

动态环境下的在线学习

张利军

南京大学

计算机科学与技术系



目录

➤ 研究背景

- 机器学习

- 大数据

➤ 在线学习

- 遗憾

- 在线凸优化

➤ 动态环境

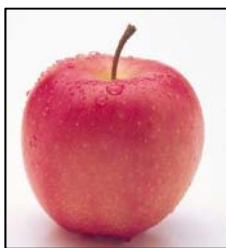
- 动态遗憾

- 自适应算法

➤ 总结

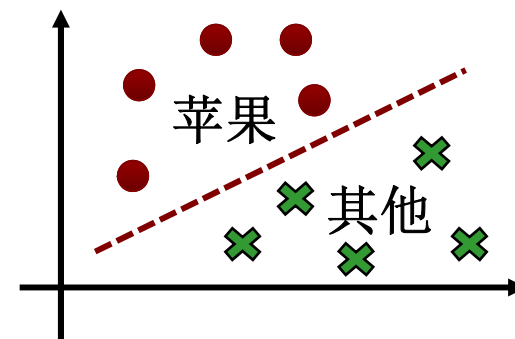
机器学习

➤ 监督学习



苹果

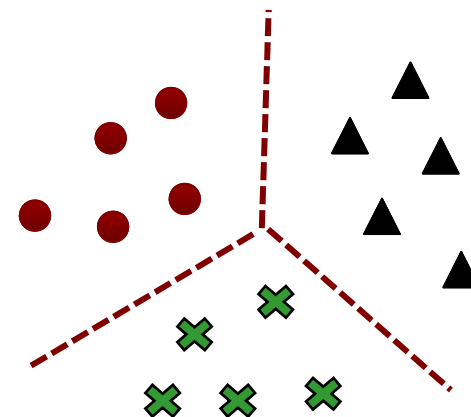
$$\begin{matrix} (\mathbf{x}_1, y_1) \\ \dots \\ (\mathbf{x}_n, y_n) \end{matrix} \Rightarrow y \approx h(\mathbf{x})$$



➤ 无监督学习



$$\begin{matrix} \mathbf{x}_1 \\ \dots \\ \mathbf{x}_n \end{matrix} \Rightarrow h(\mathbf{x})$$



➤ 经验风险最小化

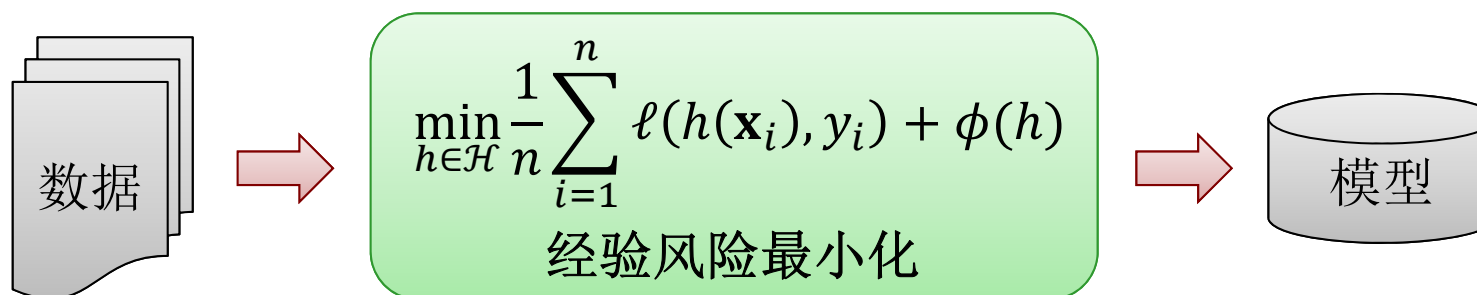
$$\min_{h \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n \ell(h(\mathbf{x}_i), y_i) + \phi(h)$$

□ $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \in \mathbb{R}^d \times \mathbb{R}$ 是训练数据

□ n 是样本数量、 d 是数据维度、 $\mathcal{H}$ 为假设空间

□ $\ell(\cdot, \cdot)$ 为损失函数、 $\phi(\cdot)$ 为正则化因子

➤ 批量学习



➤ 支持向量机

$$\min_{\mathbf{w}} \frac{1}{n} \sum_{i=1}^n \max(0, 1 - y_i \mathbf{w}^\top \mathbf{x}_i) + \frac{\lambda}{2} \|\mathbf{w}\|^2$$

□ 铰链损失: $\ell(u, v) = \max(0, 1 - uv)$

□ 线性函数空间: $\mathcal{H} = \{h(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}\}$

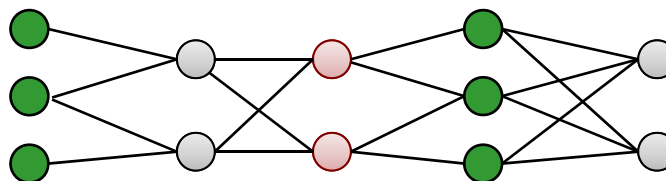
□ 正则化: $\lambda \|\mathbf{w}\|^2 / 2$

➤ 深度学习

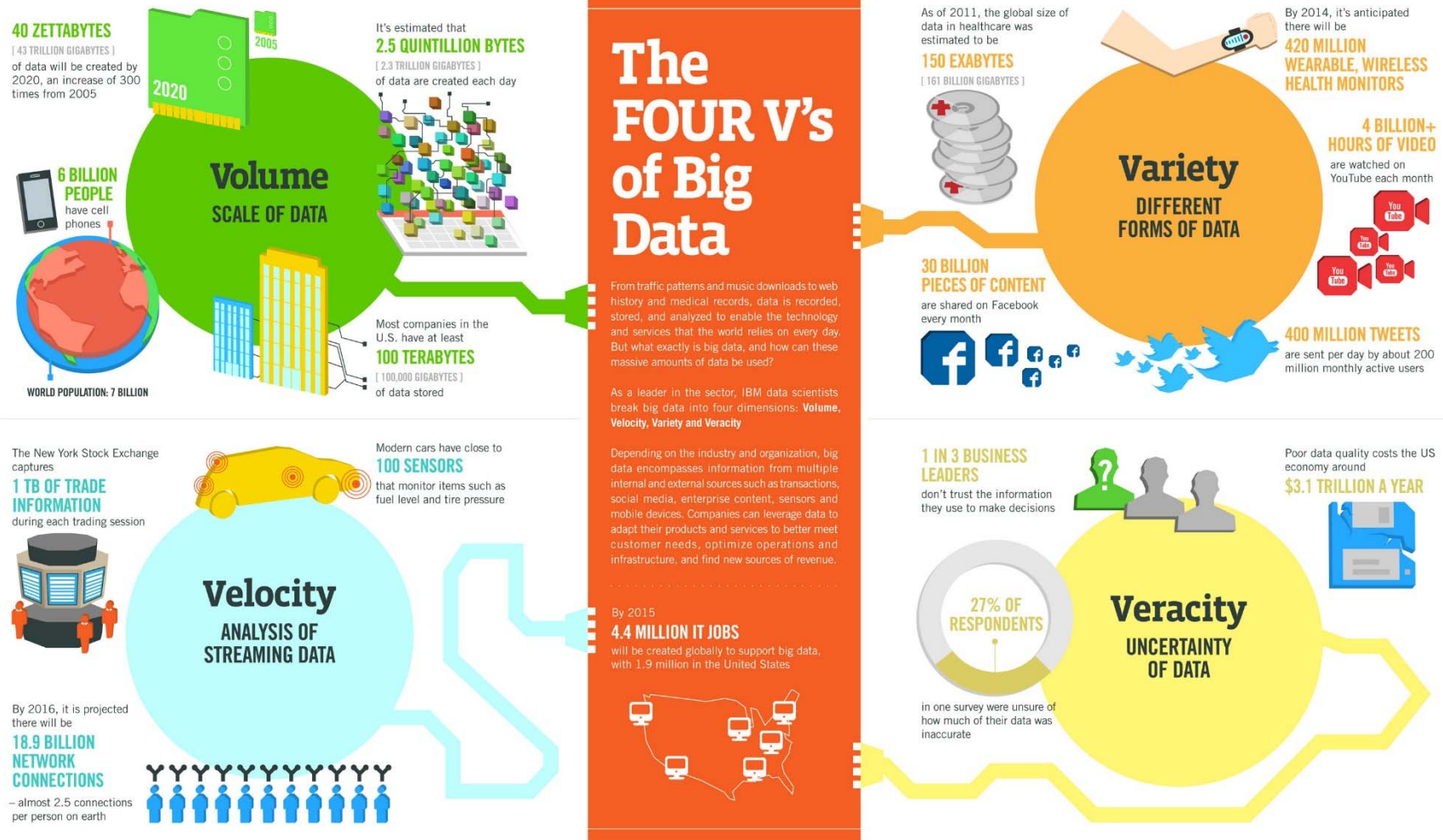
□ 函数空间 $\mathcal{H}$

✓ 神经网络

$$\min_{h \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n \ell(h(\mathbf{x}_i), y_i)$$



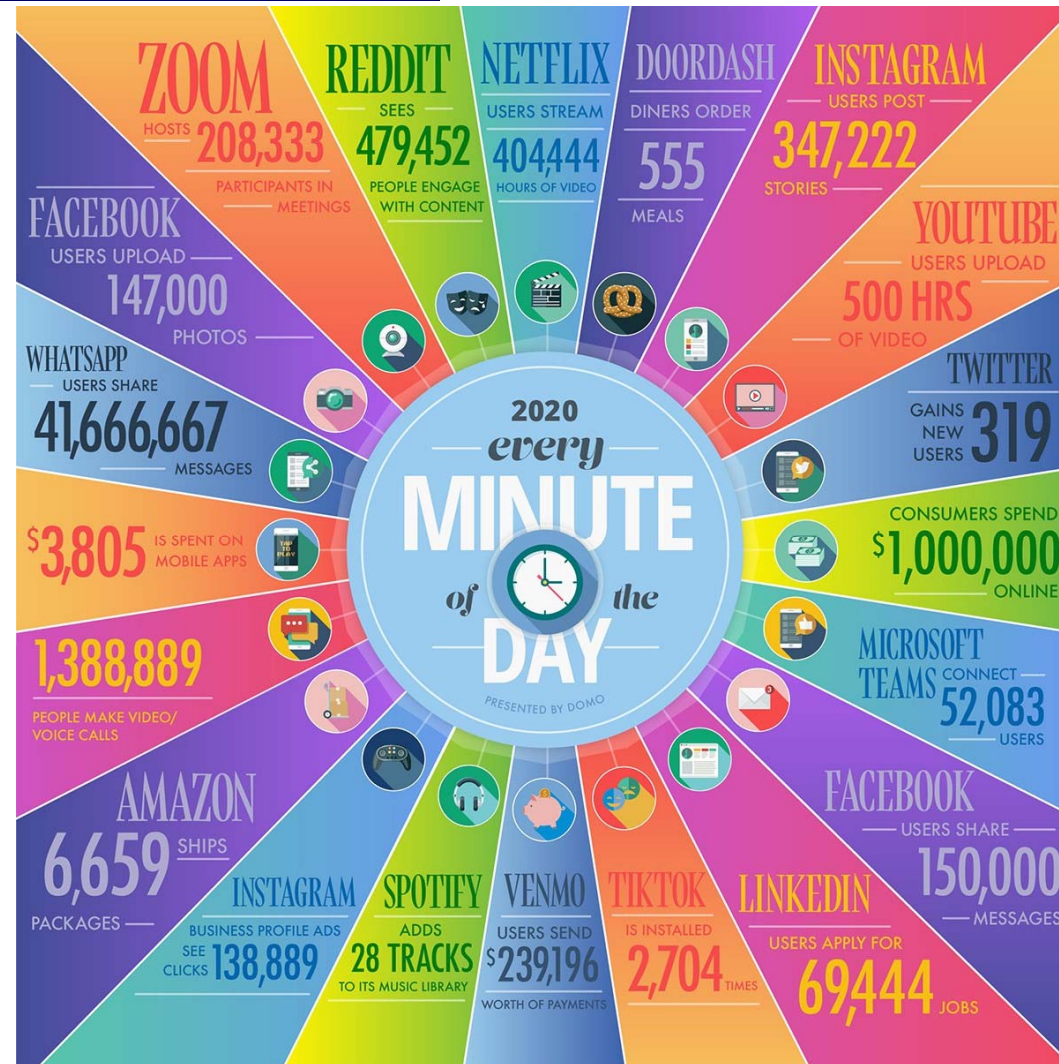
大数据



Sources: McKinsey Global Institute, Twitter, Cisco, Gartner, EMC, SAS, IBM, MEPTec, QAS

快速增长

- Facebook
 - 14.7万图片
- Amazon
 - 6659个包裹
- Zoom
 - 20.8万用户开会
- YouTube
 - 500小时的视频



大数据带来的挑战

➤ 经验误差最小化

$$\min_{h \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n \ell(h(\mathbf{x}_i), y_i) + \phi(h)$$

□ n 是训练样本的数量

□ 训练复杂度: $O(n)$ 、 $O(n^2)$

➤ 大数据“规模大”

□ 训练过程计算量大

➤ 大数据“增长快”

□ 频繁需要重新训练

在线学习

目录

- 研究背景
 - 机器学习
 - 大数据
- 在线学习
 - 遗憾
 - 在线凸优化
- 动态环境
 - 动态遗憾
 - 自适应算法
- 总结

在线学习

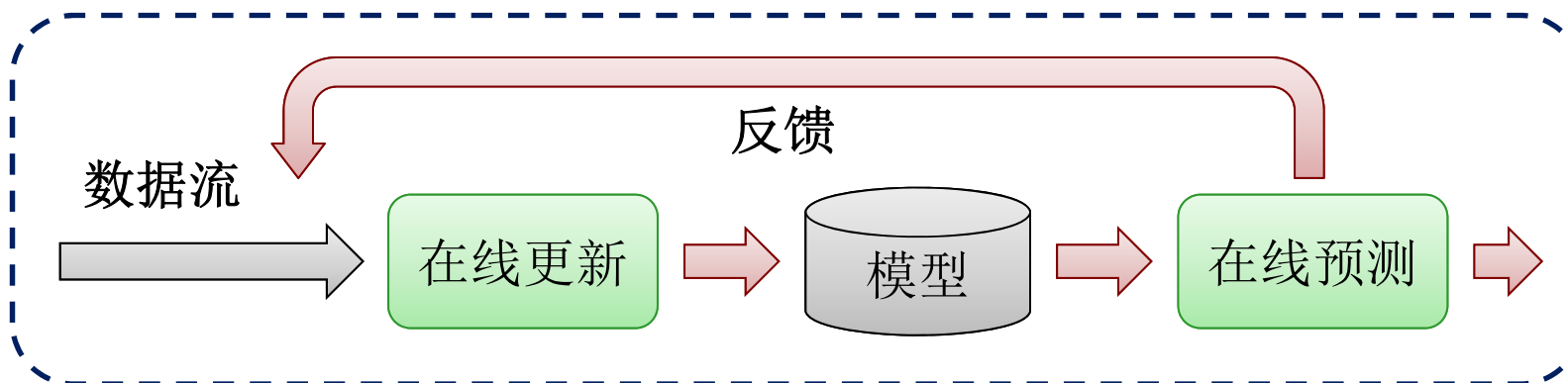
➤ 定义 [Shalev-Shwartz, 2011]

Online learning is the process of answering a sequence of questions given (**maybe partial**) knowledge of answers to previous questions and possibly additional information.

❑ 完全信息在线学习

❑ 赌博机在线学习

➤ 学习流程



数学定义

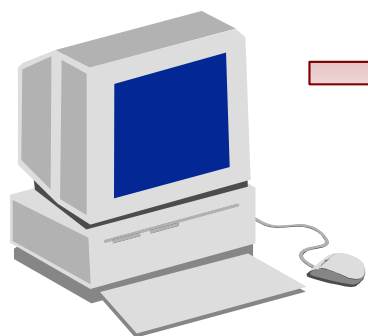
➤ 在线学习 [Cesa-Bianchi and Lugosi, 2006]

for $t = 1, 2, \dots, T$ **do**

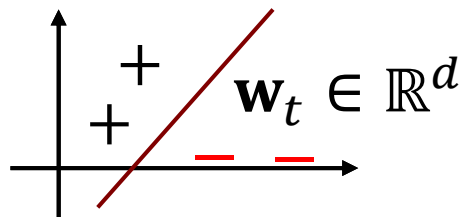
1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

end for



分类器



样本 $(\mathbf{x}_t, y_t) \in \mathbb{R}^d \times \{\pm 1\}$

损失 $f_t(\mathbf{w}) = \max(1 - y_t \mathbf{w}^\top \mathbf{x}_t, 0)$



数学定义

➤ 在线学习 [Cesa-Bianchi and Lugosi, 2006]

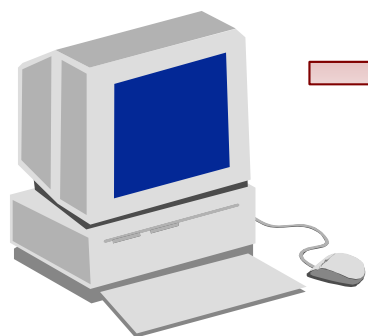
for $t = 1, 2, \dots, T$ **do**

1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

2. 学习器遭受损失 $f_t(\mathbf{w}_t)$ ，并更新 $\mathbf{w}_t$

end for



遭受损失 $f_t(\mathbf{w}_t)$ ，更新 $\mathbf{w}_t$



数学定义

➤ 在线学习 [Cesa-Bianchi and Lugosi, 2006]

for $t = 1, 2, \dots, T$ **do**

1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

2. 学习器遭受损失 $f_t(\mathbf{w}_t)$ ，并更新 $\mathbf{w}_t$

end for

➤ 累积损失

$$\sum_{t=1}^T f_t(\mathbf{w}_t)$$

➤ 最小化累积损失

数学定义

➤ 在线学习 [Cesa-Bianchi and Lugosi, 2006]

for $t = 1, 2, \dots, T$ **do**

1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

2. 学习器遭受损失 $f_t(\mathbf{w}_t)$ ，并更新 $\mathbf{w}_t$

end for

➤ 遗憾 (Regret)

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w})$$

批量学习

➤ 最小化遗憾

数学定义

➤ 在线学习 [Cesa-Bianchi and Lugosi, 2006]

for $t = 1, 2, \dots, T$ **do**

1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

2. 学习器遭受损失 $f_t(\mathbf{w}_t)$ ，并更新 $\mathbf{w}_t$

end for

➤ 遗憾 (Regret)

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w}) = o(T)$$

➤ 最小化遗憾

发展历史

- 感知机（Perceptron） [Rosenblatt, 1958]
 - 在线分类
- 基于专家建议的在线学习 [Littlestone and Warmuth, 1989]
 - 存在 K 个离散专家
 - 集成专家的预测
- 在线凸优化 [Zinkevich, 2003]
 - 损失函数为凸、决策集为凸
 - 在线分类，例如在线SVM
 - 在线回归，例如在线最小二乘

在线凸优化

➤ 在线学习

for $t = 1, 2, \dots, T$ **do**

1. 学习器选择决策 $\mathbf{w}_t \in \mathcal{W}$

同时，对手选择函数 $f_t(\cdot)$

2. 学习器遭受损失 $f_t(\mathbf{w}_t)$ ，并更新 $\mathbf{w}_t$

end for

□ $f_t(\cdot)$ 是凸函数、 $\mathcal{W}$ 为凸集合

➤ 在线梯度下降 [Zinkevich, 2003]

□ $\eta_t \geq 0$ 是步长

□ $\nabla f_t(\mathbf{w}_t)$ 为梯度

$$\mathbf{w}'_{t+1} = \mathbf{w}_t - \eta_t \nabla f_t(\mathbf{w}_t)$$

$$\mathbf{w}_{t+1} = \Pi_{\mathcal{W}}[\mathbf{w}'_{t+1}]$$

➤ 凸函数 [Zinkevich, 2003]

□ 在线梯度下降, $\eta_t = O(1/\sqrt{t})$

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w}) = O(\sqrt{T})$$

➤ 强凸函数 [Hazan et al, 2007]

□ 在线梯度下降, $\eta_t = O(1/[\lambda t])$

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w}) = O(\log T)$$

➤ 指数凹函数 [Hazan et al, 2007]

□ 在线牛顿法 $\text{Regret} = O(d \log T)$

目录

- 研究背景
 - 机器学习
 - 大数据
- 在线学习
 - 遗憾
 - 在线凸优化
- 动态环境
 - 动态遗憾
 - 自适应算法
- 总结

动态环境

➤ 遗憾

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w})$$

动态环境

➤ 遗憾→静态遗憾

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \min_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w}) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{w}^*)$$

□ $\mathbf{w}^* = \operatorname{argmin}_{\mathbf{w} \in \mathcal{W}} \sum_{t=1}^T f_t(\mathbf{w})$ 为批量学习最优解

□ 局限：假设最优模型固定不变

➤ 动态环境下的在线学习

□ 最优模型漂移

1. 推荐系统：用户兴趣
2. 金融理财：投资策略



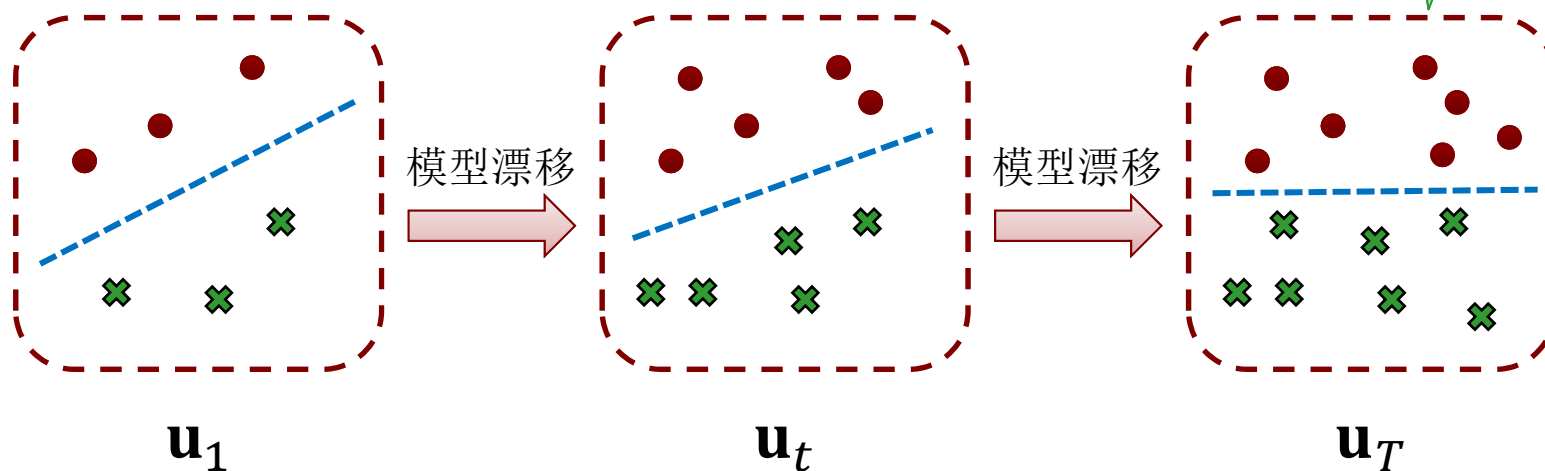
动态遗憾

➤ 动态遗憾（Dynamic Regret） [Zinkevich, 2003]

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{u}_t)$$

□ $\mathbf{u}_1, \dots, \mathbf{u}_T$ 为动态变化的模型序列

比较序列未知!



➤ 动态遗憾 (Dynamic Regret) [Zinkevich, 2003]

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{u}_t)$$

□ $\mathbf{u}_1, \dots, \mathbf{u}_T$ 为动态变化的模型序列

➤ 最差动态遗憾 [Besbes et al., 2015; Mokhtari et al., 2016; Yang et al., 2016; Zhang et al., 2017]

$$R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{w}_t^*) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T \min_{\mathbf{w} \in \mathcal{W}} f_t(\mathbf{w})$$

□ $\mathbf{w}_t^* = \operatorname{argmin}_{\mathbf{w} \in \mathcal{W}} f_t(\mathbf{w})$ 为瞬时最优解

最差动态遗憾 (1)

➤ 函数变化 [Besbes et al., 2015]

$$V_T = \sum_{t=1}^T \sup_{\mathbf{w} \in \mathcal{W}} |f_{t+1}(\mathbf{w}) - f_t(\mathbf{w})| \leq B_T$$

□ 重新启动的在线梯度下降

$$R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = \begin{cases} O(B_T^{1/3} T^{2/3}) & \text{凸函数} \\ O(\log T \sqrt{B_T T}) & \text{强凸函数} \end{cases}$$

□ 当 $B_T = o(T)$ 时, $R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = o(T)$

□ 局限: 算法需要知道先验知识 B_T

最差动态遗憾 (2)

➤ 路径长度

$$P_T^* = \sum_{t=2}^T \|\mathbf{w}_t^* - \mathbf{w}_{t-1}^*\|$$

$$R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = \begin{cases} O(P_T^*) & \text{强凸\&平滑 [Mokhtari et al., 2016]} \\ O(P_T^*) & \text{凸\&平滑\&梯度为0 [Yang et al., 2016]} \\ O(\sqrt{TB_T}) & \text{凸\&} P_T^* \leq B_T \end{cases}$$

➤ 路径平方长度 [Zhang et al., 2017]

□ 强凸且平滑函数

□ 半强凸且平滑函数

□ 自和谐函数

$$S_T^* = \sum_{t=2}^T \|\mathbf{w}_t^* - \mathbf{w}_{t-1}^*\|^2$$

$$R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = O(\min(P_T^*, S_T^*))$$

过拟合静态环境!

➤ 动态遗憾 (Dynamic Regret) [Zinkevich, 2003]

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{u}_t)$$

□ $\mathbf{u}_1, \dots, \mathbf{u}_T$ 为动态变化的模型序列

1. 静态遗憾

$$\text{Regret} = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{w}^*)$$

2. 最差动态遗憾

$$R(\mathbf{w}_1^*, \dots, \mathbf{w}_T^*) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{w}_t^*)$$

➤ 动态遗憾（Dynamic Regret） [Zinkevich, 2003]

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = \sum_{t=1}^T f_t(\mathbf{w}_t) - \sum_{t=1}^T f_t(\mathbf{u}_t)$$

□ $\mathbf{u}_1, \dots, \mathbf{u}_T$ 为动态变化的模型序列

□ 在线梯度下降

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = O\left(\sqrt{T}(1 + P_T)\right)$$

✓ 其中

$$P_T = \sum_{t=2}^T \|\mathbf{u}_t - \mathbf{u}_{t-1}\|$$

□ 优势：当比较序列变化缓慢，动态遗憾自动变小

我们的贡献

➤ 动态遗憾（Dynamic Regret） [Zinkevich, 2003]

□ 在线梯度下降

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = O\left(\sqrt{T}(1 + P_T)\right)$$

其中 $P_T = \sum_{t=2}^T \|\mathbf{u}_t - \mathbf{u}_{t-1}\|$ 为路径长度

➤ 我们的贡献 [Zhang et al, 2018]

□ 首次揭示了动态遗憾的下界

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = \Omega\left(\sqrt{T(1 + P_T)}\right)$$

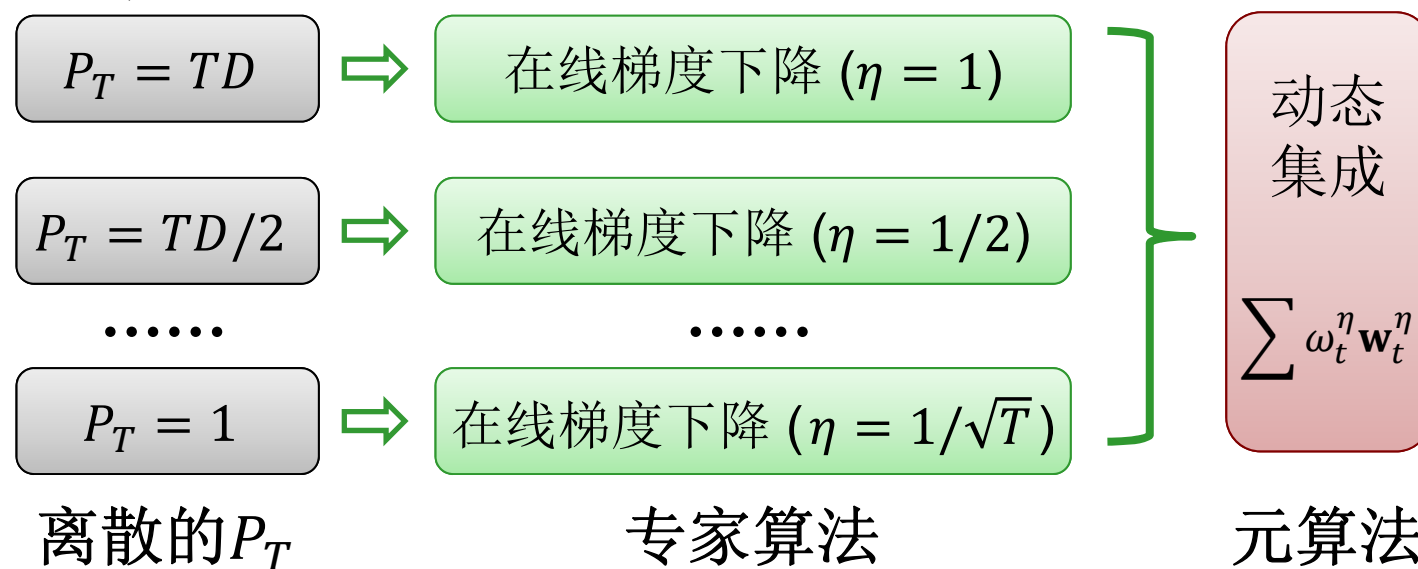
□ 设计了最优的在线学习算法——Ader

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = O\left(\sqrt{T(1 + P_T)}\right)$$

➤ 核心思想

1. 离散化路径长度 $P_T = \sum_{t=2}^T \|\mathbf{u}_t - \mathbf{u}_{t-1}\| \in [0, TD]$
2. 为 P_T 的每一个离散值设计最优的子算法，即专家
3. 通过元算法动态集成所有专家算法的输出

➤ 算法流程



问题相关的理论

➤ Ader [Zhang et al, 2018]

$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = O\left(\sqrt{T(1 + P_T)}\right)$$

□ T 为迭代次数、 $P_T = \sum_{t=2}^T \|\mathbf{u}_t - \mathbf{u}_{t-1}\|$ 为路径长度

□ 局限：没有反映问题本身的特点

➤ Sword_{best} [Zhao et al, 2020]

□ 利用平滑性质

1. 问题容易： F_T 小
2. 函数变化慢： G_T 小

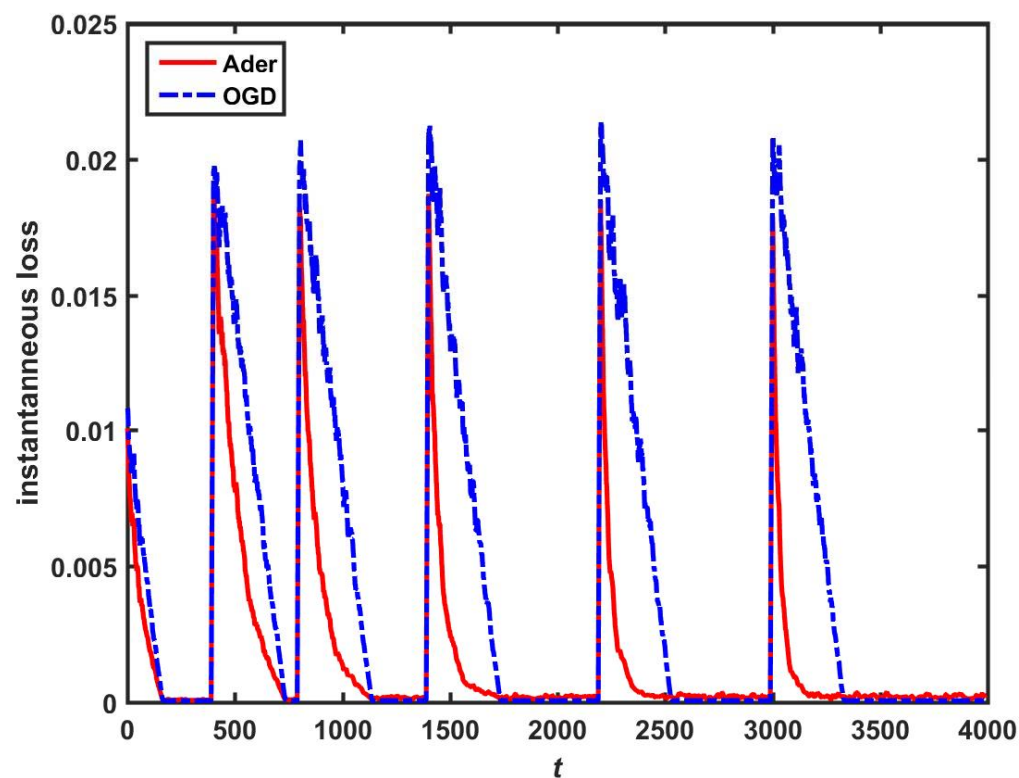
$$R(\mathbf{u}_1, \dots, \mathbf{u}_T) = O\left(\sqrt{(1 + P_T + \min(F_T, G_T))(1 + P_T)}\right)$$

$$F_T = \sum_{t=1}^T f_t(\mathbf{u}_t), \quad G_T = \sum_{t=2}^T \sup_{\mathbf{w} \in \mathcal{W}} \|\nabla f_t(\mathbf{w}) - \nabla f_{t-1}(\mathbf{w})\|_2^2$$

实验 (1)

➤ 凸损失函数 $f_t(\mathbf{w}) = |\mathbf{w}^\top \mathbf{x}_t - y_t|$

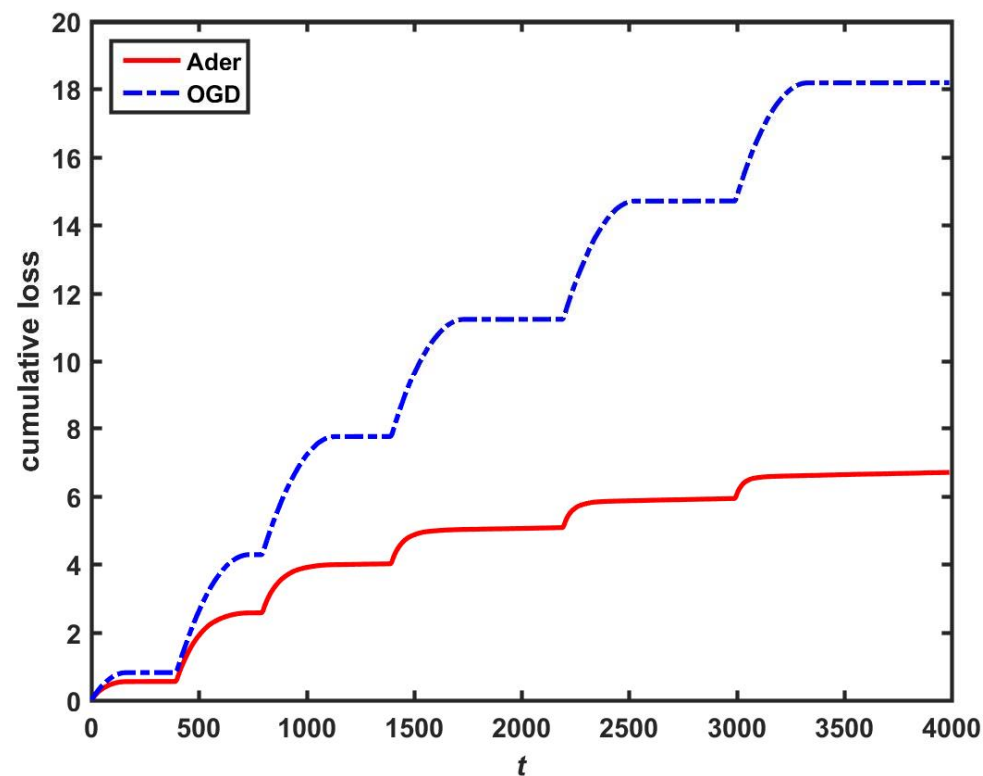
□ 瞬时损失



实验 (2)

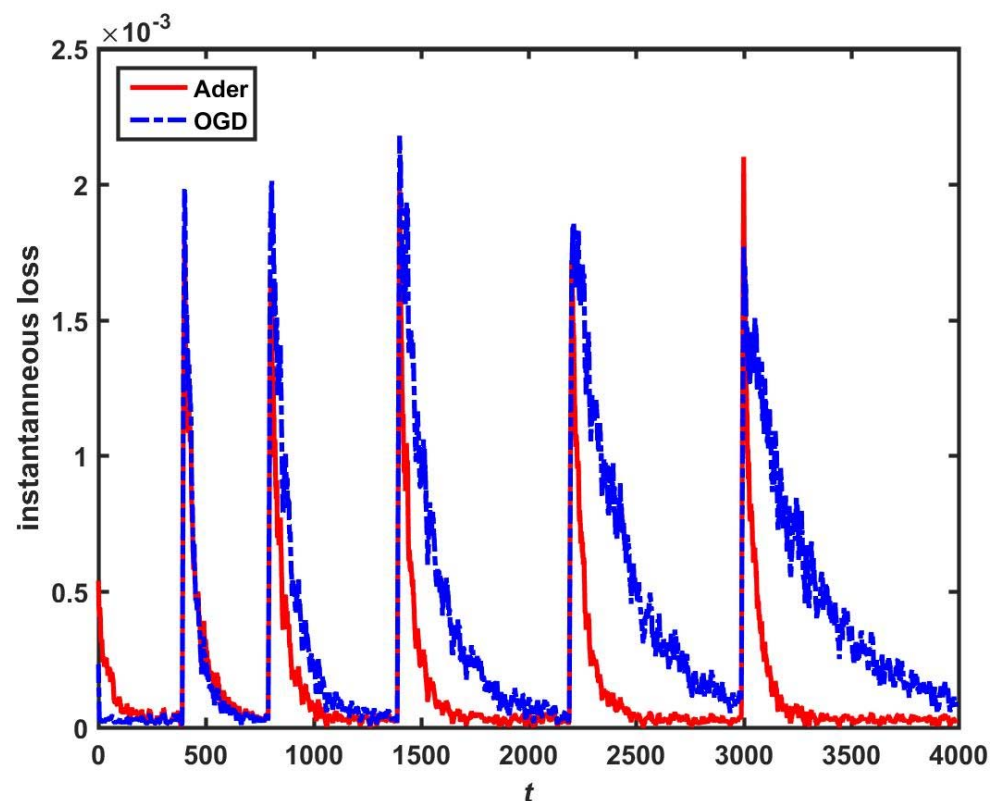
➤ 凸损失函数 $f_t(\mathbf{w}) = |\mathbf{w}^\top \mathbf{x}_t - y_t|$

□ 累计损失



LAMDA
Learning And Mining from Data
<http://lamda.nju.edu.cn>

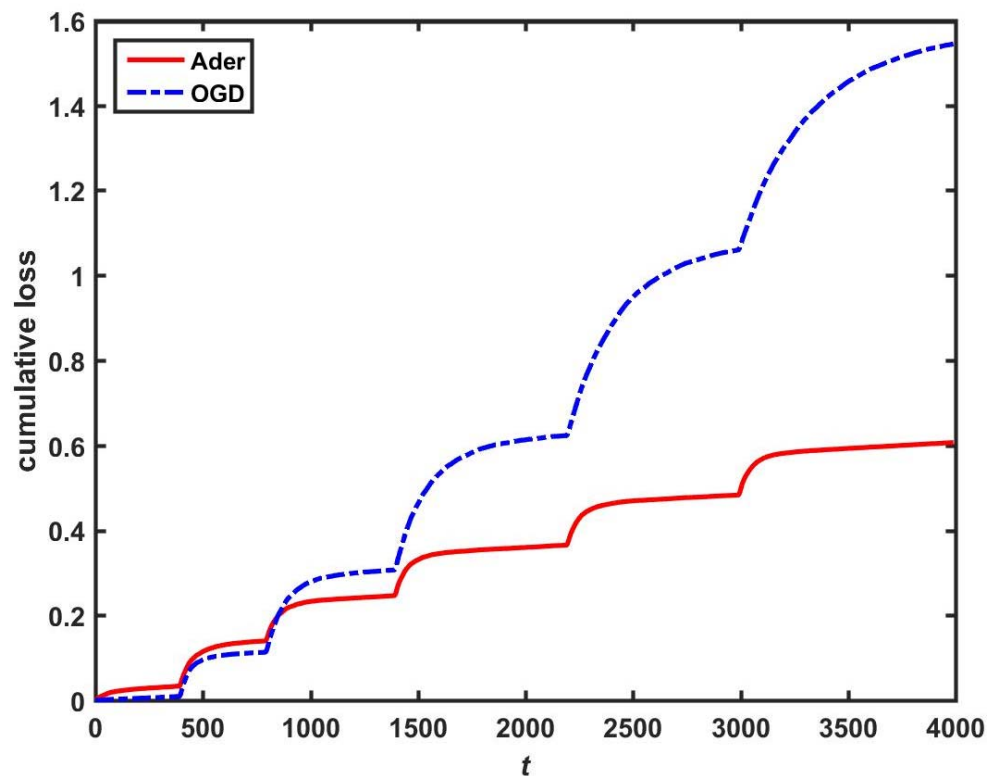
□ 瞬时损失



实验 (4)

➤ 强凸损失函数 $f_t(\mathbf{w}) = (\mathbf{w}^\top \mathbf{x}_t - y_t)^2 + \frac{\lambda}{2} \|\mathbf{w}\|^2$

□ 累计损失



目录

- 研究背景
 - 机器学习
 - 大数据
- 在线学习
 - 遗憾
 - 在线凸优化
- 动态环境
 - 动态遗憾
 - 自适应算法
- 总结

总结与展望

➤ 动态环境下的在线学习

- 新的度量准则：动态遗憾

- 首个下界 $\Omega\left(\sqrt{T(1 + P_T)}\right)$ 、最优算法 Ader [Zhang et al, 2018]

- 问题相关的算法 $\text{Sword}_{\text{best}}$ [Zhao et al, 2020]

➤ 未来工作

- 利用函数额外的性质提升动态遗憾

 - ✓ 强凸、指数凹等

- Dynamic Regret with Switching Cost [Chen et al., 2018]

$$\sum_{t=1}^T (f_t(\mathbf{w}_t) + \|\mathbf{w}_t - \mathbf{w}_{t-1}\|) - \sum_{t=1}^T (f_t(\mathbf{u}_t) + \|\mathbf{u}_t - \mathbf{u}_{t-1}\|)$$

参考文献 (1)

谢谢大家！

- Shalev-Shwartz, S. (2011)
Online learning and online convex optimization.
Foundations and Trends in Machine Learning, 4(2):107–194.
- Cesa-Bianchi, N. and Lugosi, G. (2006)
Prediction, Learning, and Games.
Cambridge University Press.
- Rosenblatt, F. (1958)
The perceptron: a probabilistic model for information storage and organization in the brain.
Psychological Review, 65:386–407.
- Littlestone, N. and Warmuth, M. K. (1989)
The weighted majority algorithm.
In Proceedings of the 30th Annual IEEE Symposium on Foundations of Computer Science, pages 256–261.

参考文献（2）

■ Zinkevich, M. (2003)

Online convex programming and generalized infinitesimal gradient ascent.
In Proceedings of the 20th International Conference on Machine Learning (ICML).

■ Hazan, E., Agarwal, A., and Kale, S. (2007)

Logarithmic regret algorithms for online convex optimization.
Machine Learning, 69:169–192.

■ Besbes, O., Gur, Y., and Zeevi, A. (2015)

Non-stationary stochastic optimization.
Operations Research, 63(5):1227–1244.

■ Mokhtari, A., Shahrampour, S., Jadbabaie, A., and Ribeiro, A. (2016)

Online optimization in dynamic environments: Improved regret rates for strongly convex problems.
In Proceedings of the 55th IEEE Conference on Decision and Control, pages 7195–7201.

■ Chen, N., Goel, G., and Wierman, A. (2018)

Smoothed online convex optimization in high dimensions via online balanced descent.
In Proceedings of the 31st Conference on Learning Theory, pages 1574–1594.

参考文献 (3)

- Yang, T., **Zhang, L.**, Jin, R., and Yi, J. (2016)
Tracking slowly moving clairvoyant: Optimal dynamic regret of online learning with true and noisy gradient.
In Proceedings of the 33rd International Conference on Machine Learning, pages 449–457.
- **Zhang, L.**, Yang, T., Yi, J., Jin, R., and Zhou, Z.-H. (2017)
Improved dynamic regret for non-degenerate functions.
In Advances in Neural Information Processing Systems 30, pages 732–741.
- **Zhang, L.**, Yang, T., Jin, R., and Zhou, Z.-H. (2018)
Dynamic regret of strongly adaptive methods.
In Proceedings of the 35th International Conference on Machine Learning.
- Zhao, P., Zhang, Y.-J., **Zhang, L.**, and Zhou, Z.-H. (2020).
Dynamic regret of convex and smooth functions.
In Advances in Neural Information Processing Systems 33.