2010)

Learning to play a game (1956)

Play a game repeatedly against a possibly suboptimal opponent

Zero-sum 2-person games played more than once

	1	2	...	M
1	$\ell(1,1)$	$\ell(1,2)$	...	
2	$\ell(2,1)$	$\ell(2,2)$	...	
$\vdots$	$\vdots$	$\vdots$	$\ddots$	
N				

$N \times M$ known loss matrix

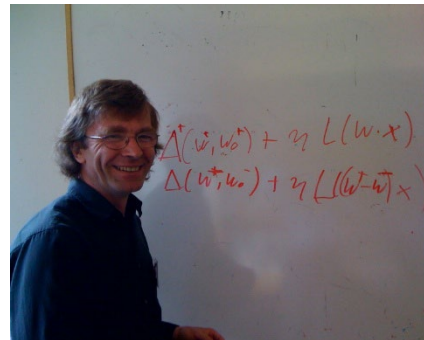
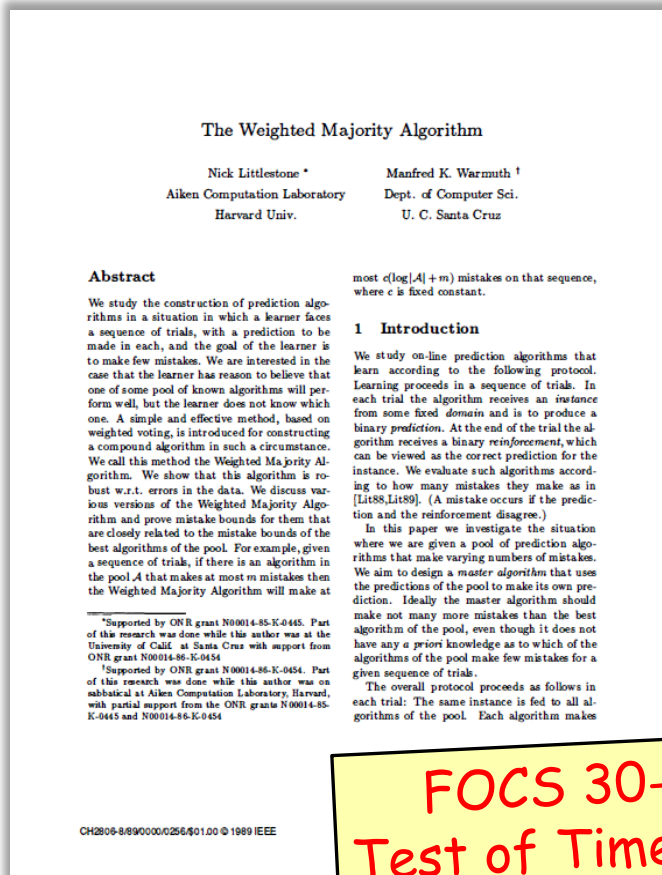
- Row player (**player**) has N actions
- Column player (**opponent**) has M actions

For each game round $t = 1, 2, \dots$

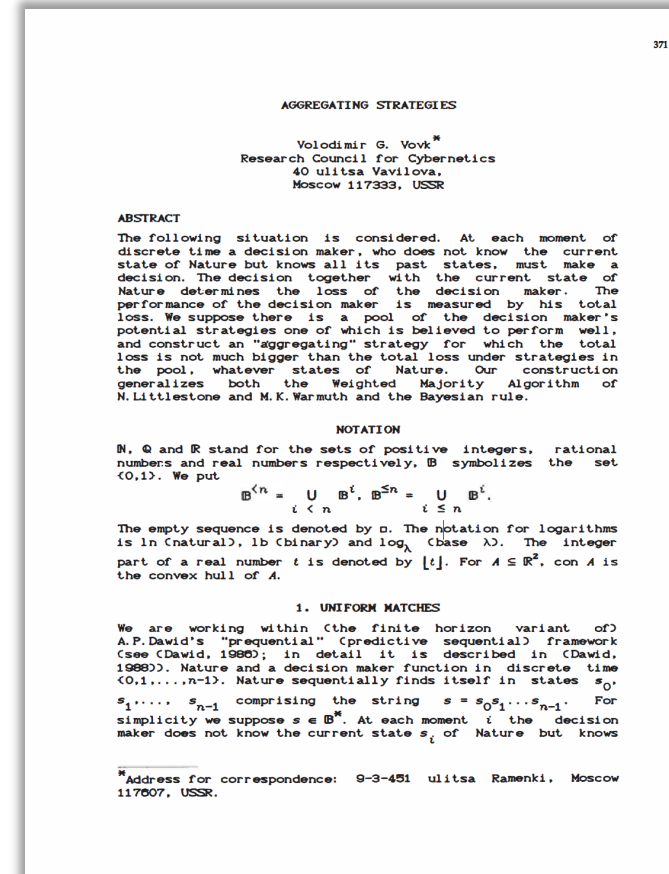
- Player chooses action i_t and opponent chooses action y_t
- The player suffers loss $\ell(i_t, y_t)$ (= gain of opponent)

Player can learn from opponent's history of past choices $y_1, \dots, y_{t-1}$

History: Prediction with Expert Advice



Manfred Warmuth
UC Santa Cruz



Volodimir G. Vovk
Royal Holloway,
University of London

Nick Littlestone and Manfred K. Warmuth.
"The Weighted Majority Algorithm." FOCS 1989: 256-261.

Volodimir G. Vovk. "Aggregating Strategies." COLT 1990: 371-383.

Summary

ONLINE OPTIMIZATION

Interaction optimization

Game-theoretic language

ONLINE-TO-BATCH CONVERSION

Online-to-batch conversion

Weighted O2B conversion

SGD, stochastic optimization

ONLINE CONVEX OPTIMIZATION

Problem formulation

Regret measure

Convex functions: OGD

Strongly convex functions: OGD

Q & A

Thanks!