

Online optimization and machine learning: Applications to resource allocation problems in wireless networks

E. Veronica Belmega^{1,2}

¹Université Gustave Eiffel, CNRS, LIGM

²ETIS, CY Cergy-Paris University, ENSEA, CNRS

January 6th, 2024



Collaborators and students

P. Mertikopoulos (CNRS, LIG)
I. Fijalkow (ENSEA, ETIS)
A. Savard (IMT Nord Europe)
R. Negrel (ESIEE Paris, LIGM)
M. Debbah (Khalifa Univ. of Sci. and Technol.)

A. Marcastel (PhD)
H. El Hassani (PhD)
I. Chafaa (PhD)
O. Bilenne (PostDoc)

Supported by

Chair IoT Orange with Univ. Cergy Pontoise foundation, ANR ORACLESS, ANR ELIOT, GdR ISIS, ENSEA, . . .

Based on



E.V. Belmega, P. Mertikopoulos, and R. Negrel, “Online convex optimization in wireless networks and beyond: The feedback - performance trade-off”, *invited paper at RAWNET intl. workshop in conjunction with WiOpt*, Sep. 2022.



I. Chafaa, R. Negrel, E.V. Belmega, and M. Debbah, “Self-supervised deep learning for mmWave beam steering exploiting sub-6 GHz channels”, *IEEE Trans. on Wireless Commun.*, Mar. 2022.

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- First order (gradient) feedback
- Zeroth-order (scalar) feedback
- 1-bit feedback

II. Deep learning vs. online policies

- mmWave beamforming from sub-6GHz channels

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- First order (gradient) feedback
- Zeroth-order (scalar) feedback
- 1-bit feedback

II. Deep learning vs. online policies

- mmWave beamforming from sub-6GHz channels

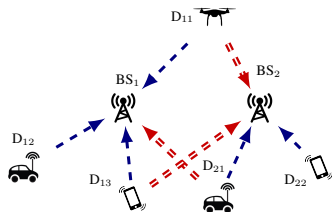
Distributed and dynamic wireless networks

B5G, 6G, IoT

[Saad'19, Tataria'21]

- **Characteristics and requirements:** dense, heterogeneous, autonomous, decentralized, energy-efficient, . . .
- **Challenges**
 - Highly *mobile* radios & networks
 - Unpredictable* and *arbitrary* connectivity patterns
 - Limited network/channel knowledge (*potentially outdated*)
 - Limited processing capabilities
- Static or stochastic (optimal or Nash) solutions: not suitable
⇒ **dynamic, non-equilibrium** solutions

Obj: Design **energy-efficient** resource allocation policies coping with
arbitrary network dynamics and **scarce** feedback



M transmitting devices, N receivers, S orthogonal bands

$$\text{Shannon rate of an arbitrary device } R_t(\mathbf{p}(t)) = \frac{1}{S} \sum_{s=1}^S \log(1 + \underbrace{w_s(t)p_s(t)}^{SNR_s})$$

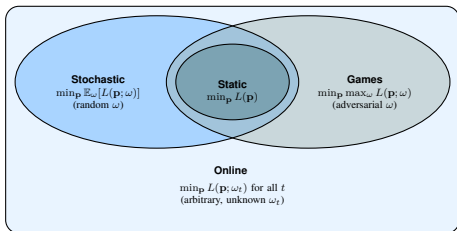
$\mathbf{p}(t)$ - power allocation vector

$$\text{Effective channel gain } w_s(t) = \frac{|h_s(t)|^2}{\sigma^2 + \underbrace{\sum_j |h_s^j(t)|^2 p_s^j(t)}_{\text{interference}}}$$

$$\begin{aligned} \text{minimize} \quad & L_t(\mathbf{p}(t)) \triangleq \underbrace{\sum_{s=1}^S p_s(t)}_{\text{power consumption}} + \underbrace{\lambda [R_{\min} - R_t(\mathbf{p}(t))]^+}_{\text{minimum rate penalty}} \\ \text{s.t.} \quad & p_s(t) \geq 0, \forall s, \sum_{s=1}^S p_s(t) \leq P_{\max} \end{aligned}$$

Rk: **Arbitrary time-varying objective**, unknown at the transmission instant

- How to optimize an unknown $L(\mathbf{p}, \omega_t)$ objective at the decision time t ?



Online iterative process

For $t = 1$ to T

- Choose policy $\mathbf{p}(t)$
- Incur loss $L(\mathbf{p}(t), \omega_t)$
- Receive **feedback**
- Update:** $\mathbf{p}(t+1) \leftarrow \mathbf{p}(t)$ based on feedback

- Goal:** develop *efficient* online processes based on strictly causal feedback
- Link with machine learning: **multi-armed bandits** (MABs) reinforcement learning
- Assumptions:** on the objective function $L(\mathbf{p}, \omega_t)$ w.r.t. $\mathbf{p}$ but **not** on the underlying dynamics of ω_t

- Ideal benchmark

$$\forall t, \quad L_t(\mathbf{p}(t)) - \underbrace{\min_{\mathbf{q} \in \mathcal{P}} L_t(\mathbf{q})}_{\text{dynamic optimal loss}}$$

Too ambitious under the worse-case assumption of an **arbitrarily time-varying network**!

- Less ambitious benchmark: **regret** [Hannan'57]

$$\underbrace{Reg_T}_{\text{overall regret}} \triangleq \sum_{t=1}^T L_t(\mathbf{p}(t)) - \underbrace{\min_{\mathbf{q} \in \mathcal{P}} \sum_{t=1}^T L_t(\mathbf{q})}_{\text{optimal overall loss of fixed policies}}$$

- **Rk**: both benchmarks require knowledge of the future!

Objective: design online policies $\mathbf{p}(t)$ that lead to **no regret**:

$$\limsup_{T \rightarrow \infty} \frac{1}{T} Reg_T \leq 0.$$

No-regret policies

Static convex optimization $L_t(\mathbf{p}) = L(\mathbf{p}), \forall t$

$\mathbf{p}(t)$ no-regret policy $\Rightarrow \bar{\mathbf{p}}(T) = \frac{1}{T} \sum_{t=1}^T \mathbf{p}(t)$ converges to the optimal solution

Stochastic convex optimization $L_t(\mathbf{p}) = L(\mathbf{p}, \omega_t), \forall t$

$\mathbf{p}(t)$ no-regret policy $\Rightarrow \bar{\mathbf{p}}(T) = \frac{1}{T} \sum_{t=1}^T \mathbf{p}(t)$ converges to the optimal solution

Non-cooperative games

[Viossat'13, Mertikopoulos'19]

A no-regret policy (cumulative) converges to the Nash equilibrium in
two-player (discrete or convex) **zero-sum** games, potential (discrete or convex) games, ...

Adversarial (non stochastic) MABs

- S arms of slot machines
- **Variable:** probability distribution $\mathbf{p}(t)$

$$\mathbf{p}(t) \in \mathcal{S} \triangleq \left\{ \mathbf{p} \in \mathbb{R}_+^S : \sum_{s=1}^S p_s = 1 \right\}$$

- **Objective:** maximize the linear expected gains
 $U_t(\mathbf{p}(t)) = \mathbf{r}(t)^T \mathbf{p}(t)$
 $\mathbf{r}(t)$ - vector of arms' rewards at time t
- **Optimal policy:** exponential/multiplicative weights
 [Auer'02]

Our online power allocation problem

- S channels or frequency subcarriers
- **Variable:** power allocation vector $\mathbf{p}(t)$

$$\mathbf{p}(t) \in \mathcal{P} \triangleq \left\{ \mathbf{p} \in \mathbb{R}_+^S : \sum_{s=1}^S p_s \leq P_{\max} \right\}$$

- **Objective:** minimize the convex loss $L_t(\mathbf{p}(t))$ - tradeoff
 power consumption vs. rate
- **Hint:** adapted version of exponential weights

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- **First order (gradient) feedback**
- Zeroth-order (scalar) feedback
- 1-bit feedback

II. Deep learning vs. online policies

- mmWave beamforming from sub-6GHz channels

Gradient-based online policies

Low complexity, parallel computing

Online Gradient Descent (OGD)

[Zinkevich'03]

- Gradient feedback: $\mathbf{v}(t) = -\nabla L_t(\mathbf{p}(t))$
- Update: Euclidean projection of $\mathbf{p}(t) + \gamma \mathbf{v}(t)$

$$\begin{aligned}\mathbf{p}(t+1) &\triangleq \arg \min_{\mathbf{q} \in \mathcal{P}} \frac{1}{2} \|\mathbf{p}(t) + \gamma \mathbf{v}(t) - \mathbf{q}\|^2 \\ &= \arg \min_{\mathbf{q} \in \mathcal{P}} \left\{ \gamma \mathbf{v}(t)^T (\mathbf{p}(t) - \mathbf{q}) + \frac{1}{2} \|\mathbf{p}(t) - \mathbf{q}\|^2 \right\}\end{aligned}$$

Based on the mirror descent [Nemirovski'83]

Online Mirror Descent (OMD)

[Shalev-Shwartz'07]

- Gradient feedback: $\mathbf{v}(t) = -\nabla L_t(\mathbf{p}(t))$
- Update: Mirror-mapping on $\mathbf{p}(t) + \gamma \mathbf{v}(t)$

$$\mathbf{p}(t+1) \triangleq \arg \min_{\mathbf{q} \in \mathcal{P}} \left\{ \gamma \mathbf{v}(t)^T (\mathbf{p}(t) - \mathbf{q}) + D_h(\mathbf{q}, \mathbf{p}(t)) \right\}$$

$D_h(\mathbf{q}, \mathbf{p}) \triangleq h(\mathbf{q}) - h(\mathbf{p}) - \nabla h(\mathbf{p})^T (\mathbf{q} - \mathbf{p})$ the Bregman divergence of regularizer $h(\cdot)$

Online Mirror Descent (OMD)

- Particular cases

- ▶ OGD: Euclidean regularizer $h(\mathbf{p}) = \frac{1}{2} \|\mathbf{p}\|^2$, $D_h(\mathbf{q}, \mathbf{p}) = \frac{1}{2} \|\mathbf{q} - \mathbf{p}\|^2$

- ▶ Exponential weights for MAB:

entropic regularizer $h(\mathbf{p}) = \sum_{s=1}^S p_s \log p_s$, Kullback-Leibler $D_h(\mathbf{q}, \mathbf{p}) = \sum_{s=1}^S q_s \log \frac{q_s}{p_s}$

- MAB (simplex) regret performance: **scalability**

- ▶ OGD: $Reg_T = \mathcal{O}(\sqrt{TS})$ [Zinkevich'03]

- ▶ Exponential weights: $Reg_T = \mathcal{O}(\sqrt{T \log S})$ [Auer'02]

OMD allows to design algorithms that are **tailored to the feasible set** compared with OGD.

OXL based on gradient feedback

Main idea: we exploit the similarity of $\mathcal{P}$ with the simplex and the entropic regularizer

Online exponential learning (OXL)

- Gradient feedback: $\mathbf{v}(t) = -\nabla L_t(\mathbf{p}(t))$
- Update:

$$\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \mathbf{v}(t) \quad \text{cumulative gradient score}$$

$$p_s(t+1) = P_{\max} \frac{\exp(y_s(t+1))}{1 + \sum_{i=1}^S \exp(y_i(t+1))}, \quad \forall s \quad \text{exponential mapping}$$

- Similar to the exponential weights for MAB
- Closed-form update as opposed to OGD, which requires an Euclidean projection
- We proved: **no-regret** property, $Reg_T \leq P_{\max} \sqrt{2V \log(1+S)T}$, $V = \max \|v(t)\|^2$

Extends to the imperfect gradient feedback case: $\tilde{\mathbf{v}}(t) = \mathbf{v}(t) + \mathbf{error}$

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- First order (gradient) feedback
- Zeroth-order (scalar) feedback
- 1-bit feedback

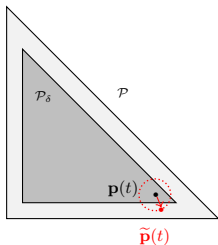
II. Deep learning vs. online policies

- mmWave beamforming from sub-6GHz channels

Can the feedback be reduced to a scalar?

- **Received feedback:** value of incurred loss $L_t(\mathbf{p})$ as in MABs with bandit feedback
- Stochastic gradient approximation: $\tilde{\mathbf{v}}(t) = \frac{S}{\delta} L_t(\mathbf{p} + \delta \mathbf{u}) \mathbf{u}$,
 $\mathbf{u}$ - uniformly drawn over the unit S -dimensional Sphere
- To estimate $\nabla L_t(\mathbf{p}(t))$, transmit at $\tilde{\mathbf{p}}(t) = \mathbf{p}(t) + \delta \mathbf{u} \rightarrow \tilde{\mathbf{p}}(t)$ may fall outside $\mathcal{P}$!
- Our solution: **shrink the feasible set** s.t. $\forall \mathbf{p}(t) \in \mathcal{P}_\delta \subset \mathcal{P} \Rightarrow \tilde{\mathbf{p}}(t) \in \mathcal{P}$

[Spall'97, Flaxman'05]



$$\mathcal{P}_\delta \triangleq \left\{ \mathbf{p} \in \mathbb{R}_+^S : p_s \geq \delta, \sum_{s=1}^S p_s \leq P_{\max} - \sqrt{S}\delta \right\}$$

Our modified OXL₀OXL₀ algorithm

- Transmit at $\tilde{\mathbf{p}}(t) = \mathbf{p}(t) + \delta \mathbf{u}$
- Receive **scalar feedback**: $L_t(\tilde{\mathbf{p}}(t))$
- Gradient estimation: $\tilde{\mathbf{v}} = -\frac{S}{\delta} L_t(\tilde{\mathbf{p}}(t)) \mathbf{u}$

• Update:

$$\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \tilde{\mathbf{v}}(t)$$

cumulative score

$$p_s(t+1) = \delta + P_{\max}(1 - C_\delta) \frac{\exp(y_s(t+1))}{1 + \sum_{i=1}^S \exp(y_i(t+1))}, \quad \forall s$$

modified exponential mapping

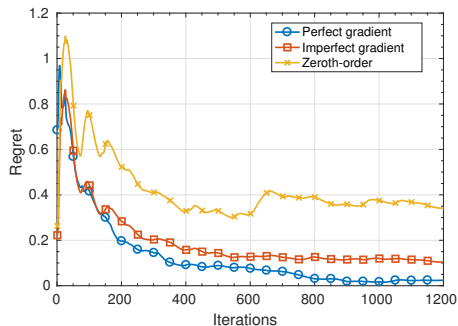
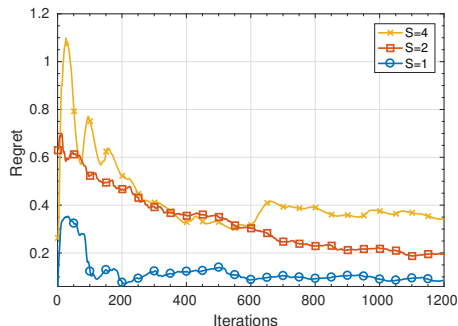
adapted to the shrunk set $\mathcal{P}_\delta$

$$C_\delta = \frac{\delta}{P_{\max}} (S + \sqrt{S})$$

- We proved: **no-regret property**, $E\text{Reg}_T = \mathcal{O}(T^{3/4})$
- Zeroth-order feedback greatly impacts the decay rate of the average regret

Tradeoff: regret decay rate vs. amount of feedback!

Impact of reducing the feedback

OXL vs. OXL₀OXL₀ and S - problem dimension

Zeroth-order information reduces the decay rate of the average regret.
The decay rate becomes worse by increasing S - the problem dimensionality.

COST-HATA wireless channels, $M = 10$ transmitting devices, $S \in \{1, 2, 4\}$ bands, $N = 1$ receiver, $F_{\max} \in [0.5, 2]$, $R_{\min} \in [0.5, 3]$ bps/Hz, $\lambda \in [0.5, 10]$, average regret over 200 random draws of $\mathbf{u}$

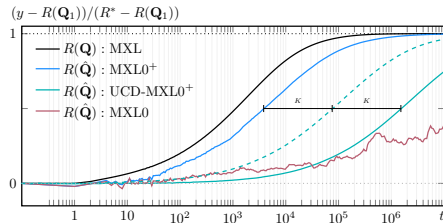
Zeroth-order feedback with callbacks

- **Issues:** slow decay rate $\mathcal{O}(T^{3/4})$, poor scaling with problem dimensionality
 —→ MXL0 performs poorly in MIMO networks
- Motivation: exponential weights yield $\mathcal{O}(\sqrt{T})$ regret in bandit feedback MABs
- Our idea: exploit the current feedback R_t jointly with the past R_{t-1}
 —→ improved two-point gradient estimator
- For **static convex objectives**, we can recover the $\mathcal{O}(\sqrt{T})$ regret as with perfect gradient feedback

K -user MIMO multiple access channel

$$R(\mathbf{Q}) = \log \det \left(\mathbf{I} + \sum_{k=1}^K \mathbf{H}_k \mathbf{Q}_k \mathbf{H}_k^\dagger \right)$$

$$\mathbf{Q}_k \succeq 0, \quad \text{Tr}(\mathbf{Q}_k) \leq R_{\max}$$



Arbitrary static channels, $K = 20$ users, 3 transmit and 16 receive antennas

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- First order (gradient) feedback
- Zeroth-order (scalar) feedback
- 1-bit feedback

II. Deep learning vs. online policies

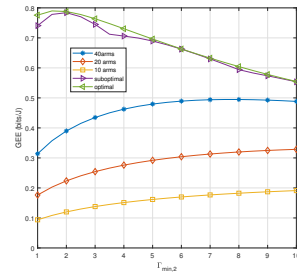
- mmWave beamforming from sub-6GHz channels

Can the feedback be reduced to one bit?

- **Our idea:** exploit MABs with outage (QoS) based rewards and ACK/NACK-type feedback
- Problems: beam alignment in mmWave networks, NOMA with no CSIT/CDIT
- **Issue:** feasible set needs to be quantized
→ performance loss

2-user downlink NOMA

$$\text{GEE}(i, \mathbf{p}) \triangleq \frac{(R_{\min,i} + R_{\min,j})(1 - \mathbb{P}_{\text{out}}(i, \mathbf{p}))}{p_i + p_j + P_c}$$



$$\Gamma_{\min,2} = 2^{R_{\min,2}} - 1, \sigma_k^2 = 0.1, R_{\min,1} = 1 \text{ bpcu}, P_c = 1 \text{ W}, \\ P_{\max} = 100 \text{ W}, T = 5000, 10^3 \text{ channel (Rayleigh) realizations}$$

- **Online optimization and no-regret learning:** a suitable framework for resource allocation problems in dynamic and unpredictable networks
- **Assumptions:** Convex (extends to fractional programs) and Lipschitz objectives $L(\mathbf{x}, \omega_t)$ w.r.t. $\mathbf{x}$, convex feasible sets

⇒ Adaptive online algorithms

Decentralized and reinforcing

Cope with **arbitrary network dynamics**

Rely on **strictly causal** information

Imperfect and **reduced feedback**

Regret-based theoretical **guarantees**

- Tradeoff: performance vs. available feedback
- Applies to semi-definite programming: in **massive MIMO** networks reducing the (matrix) gradient feedback is crucial

- Zeroth-order feedback with callbacks for **online** problems
- One-bit feedback policies: improve the performance, NOMA cope with continuous and discrete variables
- Beyond convex, Lipschitz objectives: more realistic (**non convex**) energy-efficiency measures
- Compare online algorithms with **deep learning** ones

I. Online resource optimization policies: feedback vs. performance

- Preliminaries
- First order (gradient) feedback
- Zeroth-order (scalar) feedback
- 1-bit feedback

II. Deep learning vs. online policies

- mmWave beamforming from sub-6GHz channels

Model-based approaches tradeoff

simple (unrealistic) but tractable problems vs. practical but not tractable problems

⇒ machine (deep) learning has the potential to bridge this gap

Deep learning based on neural networks

[Zappone'19, Tataria'21]

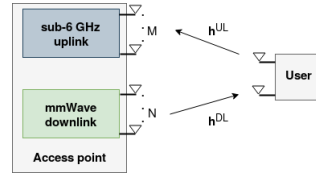
- + Powerful toolbox: *universal approximators* of complex relationships
- + Pervasive in 6G network design, operation, ...
- Large and relevant training datasets
- High computing/processing capabilities

mmWave communications

- Exploit high frequency spectrum
- MIMO systems and **beam alignment** against severe pathloss
- **Challenges:** channel estimation, device mobility and network dynamics

—→ online optimization (MABs) and deep learning
[Hashemi'18, Alrabeiah'20]

- **Idea:** Mapping sub-6GHz channels into mmWave beamforming vectors [Alrabeiah'20]



- Sub-6 GHz channels provide multipath signature, easier to estimate
- $\mathbf{h}^{UL} \rightarrow \mathbf{f}^{DL}$ highly complex and non linear mapping
—→ deep learning

Network architecture and DeepMIMO

Neural network architecture

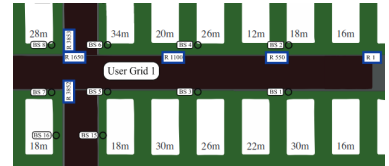
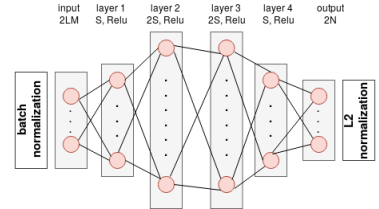
- Input: uplink sub-6 GHz channels
- Output: mmWave beamforming vectors
- **Fully-connected:** structure agnostic

Training

- DeepMIMO dataset: $\left(\{ \mathbf{h}_i[\ell]^{\text{UL}} \}_{\ell=1}^L, \{ \mathbf{h}_i[\ell]^{\text{DL}} \}_{\ell=1}^L \right)_i$
- Custom loss function: $\mathcal{L} = -1/\mathcal{B} \sum_{i=1}^{\mathcal{B}} \mathcal{R}_i$

$$\mathcal{R}_i = \frac{1}{L} \sum_{\ell=1}^L \log_2 \left(1 + \frac{P^{\text{DL}}}{L (\sigma^{\text{DL}})^2} | \mathbf{h}_i^{\text{DL}\dagger}[\ell] \mathbf{f}_i |^2 \right)$$

$\mathbf{f}_i$: predicted beamformer, P^{DL} transmit power, $\mathcal{B}$ size of mini-batch, L subcarriers, $(\sigma^{\text{DL}})^2$ noise variance



DeepMIMO setup [Alkhateeb'19]

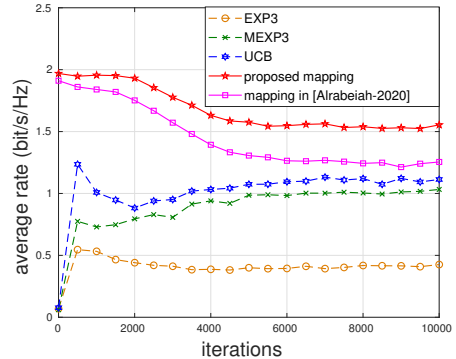
Deep learning or MABs?

MABs

- + No *offline* training, online learning
- + Low complexity $\sim 10^2$ ops/iteration
- + 1-bit feedback (ACK/NACK)
- Quantized beams $\rightarrow$ performance loss
- Trial and error $\rightarrow$ transitory phase

Deep learning

- Offline training on large and relevant datasets, *DeepMIMO* [Alkhateeb'19]
- Higher complexity $\sim 10^6 - 10^7$ ops/iteration
- sub-6 GHz channel information
- + Regression (no quantization)
- + Better running performance (no transitory phase)



MABs: EXP3, MEXP3, UCB

Deep learning [Alrabeiah'20]: classification

Depends on the target application, its characteristics and requirements

I. Chafaa, R. Negrel, E.V. Belmega, and M. Debbah, "Self-supervised deep learning for mmWave beam steering exploiting sub-6 GHz channels", *IEEE Trans. on Wireless Commun.*, 2022.

I. Chafaa, E.V. Belmega, and M. Debbah, "One-bit Feedback exponential learning for beam alignment in mobile mmWave", *IEEE Access*, Oct. 2020.

- Deep learning approaches capture **complex relationships**:
sub-6GHz uplink channels – downlink mmWave beams
- Extension to multi-cell multi-user networks
- **Robustness** to data imperfections
- **Generalization** capability: other wireless settings, problems
- Recurrent networks, reinforcement and online deep learning to capture the temporal network dynamics

AI-enhanced highly mobile and unpredictable IoT networks

- OBJ-1** Design efficient **online optimization** algorithms requiring **low-cost** and energy-efficient communication feedback
- OBJ-2** Design online algorithms maximizing **non convex** energy efficiency exploiting non convex online optimization and deep learning

Sustainable wireless communications: low-energy, low-cost and zero added electromagnetic waves

- OBJ-1** Derive the fundamental **achievable Shannon rates** in multi-user, multi-backscatter/RIS networks
- OBJ-2** Develop **efficient algorithms** that jointly tune the transmit strategy and the backscattering/RIS strategy

More intel on my webpage:

<https://sites.google.com/site/evbelmega>

-  E.V. Belmega, P. Mertikopoulos, and R. Negrel, “Online convex optimization in wireless networks and beyond: The feedback - performance trade-off”, *invited paper at RAWNET intl. workshop in conjunction with WiOpt*, Sep. 2022.
-  O. Bilenne, P. Mertikopoulos, and E.V. Belmega, “Fast Gradient-Free Optimization in Distributed Multi-User MIMO Systems”, *IEEE Trans. on Signal Processing*, Oct. 2020.
-  A. Marcastel, E. V. Belmega, P. Mertikopoulos, and I. Fijalkow, “Gradient-free online resource allocation algorithms for dynamic wireless networks”, *invited paper to IEEE SPAWC*, 2019.
-  A. Marcastel, E.V. Belmega, P. Mertikopoulos, and I. Fijalkow, “Online power optimization in dynamic IoT networks: The impact of feedback scarcity”, *IEEE Trans. on Signal Processing*, Mar. 2019.
-  P. Mertikopoulos, and E.V. Belmega, “Learning to be green: robust energy efficiency maximization in dynamic MIMO-OFDM systems”, *IEEE Journal on Selected Areas in Communication*, Apr. 2016.
-  H. El Hassani, A. Savard, and E.V. Belmega, “Energy-efficient 1-bit feedback NOMA in wireless networks with no CSIT/CDIT”, *IEEE SSP Workshop*, Jun. 2021.
-  H. El Hassani, A. Savard, and E.V. Belmega, “Adaptive NOMA in time-varying wireless networks with no CSIT/CDIT relying on a 1-bit feedback”, *IEEE Wireless Commun. Lett.*, Apr. 2021.
-  I. Chafaa, R. Negrel, E.V. Belmega, and M. Debbah, “Self-supervised deep learning for mmWave beam steering exploiting sub-6 GHz channels”, *IEEE Trans. on Wireless Commun.*, 2022.
-  I. Chafaa, E.V. Belmega, and M. Debbah, “One-bit Feedback exponential learning for beam alignment in mobile mmWave”, *IEEE Access*, Oct. 2020.



T. Chen, S. Barbarossa, X. Wang, G. B. Giannakis, and Z.-L. Zhang, “Learning and management for Internet-of-Things: Accounting for adaptivity and scalability”, Arxiv preprint arxiv:1810.11613, 2018.



J. Hannan, “Approximation to Bayes risk in repeated play”, *Contributions to the Theory of Games, Vol. III*, fPrinceton University Press, vol. 39, pp. 97–139, 1957.



Y. Viossat and A. Zapechelnyuk, “No-regret dynamics and fictitious play”, *Journal of Economic Theory*, vol. 148, no 2, pp. 825–842, 2013.



P. Mertikopoulos, and Z. Zhou, “Learning in games with continuous action sets and unknown payoff functions”, *Mathematical Programming*, vol. 173, no 1–2, pp. 465–507, 2019.



M. Zinkevich, "Online convex programming and generalized infinitesimal gradient ascent, *ICML'03: Proceedings of the 20th International Conference on Machine Learning*, pp. 928–936, 2003.



S. Shalev-Shwartz, "Online learning: Theory, algorithms, and applications," *Ph.D. dissertation, Hebrew University of Jerusalem*, 2007.



A. S. Nemirovski and D. B. Yudin, "Problem Complexity and Method Efficiency in Optimization", *Wiley, New York, NY*, 1983.



P. Auer, N. Cesa-Bianchi, Y. Freund, and R. E. Schapire, “The non stochastic multi armed bandit problem”, *SIAM Journal on Computing*, vol. 32, pp. 48–77, 2002.



J. C. Spall, “A one-measurement form of simultaneous perturbation stochastic approximation”, *Automatica*, vol. 33, no. 1, pp. 109–112, 1997.



A. D. Flaxman, A. T. Kalai, and H. B. McMahan, "Online convex optimization in the bandit setting: gradient descent without a gradient", *SODA'05: Proceedings of the 16th annual ACM-SIAM symposium on discrete algorithms*, 2005, pp. 385–394.

- 
- C. Zhang, P. Patras, and H. Haddadi, “Deep learning in mobile and wireless networking: A survey”, *Arxiv preprint arxiv:1803.04311*, 2018.
- 
- A. Zappone, M. Di Renzo, and M. Debbah, “Wireless networks design in the era of deep learning: Model-based, AI-based, or both?”, *IEEE Transactions on Communications*, vol. 67, no. 10, pp. 7331–7376, 2019.
- 
- E. C. Strinati, S. Barbarossa, J. L. Gonzalez-Jimenez, D. Ktenas, N. Cassiau, and C. Dehos, “6G: The next frontier: From holographic messaging to artificial intelligence using subterahertz and visible light communication”, *IEEE Vehicular Technology Magazine*, vol. 14, no.3, pp.42–50, 2019.
- 
- W. Saad, M. Bennis, and M. Chen, “A vision of 6G wireless systems: Applications, trends, technologies, and open research problems”, *IEEE Network*, vol. 34, no. 3, pp. 134–142, 2019.
- 
- H. Tataria, M. Shafi, A.F. Molisch, M. Dohler, H. Sjolund, and F. Tufvesson, “6G wireless systems: Vision, requirements, challenges, insights, and opportunities”, *Proceedings of the IEEE*, vol. 109, no. 7, pp. 1166–1199, 2021.
- 
- M. Alrabeiah and A. Alkhateeb, “Deep learning for mmwave beam and blockage prediction using sub-6 GHz channels”, *IEEE Trans. Commun.*, 2020.
- 
- A. Alkhateeb, “DeepMIMO: A generic deep learning dataset for millimeter wave and massive MIMO applications”, arXivpreprint arXiv:1902.06435, 2019.
- 
- X. Song, S. Haghighatshoar, and G. Caire, “A scalable and statistically robust beam alignment technique for mm-Wave systems”, *IEEE Trans. Wireless Commun.*, vol. 17, no. 7, pp. 4792–4805, 2018.
- 
- J.B. Wang, M. Cheng, J. Y. Wang, M. Lin, Y. Wu, H. Zhu, and J. Wang, “Bandit inspired beam searching scheme for mmWave high-speed train communications”, *arXiv preprint arXiv:1810.06150*, 2018.
- 
- M. Hashemi, A. Sabharwal, C.E. Koksai, and N.B. Shroff, “Efficient beam alignment in millimeter wave systems using contextual bandits”, *IEEE Conference on Computer Communications (INFOCOM)*, 2018.

- Uplink signal: captured by AP at **sub-6 GHz**

$$\mathbf{y}^{\text{UL}}[\ell] = \mathbf{h}^{\text{UL}}[\ell] x^{\text{UL}}[\ell] + \mathbf{n}^{\text{UL}}[\ell]$$

$$\mathbf{h}^{\text{UL}}[\ell] \in \mathbb{C}^{M \times 1}; \quad \mathbb{E}[|x^{\text{UL}}[\ell]|^2] = P^{\text{UL}}/L; \quad \mathbf{n}^{\text{UL}}[\ell] \sim \mathcal{N}(0, (\sigma^{\text{UL}})^2)$$

- Downlink signal: received by the user at the **mmWave band**

$$y^{\text{DL}}[\ell] = \mathbf{h}^{\text{DL}\dagger}[\ell] \mathbf{f} x^{\text{DL}}[\ell] + n^{\text{DL}}[\ell]$$

$$\mathbf{h}^{\text{DL}}[\ell] \in \mathbb{C}^{N \times 1}; \quad \mathbf{f} \in \mathbb{C}^{N \times 1}, \|\mathbf{f}\|^2 = 1; \quad \mathbb{E}[|x^{\text{DL}}[\ell]|^2] = P^{\text{DL}}/L; \quad n^{\text{DL}}[\ell] \sim \mathcal{N}(0, (\sigma^{\text{DL}})^2)$$

- Dataset samples: $(\{\mathbf{h}_i[\ell]^{\text{UL}}\}_{\ell=1}^L, \{\mathbf{h}_i[\ell]^{\text{DL}}\}_{\ell=1}^L)_i$

L number of OFDM subcarriers, i sample index

- Sub-6 GHz channels capture wireless characteristics, easy to estimate
- Dataset split at each link: 80% for training (15% validation set) and 20% for testing

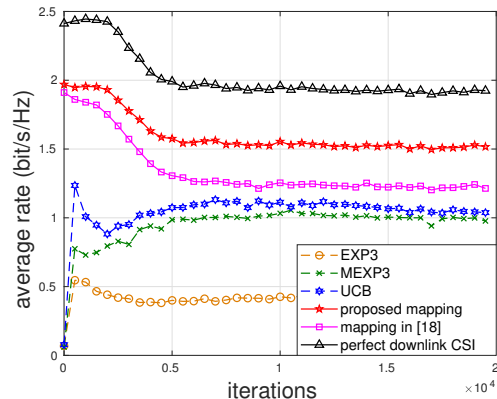
- System parameters:
 - ▶ Learning parameters: 100 epochs, ADAM optimizer (10^{-4}), $\mathcal{B} = 256$
 - ▶ Dataset parameters: 4 AP, outdoor scenario: 3.5 GHz $\rightarrow$ 28 GHz
 - ▶ Samples in each link: 108 419, 108 781, 63 350 and 117 650
 - ▶ $P^{\text{DL}} = 34$ dBm, $M = 4$, $N = 64$ and $L = 32$, mmWave bandwidth 0.5 GHz
- Comparison benchmarks:
 - ▶ Centralized learning: same neural network, full access to all data samples
 - ▶ Perfect downlink CSI: full and instantaneous knowledge of downlink channels
 - ▶ Individual learning: fully distributed, no cooperation between APs

MABs

- + No *offline* training, online learning
- + Low complexity $\sim 10^2$ ops/iteration
- + 1-bit feedback (ACK/NACK)
- Quantized beams $\rightarrow$ performance loss
- Trial and error $\rightarrow$ transitory phase

Deep learning

- Offline training on large and relevant datasets, *DeepMIMO* [Alkhateeb'19]
- Higher complexity $\sim 10^6 - 10^7$ ops/iteration
- sub-6 GHz channel information
- + Regression (no quantization)
- + Better running performance (no transitory phase)



MABs: EXP3, MEXP3, UCB

Deep learning [Alrabeiah'20]: classification

Depends on the target application, its characteristics and requirements

I. Chafaa, R. Negrel, E.V. Belmega, and M. Debbah, "Self-supervised deep learning for mmWave beam steering exploiting sub-6 GHz channels", *IEEE Trans. on Wireless Commun.*, 2022.

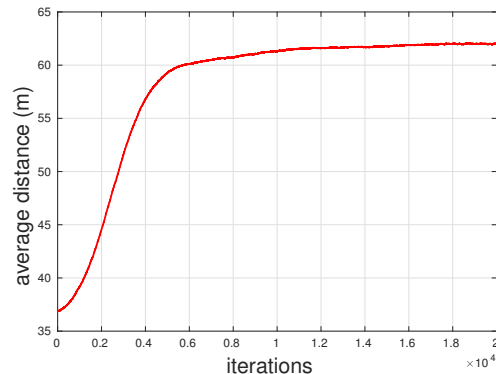
I. Chafaa, E.V. Belmega, and M. Debbah, "One-bit Feedback exponential learning for beam alignment in mobile mmWave", *IEEE Access*, Oct. 2020.

MABs

- + No *offline* training, online learning
- + Low complexity $\sim 10^2$ ops/iteration
- + 1-bit feedback (ACK/NACK)
- Quantized beams $\rightarrow$ performance loss
- Trial and error $\rightarrow$ transitory phase

Deep learning

- Offline training on large and relevant datasets, *DeepMIMO* [Alkhateeb'19]
- Higher complexity $\sim 10^6 - 10^7$ ops/iteration
- sub-6 GHz channel information
- + Regression (no quantization)
- + Better running performance (no transitory phase)



MABs: EXP3, MEXP3, UCB

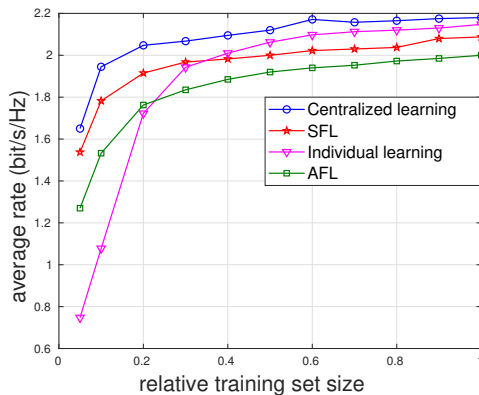
Deep learning [Alrabeiah'20]: classification

Mobility model: average distance between the AP and user increases in time

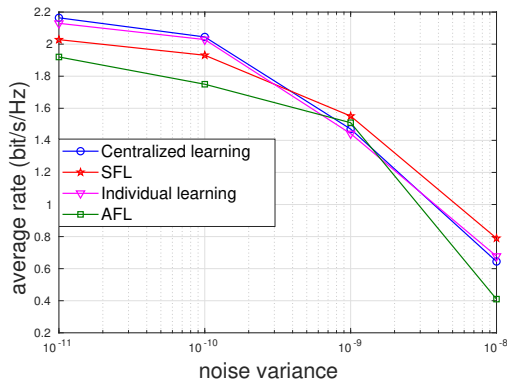
Depends on the target application, its characteristics and requirements

I. Chafaa, R. Negrel, E.V. Belmega, and M. Debbah, "Self-supervised deep learning for mmWave beam steering exploiting sub-6 GHz channels", *IEEE Trans. on Wireless Commun.*, 2022.

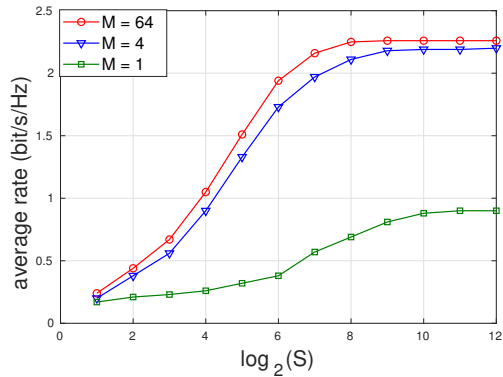
I. Chafaa, E.V. Belmega, and M. Debbah, "One-bit Feedback exponential learning for beam alignment in mobile mmWave", *IEEE Access*, Oct. 2020.



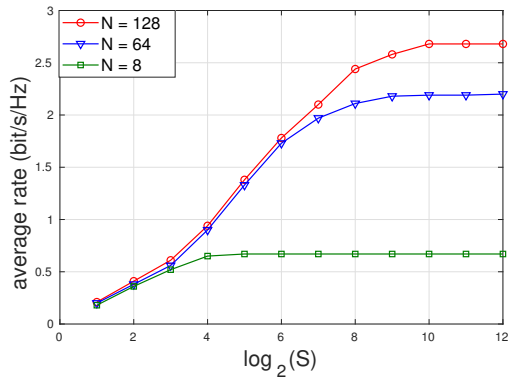
- FL approach is efficient in case of local data scarcity



- Uplink channels are contaminated with different levels of noise
- Variance of uplink channels: 10^{-9}
- SFL outperforms CL and individual learning at high noise regime



- $N = 64$ and $L = 32$
- When the input size M increases, the optimal size decreases



- $M = 4$ and $L = 32$.
- When the output size N increases, the optimal size increases as well

MABs

- UCB attains its ε -optimal performance (ε -regret) after $1/\varepsilon$ iterations (up to logarithmic factors)
- EXP3 and MEXP3 attains it after $1/\varepsilon^2$ iterations
- One iteration of these algorithms scales linearly with the codebook size: $\mathcal{O}(A)$

Deep learning

- Complexity of $\mathcal{O}(LMS + S^2 + SN)$ at each iteration

Entropic regularizer $h(\cdot)$: $D_h(\mathbf{y}, \mathbf{x}) = \sum_j y_j \log \frac{y_j}{x_j}$ (Kullback-Leibler divergence)

Exponential weights (EW / HEDGE / Multiplicative weights)

[Auer'02]

- Full information (gradient feedback): $\mathbf{v}(t) \equiv \mathbf{r}(t) = \nabla U_t(\mathbf{p}(t))$ (reward vector)
- Recursive update (multiplicative):

$$p_s(t+1) = \frac{p_s(t) \exp(\gamma r_s(t))}{\sum_{i=1}^S p_i(t) \exp(\gamma r_i(t))}, \quad \forall s$$

- Equivalent to:

$$\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \mathbf{r}(t) \quad \text{cumulative reward score}$$

$$p_s(t+1) = \frac{\exp(y_s(t+1))}{\sum_{i=1}^S \exp(y_i(t+1))}, \quad \forall s \quad \text{exponential mapping}$$

- Regret bound: $\text{Reg}_T \leq \sqrt{2T \log S}$ with $\gamma^* = \sqrt{2 \log S / T}$

- Convex conjugate of the regularizer $h(\mathbf{p})$ - strongly convex:

$$\begin{aligned}h^*(\mathbf{y}) &= \max_{\mathbf{p} \in \mathcal{S}} \mathbf{y}^T \mathbf{p} - h(\mathbf{p}) \\&= \log \left(\sum_i \exp(y_i) \right)\end{aligned}$$

- Its gradient is equivalent with the update of EW algorithm $\mathbf{p}(t+1) = \nabla h^*(\mathbf{y}(t+1))$

$$\frac{\partial h^*(\mathbf{y})}{\partial y_i} = \frac{\exp(y_i)}{\sum_k \exp(y_k)}$$

- The update is also equivalent with

$$\mathbf{p}(t+1) = \arg \max_{\mathbf{p} \in \mathcal{S}} \mathbf{y}(t)^T \mathbf{p} - h(\mathbf{p})$$

$$\text{Feasible set: } \mathcal{P} = \left\{ \mathbf{p} \in \mathbb{R}_+^S : \sum_{s=1}^S p_s \leq P_{\max} \right\}$$

$$\text{Our regularizer } h(\mathbf{p}) = \sum_i p_i \log p_i + (P_{\max} - \sum_i p_i) \log(P_{\max} - \sum_i p_i)$$

$$D_h(\mathbf{q}, \mathbf{p}) = \sum_j q_j \log \frac{q_j}{p_j} + (P_{\max} - \sum_k q_k) \log \frac{P_{\max} - \sum_k q_k}{P_{\max} - \sum_k p_k}$$

Kullback-Leibler type of divergence with additional variables:

$$q_{s+1} = P_{\max} - \sum_{k=1}^S q_k, p_{s+1} = P_{\max} - \sum_{k=1}^S p_k \text{ (the unused power of zero reward)}$$

OMD algorithm

- Gradient feedback: $\mathbf{v}(t) = -\nabla L_t(\mathbf{p}(t))$
- Update:

$$p_s(t+1) = \frac{P_{\max} p_s(t) \exp(\gamma v_s(t))}{\sum_{i=1}^S p_i(t) \exp(\gamma v_i(t)) + P_{\max} - \sum_k p_k(t)}, \quad \forall s$$

- Equivalent to OXL:

$$\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \mathbf{v}(t) \quad \text{cumulative gradient score}$$

$$p_s(t+1) = P_{\max} \frac{\exp(y_s(t+1))}{1 + \sum_{i=1}^S \exp(y_i(t+1))}, \quad \forall s \quad \text{exponential mapping}$$

- Convex conjugate of $h(\mathbf{p})$

$$\begin{aligned} h^*(\mathbf{y}) &= \max_{\mathbf{p} \in \mathcal{P}} \mathbf{y}^T \mathbf{p} - h(\mathbf{p}) \\ &= P_{\max} \log \left(1 + \sum_i \exp(y_i) \right) - P_{\max} \log P_{\max} \end{aligned}$$

- $h(\mathbf{p})$ is $1/P_{\max}$ - strongly convex wrt $\|\cdot\|_1 \implies h^*(\mathbf{y})$ is $P_{\max}$ strongly smooth wrt $\|\cdot\|_\infty$
- Its gradient is equivalent with the update of OXL algorithm: $\mathbf{p}(t+1) = \nabla h^*(\mathbf{y}(t+1))$

$$\frac{\partial h^*(\mathbf{y})}{\partial y_i} = P_{\max} \frac{\exp(y_i)}{1 + \sum_k \exp(y_k)}$$

- The update is also equivalent with

$$\mathbf{p}(t+1) = \arg \max_{\mathbf{p} \in \mathcal{P}} \mathbf{y}(t)^T \mathbf{p} - h(\mathbf{p})$$

- First-order convexity condition of $L_t(\cdot)$ (linearization), $\mathbf{v}(t) = \nabla L_t(\mathbf{p}(t))$:

$$Reg_{\mathbf{q}}(T) \leq - \sum_{t=1}^T \langle \mathbf{v}(t) | \mathbf{p}(t) - \mathbf{q} \rangle$$

- Cumulative score: $\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \mathbf{v}(t)$ and $\mathbf{y}(1) = 0$

$$Reg_{\mathbf{q}}(T) \leq - \sum_{t=1}^T \langle \mathbf{v}(t) | \mathbf{p}(t) \rangle + \frac{1}{\gamma} \langle \mathbf{y}(T+1) | \mathbf{q} \rangle$$

- Using $\mathbf{p}(t) = \nabla h^*(\mathbf{y}(t))$ and $R_{\max}$ -strong smoothness of $h^*(\mathbf{y})$

$$h^*(\mathbf{y}(t+1)) \leq h^*(\mathbf{y}(t)) + \gamma \langle \mathbf{v}(t) | \nabla h^*(\mathbf{y}(t)) \rangle + \frac{\gamma^2}{2} R_{\max} \|\mathbf{v}(t)\|_{\infty}^2,$$

$$Reg_{\mathbf{q}}(T) \leq \frac{1}{\gamma} \left[h^*(0) - h^*(\mathbf{y}(T+1)) \right] + \frac{\gamma}{2} R_{\max} \sum_{t=1}^T \|\mathbf{v}(t)\|_{\infty}^2 + \frac{1}{\gamma} \langle \mathbf{y}(T+1) | \mathbf{q} \rangle.$$

- Fenchel inequality: $h^*(\mathbf{y}(T+1)) + h(\mathbf{q}) \geq \langle \mathbf{y}(T+1) | \mathbf{q} \rangle$

$$Reg_{\mathbf{q}}(T) \leq \frac{1}{\gamma} [h(\mathbf{q}) + R_{\max} \log(1+S) - R_{\max} \log(R_{\max})] + \frac{\gamma}{2} R_{\max} VT.$$

- From $h(\mathbf{q}) \leq R_{\max} \log(R_{\max})$

$$Reg(T) \leq \frac{R_{\max} \log(1+S)}{\gamma} + \frac{\gamma}{2} R_{\max} VT$$

- Fixed step-size γ

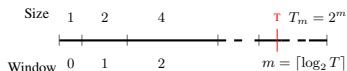
$$Reg_T \leq \underbrace{\frac{P_{\max} \log(1+S)}{\gamma}}_{\text{exploration penalty}} + \underbrace{\frac{\gamma P_{\max} VT}{2}}_{\text{exploitation penalty}}$$

$$V = \max \|\mathbf{v}(t)\|^2$$

- If T is known: optimal step-size $\gamma^* = \sqrt{\frac{2 \log(1+S)}{TV}}$

$$Reg_T \leq P_{\max} \sqrt{2V \log(1+S)T}$$

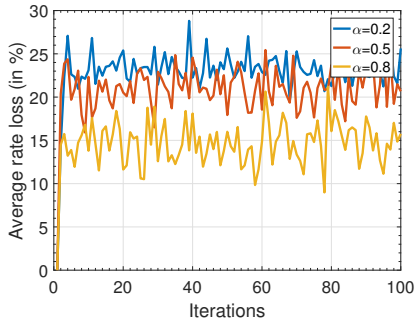
- If T is not-known: window doubling trick



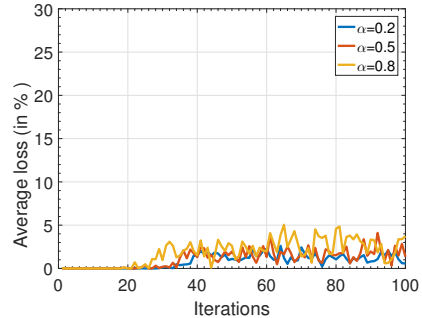
$$Reg_T \leq \frac{2}{\sqrt{2}-1} P_{\max} \sqrt{2V \log(1+S)T}$$

$$\text{Average rate loss: } \frac{[R_{\min} - R(\mathbf{p})]^+}{R_{\min}}$$

Waterfilling (energy-driven)



OXL



Having outdated feedback impacts the waterfilling performance at high channel variability.

Single device time-varying setup: $M = N = 1, S = 4, g(t+1) = \alpha g(t) + (1 - \alpha)z(t)$ with $z(t) \sim \mathcal{N}(0, \sigma_z^2)$ (average over 90 draws), $R_{\max} = 2, R_{\min} = 3, \lambda = 1$

- **Received feedback:** unbiased gradient estimation, $\tilde{\mathbf{v}}(t) = -\nabla L_t(\mathbf{p}(t)) + \text{error}$

$$\mathbb{E} [\tilde{\mathbf{v}}(t) \mid \mathcal{F}_{t-1}] = -\nabla L_t(\mathbf{p}(t)) \quad \text{no systematic errors}$$

$$\mathbb{E} [\|\tilde{\mathbf{v}}(t)\|^2 \mid \mathcal{F}_{t-1}] \leq \tilde{V} \quad \text{bounded mean square}$$

$\mathcal{F}_t$ - history of play up to t

- OXL with $\mathbf{v}(t) \leftarrow \tilde{\mathbf{v}}(t)$ and an optimal step-size $\gamma^* = \sqrt{\frac{2 \log(1+S)}{TV}}$

$$\mathbb{E} [Reg_T] \leq P_{\max} \sqrt{2\tilde{V} \log(1+S) T}$$

Rk: The expected regret is not highly impacted by unbiased errors in the gradient estimation.

- First-order convexity condition (linearization):

$$\mathbb{E}[Reg_{\mathbf{q}}(T)] \leq \mathbb{E} \left[\sum_{t=1}^T \langle \nabla L_t(\mathbf{p}(t)) | \mathbf{p}(t) - \mathbf{q} \rangle \right]$$

- Using $\nabla L_t(\mathbf{p}(t)) = -\mathbb{E}[\tilde{\mathbf{v}}(t) | \tilde{\mathbf{v}}(t-1), \dots, \tilde{\mathbf{v}}(1)]$ and the law of total expectation

$$\mathbb{E} \left[\sum_{t=1}^T \langle \nabla L_t(\mathbf{p}(t)) | \mathbf{p}(t) - \mathbf{q} \rangle \right] = -\mathbb{E} \left[\sum_{t=1}^T \langle \tilde{\mathbf{v}}(t) | \mathbf{p}(t) - \mathbf{q} \rangle \right]$$

- Using $\mathbb{E} [\|\tilde{\mathbf{v}}(t)\|_2^2] \leq \tilde{V}$

$$\mathbb{E}[Reg(T)] \leq \frac{P_{\max} \log(1+S)}{\mu} + \frac{\mu}{2} P_{\max} T \tilde{V}$$

Stochastic gradient approximation: $\tilde{\mathbf{v}}(t) = \frac{S}{\delta} L_t(\mathbf{p} + \delta \mathbf{u}) \mathbf{u}$,

[Spall'97, Flaxman'05]

by sampling the loss function not at $\mathbf{p}$, but at a nearby point $\mathbf{u}$ - uniformly drawn over the unit S-dimensional Sphere

Properties:

$$\mathbb{E} [\tilde{\mathbf{v}}] = \nabla \tilde{L}_t(\mathbf{p}) \quad \text{link with the gradient of } \tilde{L}_t(\mathbf{p})$$

$$\tilde{L}_t(\mathbf{p}) \triangleq \mathbb{E} [L_t(\mathbf{p} + \delta \mathbf{u})] \quad \tilde{L}_t(\mathbf{p}) - \text{smooth approximation of } L_t(\mathbf{p})$$

$$|L_t(\mathbf{p}) - \tilde{L}_t(\mathbf{p})| \leq L\delta$$

L - the Lipschitz constant of $L_t(\mathbf{p})$

- For fixed γ and δ

$$\mathbb{E} [\text{Reg}_T] \leq \frac{P_{\max} \log(1+S)}{2\gamma} + \gamma \mathbf{T} S^2 \left(\frac{B}{\delta} + L \right)^2 + L \mathbf{T} \delta \left(3 + P_{\max} (S + 2\sqrt{S}) \right)$$

$$B = \max L_t(\mathbf{p})$$

- Choosing $\delta^* = \frac{P_{\max}}{(S+\sqrt{S})T^{1/4}}$ and $\gamma^* = \sqrt{\frac{P_{\max} \log(1+S)}{2T}} \left[S \left(\frac{B}{\delta^*} + L \right) \right]^{-1}$, the expected regret is sublinear

$$\mathbb{E} [\text{Reg}_T] = \mathcal{O}(T^{3/4} S^3 \sqrt{\log(1+S)})$$

Rk: The zero-th order feedback greatly impacts the decay rate of the average regret.

→ **Tradeoff:** regret decay rate vs. amount of required feedback!

New regularizer: $h(\mathbf{p}) = \sum_i (p_i - \delta) \log(p_i - \delta) + (P_{\max} - \sqrt{S}\delta - \sum_i p_i) \log(P_{\max} - \sqrt{S}\delta - \sum_i p_i)$

$$D_h(\mathbf{q}, \mathbf{p}) = \sum_j (q_j - \delta) \log \frac{q_j - \delta}{p_j - \delta} + (P_{\max} - \sqrt{S}\delta - \sum_k q_k) \log \frac{P_{\max} - \sqrt{S}\delta - \sum_k q_k}{P_{\max} - \sqrt{S}\delta - \sum_k p_k}$$

Kullback-Leibler type of divergence adapted to the shrunk set

OMD algorithm

- Gradient feedback: $\mathbf{v}(t) = -\nabla L_t(\mathbf{p}(t))$

- Update:

$$p_s(t+1) = \delta + \frac{P_{\max}(1-C_\delta) (p_s(t)-\delta) \exp(\gamma v_s(t))}{\sum_{i=1}^S (p_i(t) - \delta) \exp(\gamma v_i(t)) + P_{\max} - \sqrt{S}\delta - \sum_k p_k(t)}, \quad \forall s$$

- Equivalent to OXL₀:

$$\mathbf{y}(t+1) = \mathbf{y}(t) + \gamma \tilde{\mathbf{v}}(t) \quad \text{cumulative score}$$

$$p_s(t+1) = \delta + P_{\max}(1 - C_\delta) \frac{\exp(y_s(t+1))}{1 + \sum_{i=1}^S \exp(y_i(t+1))}, \quad \forall s \quad \text{modified exponential mapping}$$

adapted to the shrunk set $\mathcal{P}_\delta$

$$C_\delta = \frac{\delta}{P_{\max}} (S + \sqrt{S})$$

- Convex conjugate of $h(\mathbf{p})$:

$$\begin{aligned} h^*(\mathbf{y}) &= \max_{\mathbf{p} \in \mathcal{P}_\delta} \mathbf{y}^T \mathbf{p} - h(\mathbf{p}) \\ &= \delta \sum_i y_i + P_{\max}(1 - C_\delta) \log \left(1 + \sum_i \exp(y_i) \right) \\ &\quad - P_{\max}(1 - C_\delta) \log(P_{\max}(1 - C_\delta)) \end{aligned}$$

- Its gradient is equivalent with the update of OXL algorithm: $\mathbf{p}(t+1) = \nabla h^*(\mathbf{y}(t))$

$$\frac{\partial h^*(\mathbf{y})}{\partial y_i} = \delta + P_{\max}(1 - C_\delta) \frac{\exp(y_i)}{1 + \sum_k \exp(y_k)}$$

- The update is also equivalent with

$$\mathbf{p}(t+1) = \arg \max_{\mathbf{p} \in \mathcal{P}_\delta} \mathbf{y}(t)^T \mathbf{p} - h(\mathbf{p})$$

- Compare $L_t(\mathbf{p}_\delta(t) + \delta \mathbf{u})$, $L_t(\mathbf{q})$ to $L_t(\mathbf{p}_\delta(t))$, $L_t(\mathbf{q}_\delta)$ using the L -Lipschitz property (move to the shrunk set)

$$\mathbb{E}[\text{Reg}_{\mathbf{q}}(T)] \leq \mathbb{E} \left[\sum_{t=1}^T L_t(\mathbf{p}_\delta(t)) - L_t(\mathbf{q}_\delta) \right] + LT\delta \left(1 + P_{\max} \left(S + 2\sqrt{S} \right) \right)$$

- Compare $L_t(\mathbf{p}_\delta(t))$, $L_t(\mathbf{q}_\delta)$ to smooth approx. $\tilde{L}_t(\mathbf{p}_\delta(t))$, $\tilde{L}_t(\mathbf{q}_\delta)$ for linking the regret with $\tilde{v}(t)$

$$\mathbb{E}[\text{Reg}_{\mathbf{q}}(T)] \leq \mathbb{E} \left[\sum_{t=1}^T \tilde{L}_t(\mathbf{p}_\delta(t)) - \tilde{L}_t(\mathbf{q}_\delta) \right] + LT\delta \left(3 + P_{\max} \left(S + 2\sqrt{S} \right) \right)$$

- First-order convexity condition of $\tilde{L}_t(\mathbf{p})$: $\mathbb{E} \left[\sum_{t=1}^T \tilde{L}_t(\mathbf{p}_\delta(t)) - \tilde{L}_t(\mathbf{q}_\delta) \right] \leq \mathbb{E} \left[\sum_{t=1}^T \langle \nabla \tilde{L}_t(\mathbf{p}_\delta(t)) | \mathbf{p}_\delta(t) - \mathbf{q}_\delta \rangle \right]$
- Using $\nabla \tilde{L}_t(\mathbf{p}_\delta(t)) = \mathbb{E}[\tilde{\mathbf{v}}(t) | \mathbf{u}(1), \dots, \mathbf{u}(t-1)]$ and the total law of expectation:

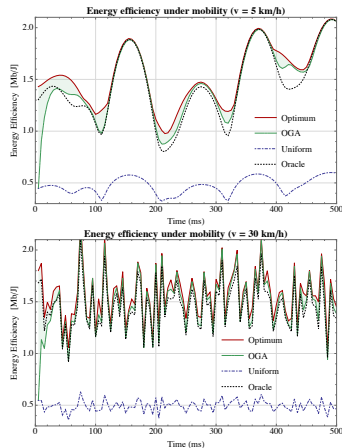
$$\mathbb{E} \left[\sum_{t=1}^T \langle \nabla \tilde{L}_t(\mathbf{p}_\delta(t)) | \mathbf{p}_\delta(t) - \mathbf{q}_\delta \rangle \right] \leq -\mathbb{E} \left[\sum_{t=1}^T \langle \tilde{\mathbf{v}}(t) | \mathbf{p}_\delta(t) - \mathbf{q}_\delta \rangle \right]$$
- Using the convex conjugate of $h(\mathbf{p}_\delta)$, the links with the update and $\|\tilde{v}(t)\|_\infty^2 \leq S^2(B/\delta + L)^2$

$$\mathbb{E}[\text{Reg}(T)] \leq \frac{P_{\max} \log(1+S)}{2\gamma} + \gamma TS^2 \left(\frac{B}{\delta} + L \right)^2 + LT\delta \left(3 + P_{\max} \left(S + 2\sqrt{S} \right) \right).$$

where $B = \max_{t, \mathbf{p}} L_t(\mathbf{p})$

- Based on **gradient feedback**, we proposed matrix gradient descent and exponential learning for rate/energy-efficiency maximization
- Exponential mapping preserves the positivity of the input covariance matrices
- Our online policies perform close to the **ideal benchmark: dynamic optimal one**

COST-HATA wireless channels, mobile users $[3 - 130]$ km/h, $K = 15$ users,
 4×8 MIMO, $S = 8$ bands



- Multi-antenna devices: M_k transmit antennas at user $k \in \{1, \dots, K\}$, N receive antennas
- Achieved Shannon sum rate

$$R(\mathbf{Q}) = \log \det \left(\mathbf{I} + \sum_{k=1}^K \mathbf{H}_k \mathbf{Q}_k \mathbf{H}_k^\dagger \right)$$

$\mathbf{H}_k \in \mathbb{C}^{N \times M_k}$ - channel matrix between transmitter k and receiver, $\mathbf{Q} = (\mathbf{Q}_1, \dots, \mathbf{Q}_K)$, $M = \max_k M_k$

- Distributed optimization: each transmitter k tunes its input covariance matrix

$$\mathbf{Q}_k \in \mathbb{C}^{M_k \times M_k} \text{ s.t.}$$

$$\mathbf{Q}_k \succeq 0, \quad \text{Tr}(\mathbf{Q}_k) = P_{\max, k}$$

- **Assumptions:** static channels, $S = 1$
- Rate maximization algorithms based on **gradient feedback**:

MXL - matrix exponential learning [MertikopoulosB'16]

IWF - iterative waterfilling [Yu'04]

$$\nabla_k R(\mathbf{Q}) = \mathbf{H}_k^\dagger \left[\mathbf{I} + \sum_{\ell=1}^K \mathbf{H}_\ell \mathbf{Q}_\ell \mathbf{H}_\ell^\dagger \right]^{-1} \mathbf{H}_k$$

- Massive MIMO systems: $N \gg M, K \gg 1$

Issue: too much feedback $\mathcal{O}(\min\{KM^2, N^2\})$ at each iteration!

MXL algorithm

[MertikopoulosB'16]

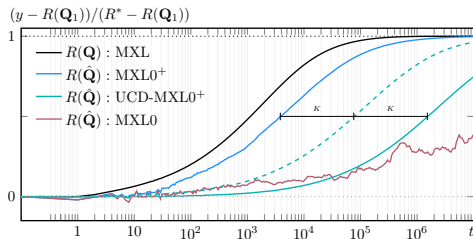
- Transmit at $\mathbf{Q}(t)$
- Receive **matrix** gradient feedback:
 $\mathbf{V}(t) = \nabla_k R(\mathbf{Q}(t))$
- **Update:**

$$\mathbf{Y}(t+1) = \mathbf{Y}(t) + \gamma \mathbf{V}(t)$$

$$\mathbf{Q}(t+1) = \Lambda(\mathbf{Y}_t), \quad \Lambda_k = P_{\max, k} \frac{\exp(\mathbf{Y}_k)}{\text{tr}(\exp(\mathbf{Y}_k))}$$

- MXL with **utility-based gradient estimator**: feedback $R(t) = R(\mathbf{Q}(t))$, slow convergence $\mathcal{O}(T^{-1/4})$, poor scaling with problem dimensionality (MXL0)
- **Idea**: exploit the current sum rate feedback $R(t)$ jointly with the previous one $R(t-1) = R(\mathbf{Q}(t-1))$ (MXL0⁺)

→ improved two-point gradient estimator with no additional feedback



Arbitrary static channels, $K = 20$ users, 3 transmit and 16 receive antennas

Remark: for **static** convex and Lipschitz objectives, MXL0⁺ recovers the convergence rate $\mathcal{O}(1/\sqrt{T})$ as MXL with perfect gradient feedback !

- First-order derivative: $\ell'(x) = \lim_{\delta \rightarrow 0} \frac{\ell(x+\delta) - \ell(x-\delta)}{2\delta}$
- A two query estimator for small δ : $\hat{v}_2 = \frac{\ell(x+\delta) - \ell(x-\delta)}{2\delta}$
- One query random estimator: $u \sim \text{Bernoulli}(1/2)$

$$\hat{v} = \frac{\ell(x + \delta u)u}{\delta}, \quad \mathbb{E}[\hat{v}] = \hat{v}_2$$

Bias: $|\mathbb{E}[\hat{v}] - v| = \mathcal{O}(\delta)$, variance: $|\hat{v}| = \mathcal{O}(1/\delta)$

- One query with callback ρ :

$$\hat{v}_\rho = \frac{(\ell(x + \delta u) - \rho) u}{\delta}, \quad \rho = \ell(x)$$

Bias: $|\mathbb{E}[\hat{v}_\rho] - v| = \mathcal{O}(\delta)$, variance: $|\hat{v}_\rho| \leq L$
 L - Lipschitz constant of ℓ

- Ensure the query point falls in the feasible set

$$\hat{\mathbf{Q}}_k = \mathbf{Q}_k + \frac{\delta}{r_k} (\mathbf{C}_k - \mathbf{Q}_k) + \delta \mathbf{Z}_k \in \mathcal{Q}_k$$

K users, M_k transmit antennas, N receive antennas, $M = \max_k M_k$

$\mathbf{C}_k = \mathbf{I}/M_k$, $r_k = 1/\sqrt{M_k(M_k - 1)}$, $\delta < r_k$

$\mathbf{Z}_k$ uniformly drawn on the unit sphere, $d_k = M_k^2 - 1$ dimension of $\mathcal{Q}_k$

- One-point estimator: $\hat{\mathbf{V}}_k(\mathbf{Q}) = \frac{d_k}{\delta} R(\hat{\mathbf{Q}}) \mathbf{Z}_k$
- One-point estimator with callback: $\tilde{\mathbf{V}}_{k,t}(\mathbf{Q}) = \frac{d_k}{\delta} [R(\hat{\mathbf{Q}}_t) - R(\hat{\mathbf{Q}}_{t-1})] \mathbf{Z}_{k,t}$
- **Remark:** our algorithm does not require two queries at each iteration t !
- Convergence: $\mathcal{O}(\sqrt{K^4 M^6 \log M/T})$ with $\gamma \propto 1/\sqrt{T}$, $\delta \propto 1/\sqrt{T}$

Qu: Can we design an *adaptive NOMA* scheme via MABs relying on **little feedback** and **no CSIT/CDIT**?

Main challenge: Unknown channels that vary in time

- Simple case: $K = 2$ (or user pairing)
- Decoding order is a *discrete* control variable $i \in \{1, 2\}$
 - ▶ User $i \in \{1, 2\}$ performs SIC
 - ▶ User $j \neq i$ suffers interference
- Fix the transmit rates $R_{\min,k}$ and measure the outage

$$\mathbb{P}_{\text{out}}(i, \mathbf{p}) = \mathbb{P}[R_i(t) \leq R_{\min,i} \cup \min\{R_{j \rightarrow j}(t), R_{j \rightarrow i}(t)\} \leq R_{\min,j}]$$

- Feedback: 1-ACK/NACK bit from each user

[ElHassaniSB'21a, ChafaaBD'20]

$$R_i(t) = \log \left(1 + \frac{|h_i(t)|^2 p_i}{\sigma_i^2} \right), \quad R_{j \rightarrow j}(t) = \log \left(1 + \frac{|h_j(t)|^2 p_j}{|h_j(t)|^2 p_i + \sigma_j^2} \right), \quad R_{j \rightarrow i}(t) = \log \left(1 + \frac{|h_i(t)|^2 p_j}{|h_i(t)|^2 p_i + \sigma_i^2} \right)$$

- Ideal energy-efficient target

[Zhang'20]

$$\begin{aligned} \text{maximize} \quad & \text{GEE}(i, \mathbf{p}) \triangleq \frac{(R_{\min,i} + R_{\min,j}) (1 - \mathbb{P}_{\text{out}}(i, \mathbf{p}))}{p_i + p_j + P_c} \\ \text{subject to} \quad & i \in \{1, 2\}, \quad \mathbf{p} = (p_i, p_j) \in \mathbb{R}_+^2, \quad p_i + p_j \leq P_{\max} \end{aligned}$$

- Quantization of the power allocation
 - ▶ User fairness: less power is allocated to the SIC decoding user i

$$\mathbf{p}_\beta = (0.25 \beta P_{\max}, 0.75 \beta P_{\max}), \text{ for discrete } \beta \in \mathcal{B}$$

- ▶ Uniformly quantized $\mathcal{B} \subset [0, 1]$
- An arm $\mathbf{a} \triangleq (i, \mathbf{p}_\beta) \in \mathcal{A}$ decoding and power allocation vector

$$\mathcal{A} = \{1, 2\} \times \{\mathbf{p}_\beta = (0.25 \beta P_{\max}, 0.75 \beta P_{\max}) \mid \beta \in \mathcal{B}\}$$

Rk: The quantization and the $1/4 - 3/4$ user split will imply an optimality loss.

Initialize: $t = 1, \mathbf{a}(1) = (1, 0.5)$

For $t = 1$ to T

- Play arm $\mathbf{a}(t) = (i(t), \beta(t))$
 - ▶ Inform users of their decoding: $i(t)$
 - ▶ Transmit with $\mathbf{p}_{\beta}(t) = (0.25 \beta(t) P_{\max}, 0.75 \beta(t) P_{\max})$

- Receive 1-bit ACK/NACK feedback from each user

- Compute reward

$$U_t(\mathbf{a}(t)) = \begin{cases} \frac{R_{\min,1} + R_{\min,2}}{\beta(t) P_{\max} + P_c}, & \text{if } R_i(t, \mathbf{a}(t)) \geq R_{\min,i} \text{ and} \\ & \min\{R_{j \rightarrow j}(t, \mathbf{a}(t)), R_{j \rightarrow i}(t, \mathbf{a}(t))\} \geq R_{\min,j} \\ 0, & \text{otherwise} \end{cases}$$

- Update policy $\mathbf{a}(t+1) \leftarrow \mathbf{a}(t)$ via UCB or EXP3
- $t \leftarrow t + 1$

- Both UCB and EXP3 have the no regret property

→ adaptive NOMA reaches the best expected reward

$$\max_{\mathbf{a} \in \mathcal{A}} \mathbb{E}[U_t(\mathbf{a})] \equiv \max_{i \in \{1, 2\}, \beta \in \mathcal{B}} GEE(i, \mathbf{p}_\beta)$$

- **Regret** decay rate: $Reg_T^{\text{UCB}} = \mathcal{O}(\log T/T)$ and $Reg_T^{\text{EXP3}} = \mathcal{O}(1/\sqrt{T})$
- **Complexity** of UCB and EXP3: every iteration scales as $\mathcal{O}(|\mathcal{A}|)$

Rk: There are two fundamental tradeoffs: performance vs. feedback information and performance vs. complexity.

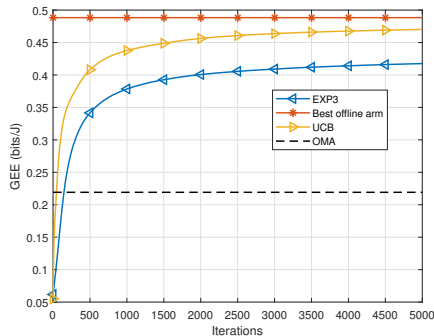
- Benchmark quantized OMA

$$R_k^{\text{OMA}}(t) = \frac{1}{2} \log \left(1 + \frac{|h_k(t)|^2 \beta P_{\max}}{\sigma_k^2} \right)$$

$$GEE^{\text{OMA}}(\beta) = \frac{(R_{\min,1} + R_{\min,2})(1 - \mathbb{P}_{\text{out}}^{\text{OMA}}(\beta))}{\beta P_{\max} + P_c}$$

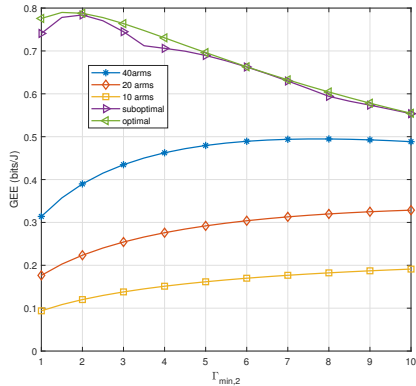
- Assumes perfect CDIT

$$\max_{\beta \in \mathcal{B}} GEE^{\text{OMA}}(\beta)$$

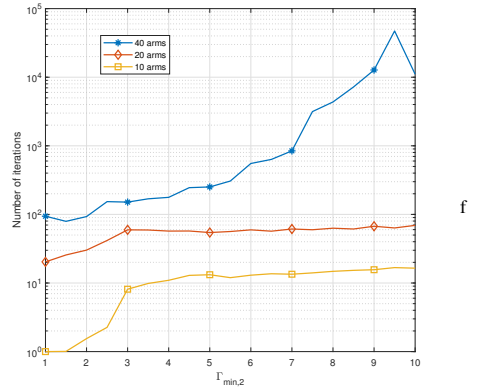


Our 1-bit feedback adaptive NOMA scheme outperforms quantized OMA with perfect CDIT.

$\sigma_k^2 = 0.1$, $R_{\min,1} = 1$ bpcu, $R_{\min,2} = 10$ bpcu, $P_c = 1$ W, $P_{\max} = 100$ W, $T = 5000$, 10^3 channel (Rayleigh) realizations



Energy efficiency as a function of $\Gamma_{\min,2} = 2^{R_{\min,2}} - 1$



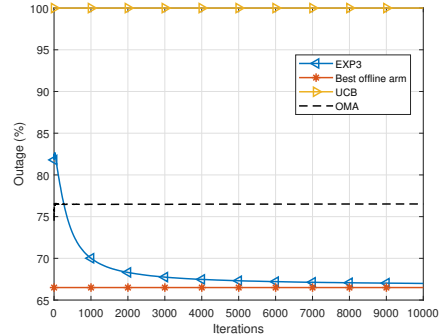
Number of iterations to reach a 10% regret level

Our adaptive 1-bit feedback NOMA scheme is sub-optimal compared to the ideal target (with CDIT).

Outage minimization problem

Malicious jammer

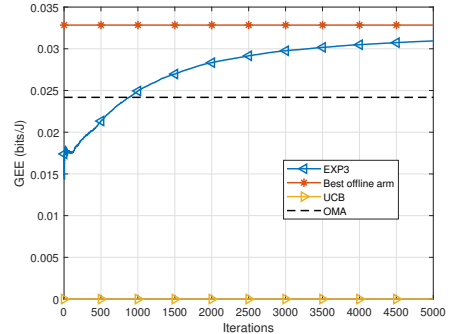
- Has access to perfect CSI
- Anticipates the arm that will be played via **deterministic** UCB
- Interferes such that the system is put systematically in outage



In non-stationary adversarial settings EXP3 reaches no regret while UCB is brought to a fault.

Malicious jammer

- Has access to perfect CSI
- Anticipates the arm that will be played via **deterministic** UCB
- Interferes such that the system is put systematically in outage



In non-stationary adversarial settings EXP3 reaches no regret while UCB is brought to a fault.

- They **tradeoff between data exploration and exploitation**

→ poorly chosen lead to **positive regret**

- From a theoretical perspective (to reach no regret asimptotically)

$$\eta^* > 2 \quad \text{and} \quad \gamma^* = \sqrt{\frac{|\mathcal{A}| \ln |\mathcal{A}|}{(e-1)T}}$$

- But these values minimize the upper-bound of the regret

→ not optimal for the actual regret

- They can be further shaped via empirical experiments

Online mirror descent for convex optimization

- **Regret guarantees** $\mathcal{O}(\sqrt{T})$ rely on Lipschitz continuity

$$|\ell(\mathbf{x}) - \ell(\mathbf{y})| \leq G\|\mathbf{x} - \mathbf{y}\| \quad \text{or} \quad \|\text{grad } \ell(\mathbf{x})\| \leq G$$

- We extend them to problems with **singular gradients** of the objective at the boundary of the feasible set
- **Idea:** generalize the Lipschitz continuity based on local Riemannian norms, $\|\mathbf{z}\|_x = \mathbf{z}^T \mathbf{G}(\mathbf{x}) \mathbf{z}$ with $\mathbf{G}(\mathbf{x}) \succeq 0$

Collaborator: P. Mertikopoulos (CNRS)

Student: K. Antonakopoulos (PhD)

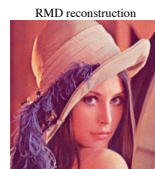
Supported by: Inria, ANR ORACLESS

Publications: 3 confs.

$$\text{e.g. } \ell(\mathbf{x}) = -\log(\mathbf{a}^T \mathbf{x}), \|\text{grad } \ell\|^2 = \frac{\sum_i a_i^2}{(\mathbf{a}^T \mathbf{x})^2}$$

$$\mathbf{G}(\mathbf{x}) = \mathbf{I} / \left(\sum_i x_i \right)^2 \Rightarrow \|\text{grad } \ell\|_x^2 \leq \frac{\sum_i a_i^2}{(\min_j a_j)^2}$$

- Applications: tomography, astronomical data
- Ill-conditioned reconstruction problems: $\hat{\mathbf{u}} = \mathbf{H}\mathbf{x} + \hat{\mathbf{z}}$ $\dim \hat{\mathbf{u}} \ll \dim \mathbf{x}$
 $\mathbf{H}$ - (Toeplitz, circulant, convoluting) optical system, $\hat{\mathbf{z}}$ - Poisson noise
 $\hat{u}_i \sim \text{Poisson}(\mathbf{H}\mathbf{x})_i$
- **Objective:** minimize Kullback-Leibler divergence
 $\ell(\mathbf{x}, \hat{\mathbf{u}}) = \sum_i [\hat{u}_i \log(\hat{u}_i / (\mathbf{H}\mathbf{x})_i) + (\mathbf{H}\mathbf{x})_i - \hat{u}_i]$
- Example: test image contaminated with Poisson noise



$$\dim \hat{\mathbf{u}} = \dim \hat{\mathbf{x}} / 2, \dim \hat{\mathbf{x}} \simeq 10^5$$

LR - Lucy-Richardson, mirror descent with entropic regularizer, [Bertero-2009]

CMD - mirror descent with log-barrier regularizer, [He-2016]

RMD - mirror descent with local Riemann norm

- Data **similarity** search and classification
Mahalanobis distance:

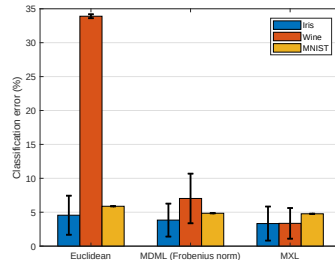
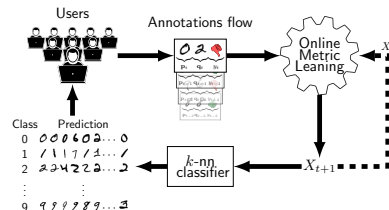
$$d_X(\mathbf{p}, \mathbf{q}) = \|\mathbf{X}^{1/2} \mathbf{p} - \mathbf{X}^{1/2} \mathbf{q}\|^2$$

$\mathbf{p}, \mathbf{q}$ - input data, e.g., images

- Objective:** learn the best discriminative linear data transformation matrix $\mathbf{X}$

$$\mathbf{X} \succeq 0, \quad \text{Tr}(\mathbf{X}) \leq c$$

- Training examples are obtained online, **on-the-fly**
- Our matrix exponential learning is competitive in a toy setup
- Other examples: online matrix completion, universal linear filtering, online dictionary learning



MDML - mirror descent with Frobenius norm regularizer [Kunapuli-2012]

- 
- R. M. Lee, M. J. Assante, and T. Conway, “Analysis of the cyber attack on the Ukrainian power grid. Defense use case”, *Electricity Information Sharing and Analysis Center*, Mar. 2016.
- 
- M. Ozay, I. Esnaola, F. T. Yarman Vural, S. R. Kulkarni, and H. V. Poor, “Machine learning methods for attack detection in the smart grid”, *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 27, no. 8, pp. 1773–1786, 2016.
- 
- A. Altman, A. Bercovici-Boden, and M. Tennenholtz, “Learning in one-shot strategic form games”, In *European Conf. on Machine Learning*, Sep. 2006.
- 
- Y. Wu, A. Khisti, C. Xiao, G. Caire, K.K. Wong, and X. Gao, “A survey of physical layer security techniques for 5G wireless networks and challenges ahead”, *IEEE J. Sel. Areas Commun.*, vol. 36, no 4, pp. 679–695, 2018.
- 
- H. Xing, K.K. Wong, A. Nallanathan, and R. Zhang, “Wireless powered cooperative jamming for secrecy multi-AF relaying networks”, *IEEE Trans. Wireless Commun.*, vol. 15, no. 12, pp. 7971–7984, 2016.
- 
- H. Zhang, J. Zhang, and K. Long, “Energy efficiency optimization for NOMA UAV network with imperfect CSI”, *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 12, pp. 2798–2809, 2020.
- 
- G. Kunapuli and J. Shavlik, “Mirror descent for metric learning: A unified approach”, *Joint European Conference on Machine Learning and Knowledge Discovery in Databases. Springer*, 2012, pp. 859–874.
- 
- W. Yu, W. Rhee, S. Boyd, and J. M. Cioffi, “Iterative water-filling for Gaussian vector multiple-access channels”, *IEEE Trans. Inf. Theory*, vol. 50, no. 1, pp. 145–152, 2004.
- 
- M. Bertero, P. Boccacci, G. Desidera, and G. Vicidomini, “Image deblurring with Poisson data: from cells to galaxies”, *Inverse Problems*, 2009.
- 
- N. He, Z. Harchaoui, Y. Wang, and L. Song, “Fast and simple optimization for poisson likelihood models”, *arXiv preprint arXiv:1608.01264*, 2016.