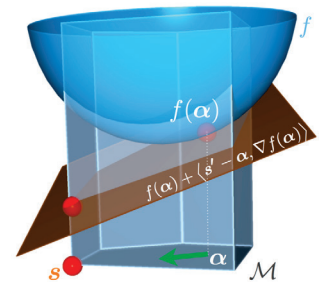


Recent Advances in Frank-Wolfe Optimization

Simon Lacoste-Julien

Université 
de Montréal



OSL 2017 – Les Houches – April 13th, 2017

Outline

- Frank-Wolfe algorithm review
- global linear convergence of FW optimization variants
 - condition number of domains & pyramidal width
- saddle point Frank-Wolfe

Frank-Wolfe algorithm [Frank, Wolfe 1956]

(aka conditional gradient)

- alg. for constrained opt.: $\min_{\alpha \in \mathcal{M}} f(\alpha)$

where:

f convex & cts. differentiable

$\mathcal{M}$ convex & compact

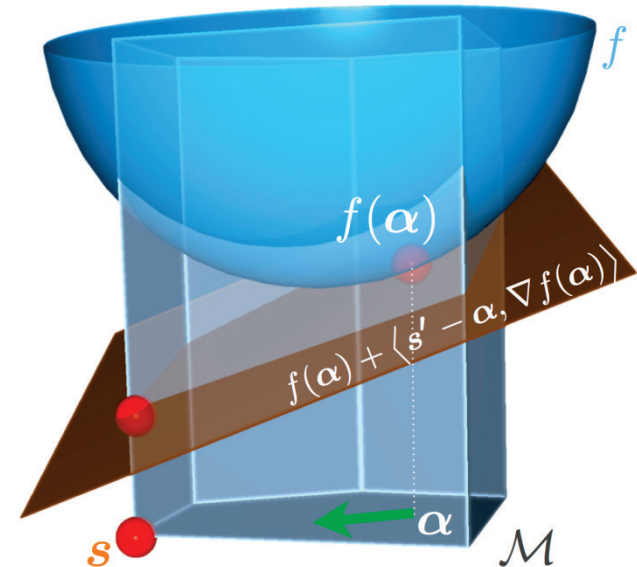
- FW algorithm – repeat:

1) Find good feasible direction by minimizing linearization of f :

$$s_{t+1} \in \arg \min_{s' \in \mathcal{M}} \langle s', \nabla f(\alpha_t) \rangle$$

2) Take convex step in direction:

$$\alpha_{t+1} = (1 - \gamma_t) \alpha_t + \gamma_t s_{t+1}$$



- Properties: $O(1/T)$ rate
 - sparse iterates
 - get duality gap $g(\alpha)$ for free
 - affine invariant
 - rate holds even if linear subproblem solved **approximately**

Frank-Wolfe: properties

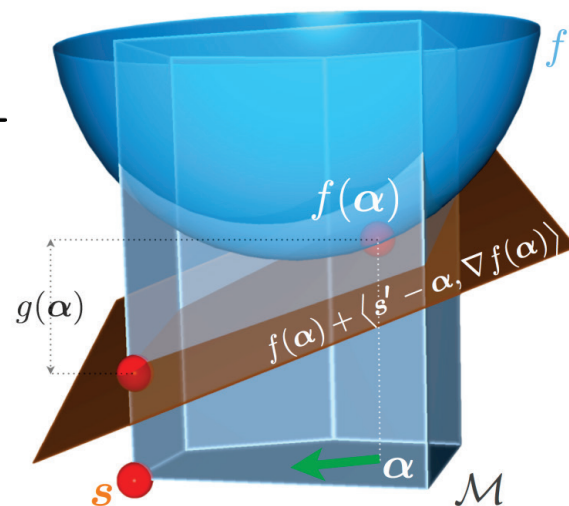
- convex steps => convex sparse combo:

$$\alpha_T = \rho_0 \alpha_0 + \sum_{t=1}^T \rho_t s_t \quad \text{where} \quad \sum_{t=0}^T \rho_t = 1$$

- get duality gap certificate for free

$$g(\alpha) := \max_{s' \in \mathcal{M}} \langle s' - \alpha, -\nabla f(\alpha) \rangle = \langle s - \alpha, -\nabla f(\alpha) \rangle$$

- (special case of Fenchel duality gap)
- also converge as $O(1/T)$!



- only need to solve linear subproblem

approximately (additive/multiplicative bound)

$\rightarrow O(1/\sqrt{T})$ for non-convex f
[L.-J. arXiv 2016]

- affine invariant!

see survey [Jaggi ICML 2013]

- numerically stable

(also [Lan arXiv 2013])

Why comeback of FW in ML?

- big data -> first order algorithm
- sparse algorithms – e.g. see references in [Locatello et. AISTATS 2017]
- structured constrained sets with cheaper LMOs:

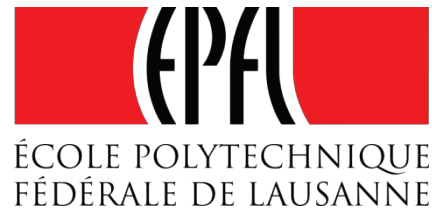
$\mathcal{X}$	Optimization Domain Atoms $\mathcal{A}$	$\mathcal{D} = \text{conv}(\mathcal{A})$	Complexity of one Frank-Wolfe Iteration $\Omega_{\mathcal{D}}^*(y) = \sup_{s \in \mathcal{D}} \langle s, y \rangle$	Complexity
$\mathbb{R}^n$	Sparse vectors	$\ \cdot\ _1$ -ball	$\ y\ _\infty$	$O(n)$
$\mathbb{R}^n$	Sign-vectors	$\ \cdot\ _\infty$ -ball	$\ y\ _1$	$O(n)$
$\mathbb{R}^n$	ℓ_p -Sphere	$\ \cdot\ _p$ -ball	$\ y\ _q$	$O(n)$
$\mathbb{R}^n$	Sparse non-neg. vectors	Simplex Δ_n	$\max_i \{y_i\}$	$O(n)$
$\mathbb{R}^n$	Latent group sparse vectors	$\ \cdot\ _{\mathcal{G}}$ -ball	$\max_{g \in \mathcal{G}} \ y_{(g)}\ _q^*$	$\sum_{g \in \mathcal{G}} g $
$\mathbb{R}^{m \times n}$	Matrix trace norm	$\ \cdot\ _{tr}$ -ball	$\ y\ _{op} = \sigma_1(y)$	$\tilde{O}(N_f / \sqrt{\varepsilon'})$ (Lanczos)
$\mathbb{R}^{m \times n}$	Matrix operator norm	$\ \cdot\ _{op}$ -ball	$\ y\ _{tr} = \ (\sigma_i(y))\ _1$	SVD
$\mathbb{R}^{m \times n}$	Schatten matrix norms	$\ (\sigma_i(\cdot))\ _p$ -ball	$\ (\sigma_i(y))\ _q$	SVD
$\mathbb{R}^{m \times n}$	Matrix max-norm	$\ \cdot\ _{\max}$ -ball		$\tilde{O}(N_f(n+m)^{1.5} / \varepsilon'^{2.5})$
$\mathbb{R}^{n \times n}$	Permutation matrices	Birkhoff polytope		$O(n^3)$
$\mathbb{R}^{n \times n}$	Rotation matrices			SVD (Procrustes prob.)
$\mathbb{S}^{n \times n}$	Rank-1 PSD matrices of unit trace	$\{x \succeq 0, \text{Tr}(x) = 1\}$	$\lambda_{\max}(y)$	$\tilde{O}(N_f / \sqrt{\varepsilon'})$ (Lanczos)
$\mathbb{S}^{n \times n}$	PSD matrices of bounded diagonal	$\{x \succeq 0, x_{ii} \leq 1\}$		$\tilde{O}(N_f n^{1.5} / \varepsilon'^{2.5})$

(table from [Jaggi ICML 2013])

On the Global Linear Convergence of Frank-Wolfe Optimization Variants

[L.-J. and Jaggi, NIPS 2015]

joint work with
Martin Jaggi



Problem setup

- We want to optimize over: $\mathcal{M} = \text{conv}(\mathcal{A})$

$$\min_{\mathbf{x} \in \mathcal{M}} f(\mathbf{x}) \quad \mathcal{A} \subseteq \mathbb{R}^d \text{ is a *finite* set of *atoms*}$$

with only access to a *linear minimization oracle*:

$$\text{LMO}_{\mathcal{A}}(\mathbf{r}) \in \arg \min_{\mathbf{x} \in \mathcal{A}} \langle \mathbf{r}, \mathbf{x} \rangle$$

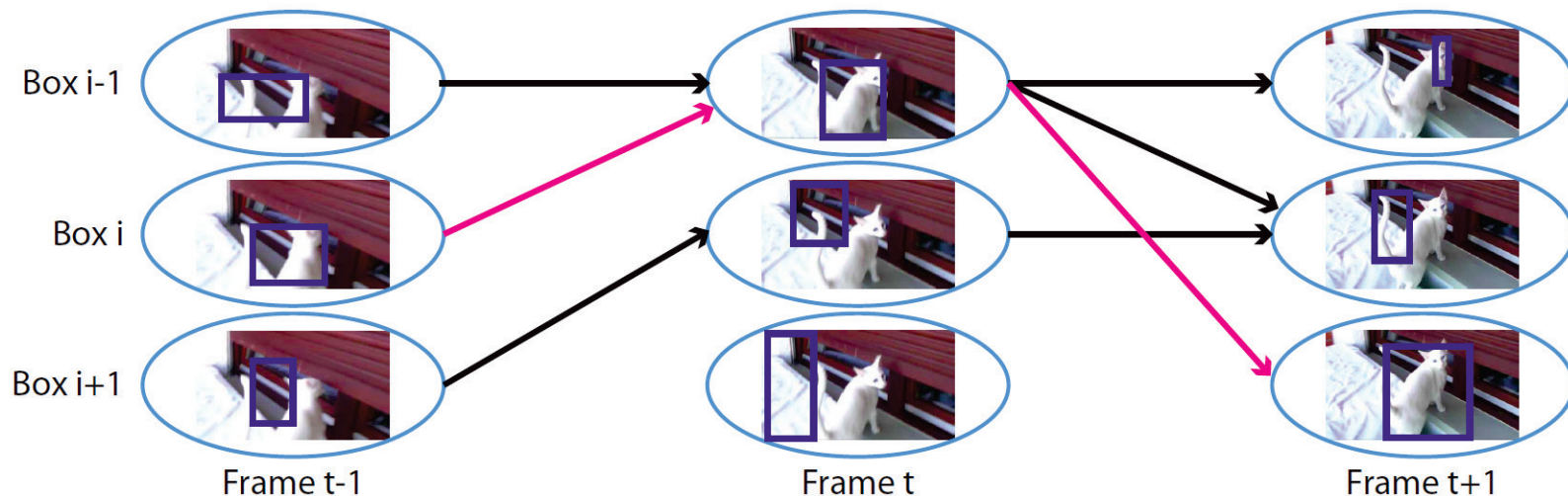
assume f is μ -strongly convex

and ∇f is L -Lipschitz continuous

Examples: QP over combinatorial polytopes

- For tracking [Chari, L.-J. et al. CVPR 15] or video co-localization [Joulin, Tang, Fei-Fei ECCV 14]

$\mathcal{M}$ is *flow polytope*, $\mathcal{A}$ contains integer flows



video co-localization

other examples...

- for structured SVM learning [L.-J., Jaggi et al. ICML 13] or approximate marginal inference [Krishnan, L.-J., Sontag NIPS 15]

$\mathcal{M}$ is *marginal polytope*,

$\mathcal{A}$ contains clique assignments in a MRF

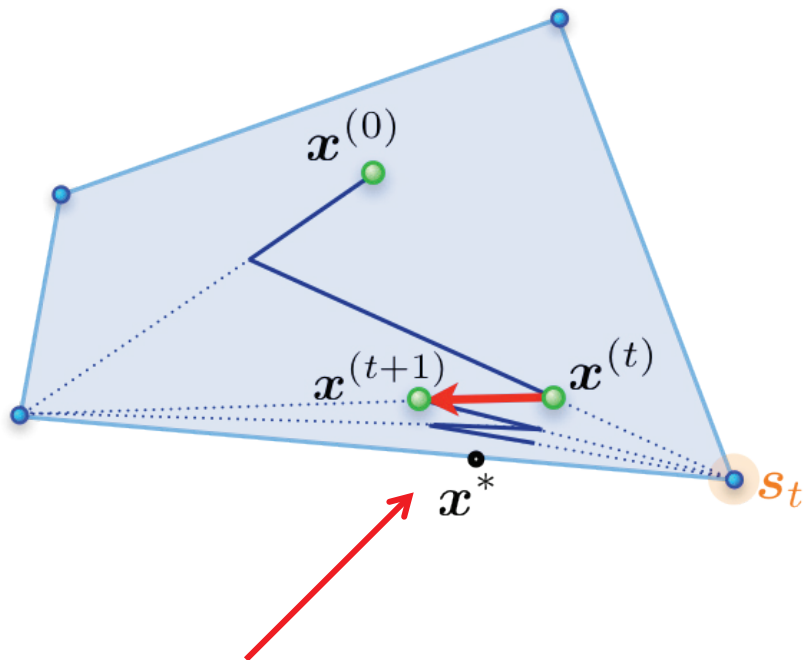
- for submodular function optimization [Bach TML 13]

$\mathcal{M}$ is *base polytope*,

$\mathcal{A}$ has one possible atom
per permutation of ground set

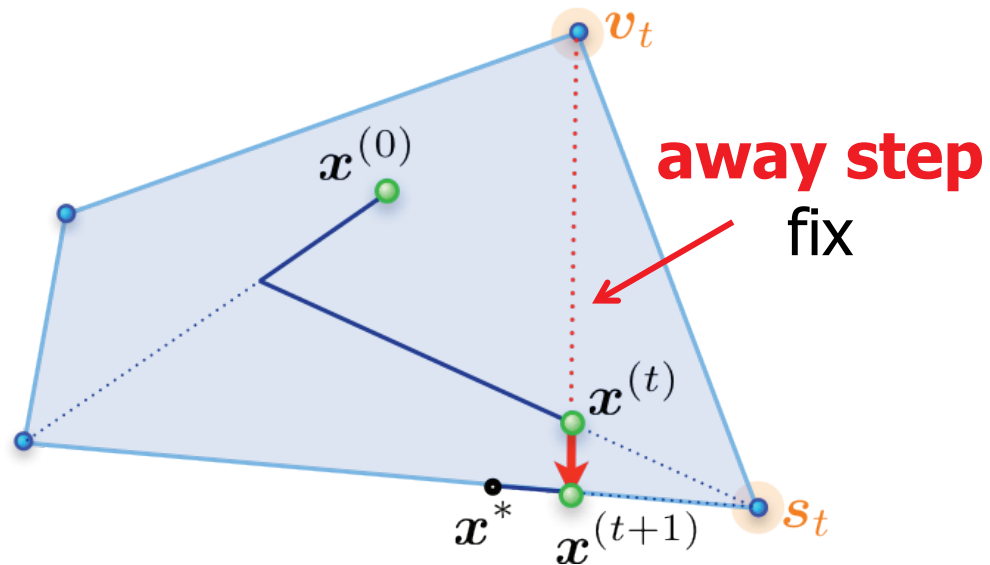
slow convergence of Frank-Wolfe...

standard FW



zig-zagging problem for FW

away-step FW



[Wolfe 1970]

[Guelat & Marcotte 1986]

Algorithm 1: Away-steps Frank-Wolfe: AFW($\mathbf{x}^{(0)}, \mathcal{A}, \epsilon$)

Let $\mathbf{x}^{(0)} \in \mathcal{A}$ and $\mathcal{S}^{(0)} := \{\mathbf{x}^{(0)}\}$; we maintain $\mathbf{x}^{(t)} = \sum_{\mathbf{v} \in \mathcal{S}^{(t)}} \alpha_{\mathbf{v}}^{(t)} \mathbf{v}$

for $t = 0 \dots T$ **do**

Let $\mathbf{s}_t := \text{LMO}_{\mathcal{A}}(\nabla f(\mathbf{x}^{(t)}))$ and $\mathbf{d}_t^{\text{FW}} := \mathbf{s}_t - \mathbf{x}^{(t)}$
(the FW direction)

Let $\mathbf{v}_t \in \arg \max_{\mathbf{v} \in \mathcal{S}^{(t)}} \langle \nabla f(\mathbf{x}^{(t)}), \mathbf{v} \rangle$ and $\mathbf{d}_t^{\text{A}} := \mathbf{x}^{(t)} - \mathbf{v}_t$
(the away direction)

if $g_t^{\text{FW}} := \langle -\nabla f(\mathbf{x}^{(t)}), \mathbf{d}_t^{\text{FW}} \rangle \leq \epsilon$ **then return** $\mathbf{x}^{(t)}$

if $\langle -\nabla f(\mathbf{x}^{(t)}), \mathbf{d}_t^{\text{FW}} \rangle \geq \langle -\nabla f(\mathbf{x}^{(t)}), \mathbf{d}_t^{\text{A}} \rangle$ **then**

$\mathbf{d}_t := \mathbf{d}_t^{\text{FW}}$, and $\gamma_{\max} := 1$ *(choose the FW direction)*

else

$\mathbf{d}_t := \mathbf{d}_t^{\text{A}}$, and $\gamma_{\max} := \alpha_{\mathbf{v}_t} / (1 - \alpha_{\mathbf{v}_t})$

(choose the away direction, and maximum feasible step-size)

end

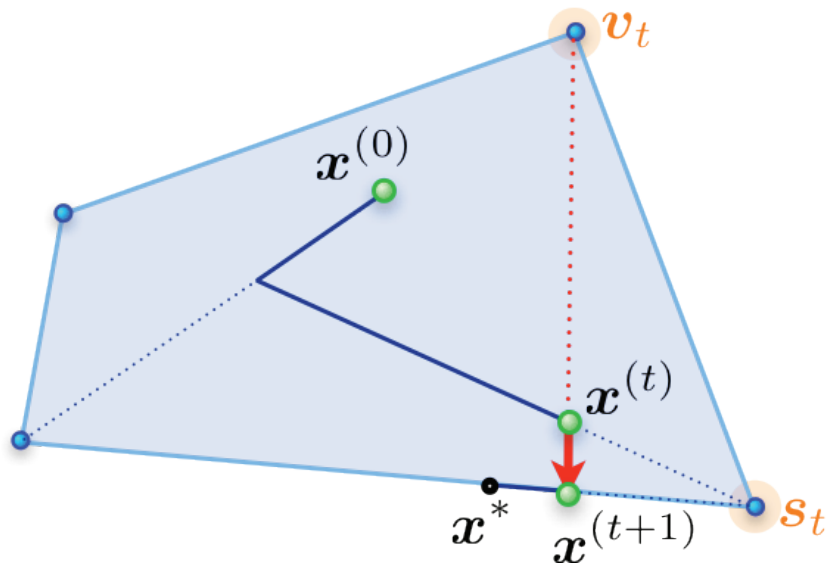
Line search: $\gamma_t \in \arg \min_{\gamma \in [0, \gamma_{\max}]} f(\mathbf{x}^{(t)} + \gamma \mathbf{d}_t)$

Update $\mathbf{x}^{(t+1)} := \mathbf{x}^{(t)} + \gamma_t \mathbf{d}_t$ *(and weights $\alpha^{(t+1)}$ accordingly)*

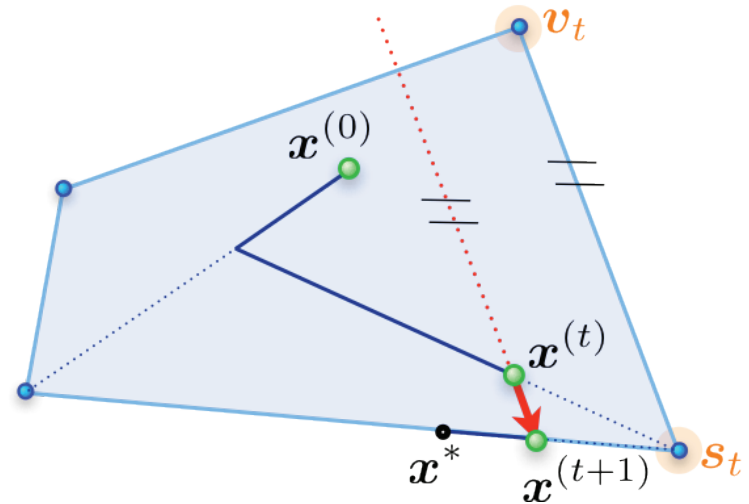
Update $\mathcal{S}^{(t+1)} := \{\mathbf{v} \text{ s.t. } \alpha_{\mathbf{v}}^{(t+1)} > 0\}$

end

other variants:



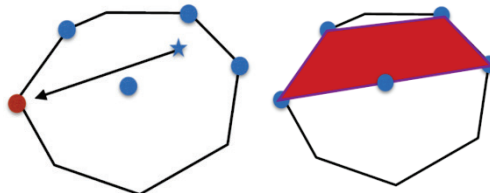
away-step FW



pairwise FW

[MDM 1974]

- fully-corrective FW (FCFW): re-optimize over convex hull of previously found vertices (correction polytope)



[Halloway 1974, Von Hohenbalken 1977, ...]

Previous convergence results

assumption: f is strongly convex (with Lipschitz gradient)

[Wolfe 70, Guélat & Marcotte 86]:

- Frank-Wolfe algorithm converges **linearly** if solution x^* is in relative **interior** of M
- Frank-Wolfe **with away steps** converges **linearly** with a constant depending on the distance between x^* and the boundary of M in the optimal face containing x^*

Problems:

- constant could be arbitrarily close to zero $\rightarrow$ not a true linear convergence result
- constant depends on unknown x^*
- analysis is **not affine invariant**

(FW alg. is invariant to affine transformations of variables)

Our contribution: [L.-J. & Jaggi NIPS 15, arxiv 13]

- we give an **affine invariant** analysis of the **global linear convergence** of Frank-Wolfe with away steps with constant bounded away from zero:

thm: $f(x^{(k)}) - f(x^*) \leq h_0 \exp(-\frac{1}{8} \rho_f k)$

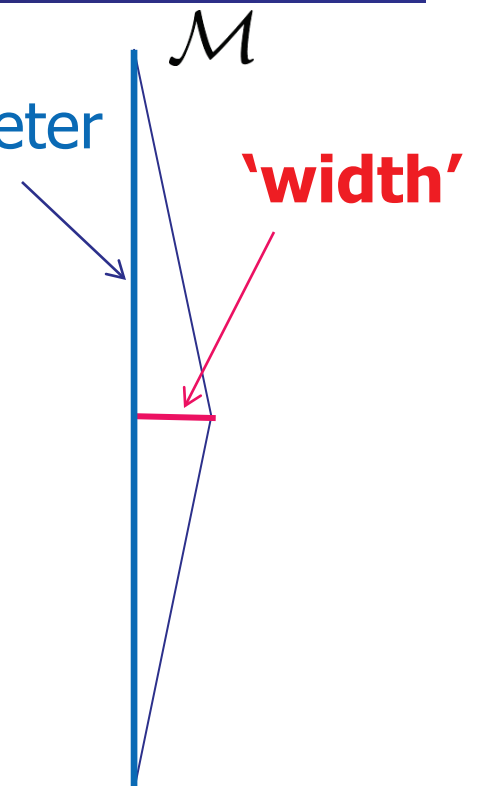
where: geometric strong convexity constant (new!)

$$\rho_f := \frac{\mu_f}{C_f} \begin{array}{l} \geq \mu(\text{pyr.width}(\mathcal{M}))^2 \\ \leq L(\text{diam}(\mathcal{M}))^2 \end{array}$$

curvature constant

f has L -Lipschitz gradient
and is μ -strongly convex

'Condition number' of domain!

$$\frac{1}{\rho_f} \leq \underbrace{\frac{L}{\mu}}_{\text{condition number of } f} \cdot \underbrace{\left(\frac{\text{diam}(\mathcal{M})}{\text{pyr.width}(\mathcal{M})} \right)^2}_{\substack{\text{'eccentricity' of } \mathcal{M} \\ \text{'condition number' of } \mathcal{M}}}$$


eccentricity in dimension d :

- probability simplex: $O(d)$
- unit cube: $O(d^2)$

Pyramidal width

- smallest directional width of pyramids built with **active set** as base, **FW point** as summit, and using a **feasible direction**

- value of $\approx 1/\sqrt{d}$ in dimension d :

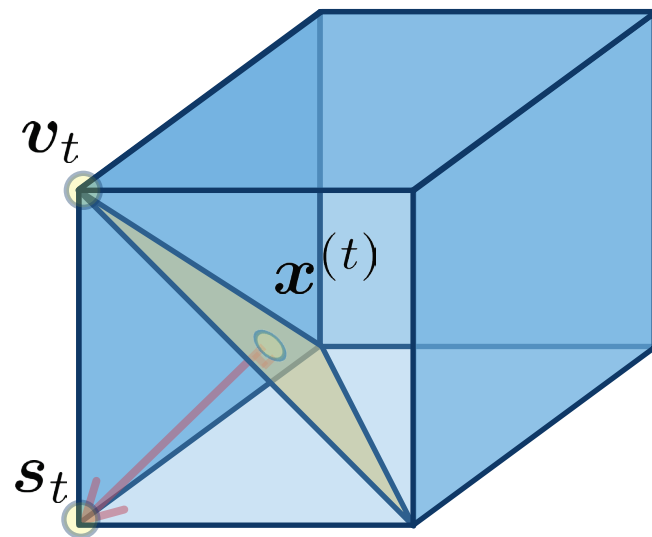
- prob. simplex
- unit cube
- l1-ball

- regular simplex has **smallest cond. number: $d/2$**

- this gives a complexity of:

$$O\left(d \frac{L}{\mu} \log \frac{1}{\epsilon}\right)$$

- unit cube has cond. number: d^2



$$\text{pyr. width(cube)} = 1/\sqrt{3}$$

[Pena & Rodriguez arXiv 2015]

shows equivalent to **facial distance**:

$$\min_{F \in \text{faces}(\mathcal{A})} \text{dist}(F, \text{conv}(\mathcal{A} \setminus F))$$

Proof elements

[from Guélat & Marcotte 86]

- as f has L -Lipschitz gradient:

$$\text{let } \mathbf{r}_t := -\nabla f(\mathbf{x}^{(t)})$$

$$f(\mathbf{x}^{(t+1)}) \leq f(\mathbf{x}^{(t)}) + \gamma \mathbf{d}_t \leq f(\mathbf{x}^{(t)}) + \gamma \langle \nabla f(\mathbf{x}^{(t)}), \mathbf{d}_t \rangle + \frac{\gamma^2}{2} L \|\mathbf{d}_t\|^2 \quad \forall \gamma \in [0, \gamma_{\max}]$$

$$(\text{suboptimality: } h_t := f(\mathbf{x}^{(t)}) - f(\mathbf{x}^*))$$

$$\text{Choose } \gamma \text{ to minimize RHS: } h_t - h_{t+1} \geq \frac{\langle \mathbf{r}_t, \mathbf{d}_t \rangle^2}{2L \|\mathbf{d}_t\|^2} = \frac{1}{2L} \langle \mathbf{r}_t, \hat{\mathbf{d}}_t \rangle^2 \quad (\text{when } \gamma_{\max} \text{ is big enough})$$

- as f is μ -strongly convex:

$$\text{with } \mathbf{e}_t := \mathbf{x}^* - \mathbf{x}^{(t)}$$

$$f(\mathbf{x}^{(t)}) + \gamma \mathbf{e}_t \geq f(\mathbf{x}^{(t)}) + \gamma \langle \nabla f(\mathbf{x}^{(t)}), \mathbf{e}_t \rangle + \frac{\gamma^2}{2} \mu \|\mathbf{e}_t\|^2$$

**angle between
negative gradient and
update direction**

Use $\gamma = 1$ on LHS; compare with lower bound of RHS:

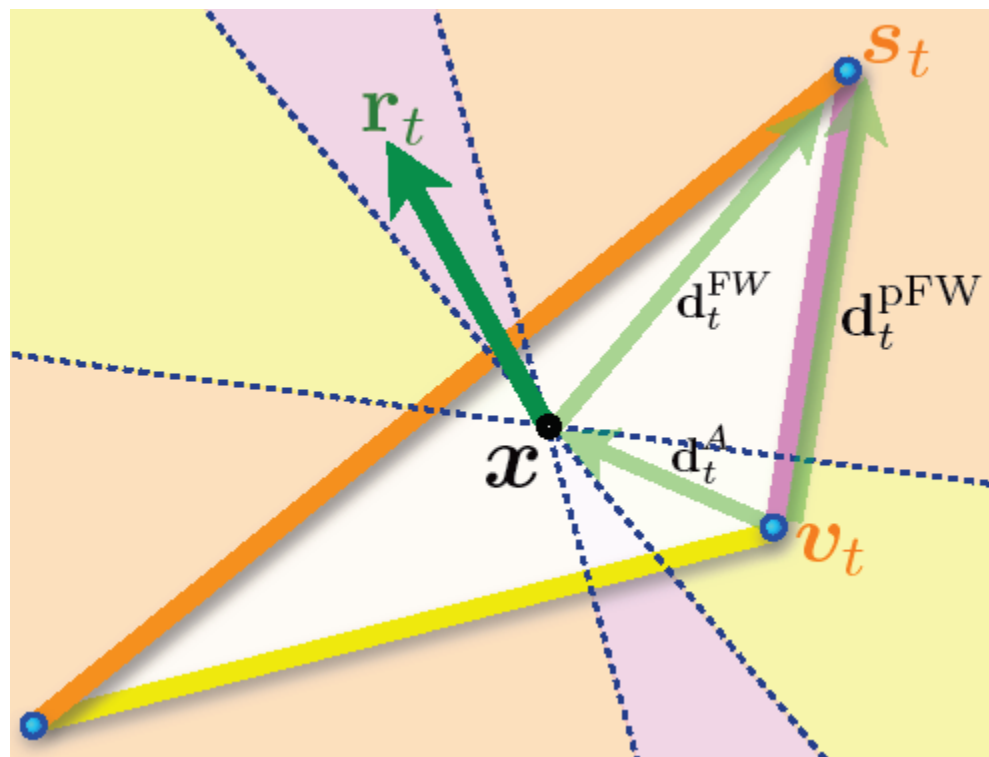
Combine to get:
(valid for any $\mathbf{d}_t$)

$$h_t - h_{t+1} \geq \frac{\mu}{L} \frac{\langle \mathbf{r}_t, \hat{\mathbf{d}}_t \rangle^2}{\langle \mathbf{r}_t, \hat{\mathbf{e}}_t \rangle^2} h_t \geq \frac{\mu}{L} \overbrace{\langle \hat{\mathbf{r}}_t, \hat{\mathbf{d}}_t \rangle^2} h_t$$

2 key insights:

$$\begin{aligned} 2\langle \mathbf{r}_t, \mathbf{d}_t \rangle &\geq \langle \mathbf{r}_t, \mathbf{d}_t^{\text{FW}} + \mathbf{d}_t^A \rangle \\ &= \langle \mathbf{r}_t, \mathbf{d}_t^{\text{PFW}} \rangle \text{ for AFW} \end{aligned}$$

$$\left\langle \hat{\mathbf{r}}, \hat{\mathbf{d}}^{\text{PFW}}(\hat{\mathbf{r}}) \right\rangle \geq \frac{PWidth(\mathcal{M})}{\text{diam}(\mathcal{M})}$$



(illustration showing possible PFW directions as $\mathbf{r}$ varies)

Important inequality

- key inequality which has been re-used several times:

$$g_t^{\text{PFW}} := \langle -\nabla f(\mathbf{x}_t), \mathbf{s}_t - \mathbf{v}_t \rangle$$

$$f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq \frac{1}{2\mu_f} (g_t^{\text{PFW}})^2$$

here is an unconstrained analog:

$$\frac{1}{2L} \|\nabla f(\mathbf{x}_t)\|^2 \leq f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq \frac{1}{2\mu} \|\nabla f(\mathbf{x}_t)\|^2$$

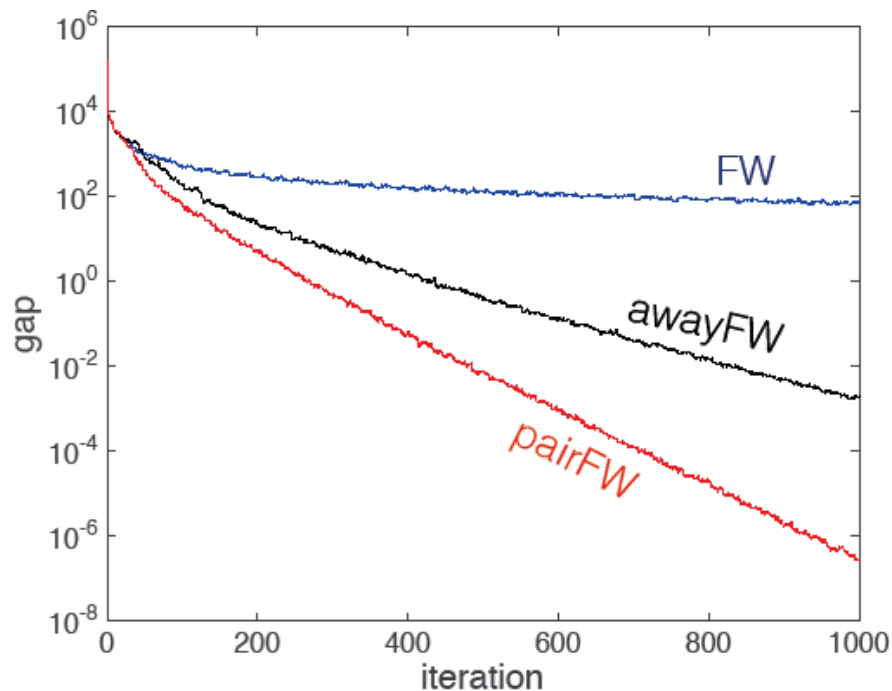
- used for:
 - ADMM + FW: [Yen et al. ICML 2016]
 - bandits [Berthet & Perchet arXiv 2017]
 - saddle pt. FW [Gidel et al. AISTATS 2017] (see 2nd part)
 - etc...

Illustrative experiments

Lasso regression:

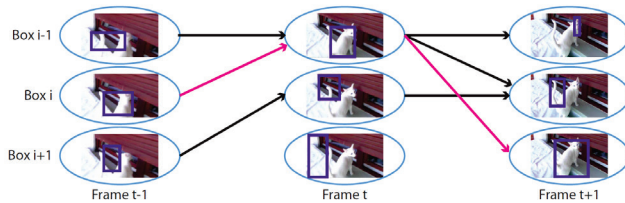
$$\min_{\|x\|_1 \leq 20} \|Ax - b\|_2^2$$

- Gaussian matrix $A \in \mathbb{R}^{200 \times 500}$
- $b = Ax^* + 0.1 \cdot \text{Gaussian noise}$
where x^* has 50 entries of ± 1 's

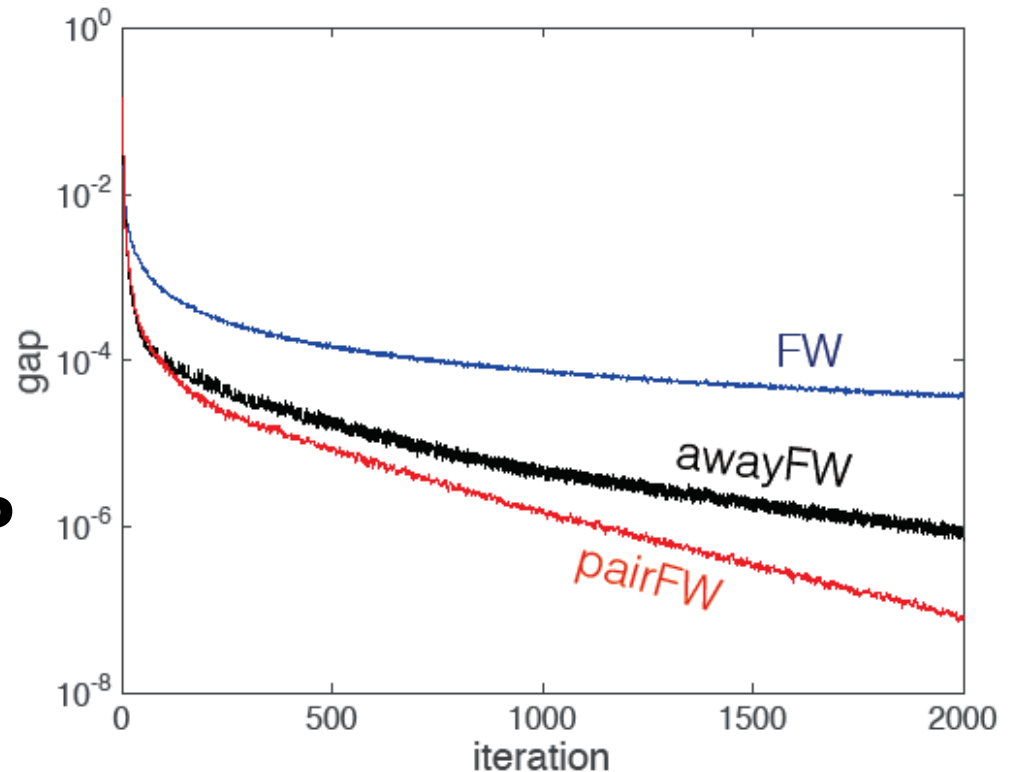


Video co-localization

- problem from [Joulin, Tang, Fei-Fei ECCV 14]

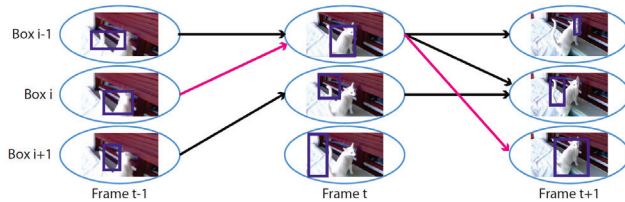


- QP over flow-polytope
- $d = 660$
- $\text{LMO}_{\mathcal{A}}$ can be solved using **shortest path DP** algorithm over network

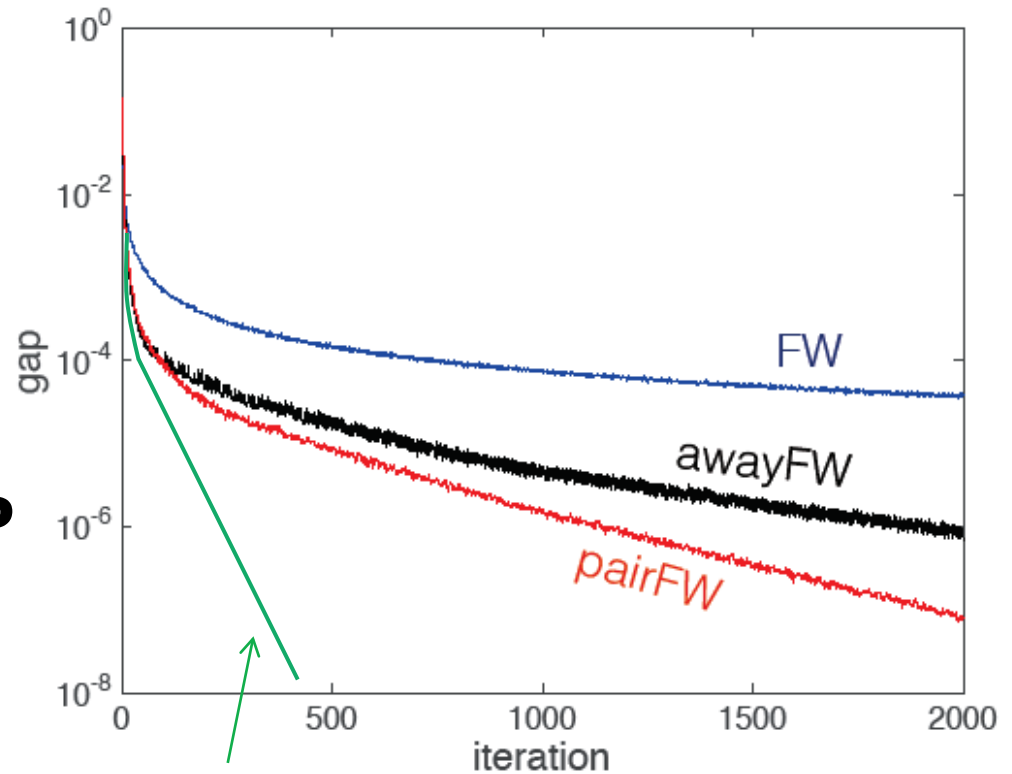


Video co-localization

- problem from [Joulin, Tang, Fei-Fei ECCV 14]



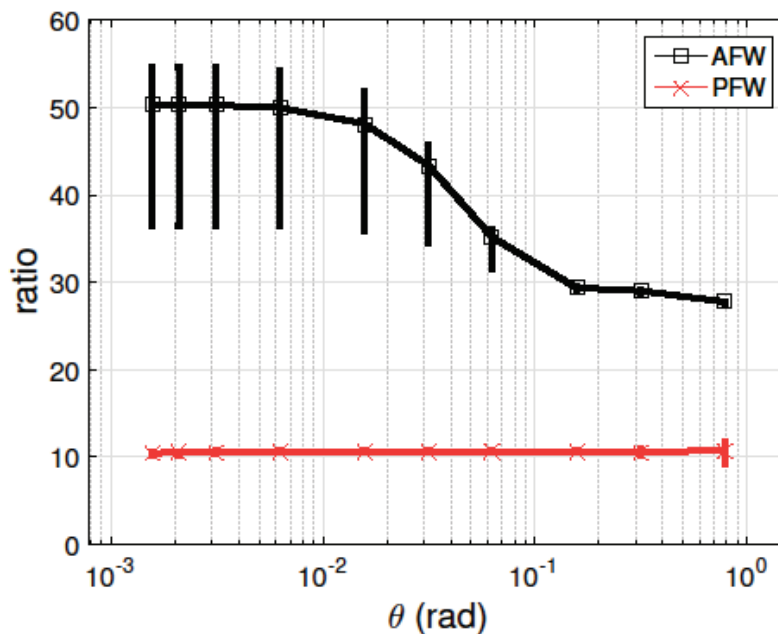
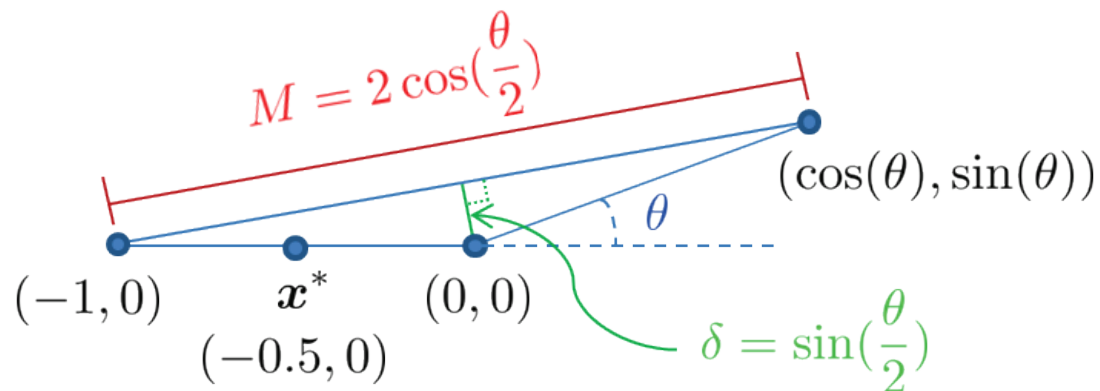
- QP over flow-polytope
- $d = 660$
- $\text{LMO}_{\mathcal{A}}$ can be solved using **shortest path DP** algorithm over network



pairFW + LMO away corner

[Garber & Meshi NIPS 2016]

Rate is empirically tight!



Discussion

- FW and variants popular in machine learning for optimization over structured polytopes
- Provide first truly global linear convergence rate for a Frank-Wolfe type algorithm *which does not need to compute any constants* (vs. [Garber & Hazan 13])
 - and *analysis is affine invariant*
 - can bound constant with condition number and purely geometric quantity 'eccentricity' -> **condition number for M**
- give first linear rate for FCFW, PFW and MNP
- extensions: used for ADMM / FW alg.; saddle point FW, etc.
 - reduce dependence to **~dimension of optimal face?**
 - -> YES: [Garber & Meshi NIPS 2016] for special 0-1 polytopes
 - AFW also linear rate for strongly convex sets
 - but general infinite number of atoms -> **still open question**

Other FW extensions / applications

- block-coordinate FW (for structured SVMs)

[L.-J. et al. ICML 2013]

AFW -> [Osokin et al. ICML 2016]

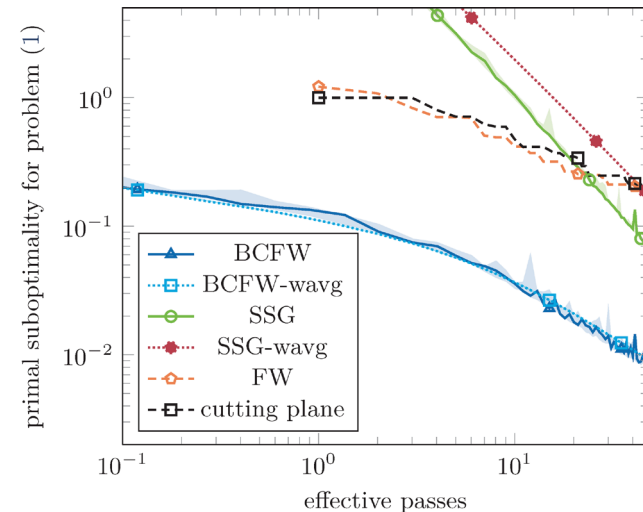
- barrier FW

[Krishnan, L.-J. & Sontag NIPS 2015]

- FW quadrature

[Bach, L.-J., Obozinski ICML 2012],

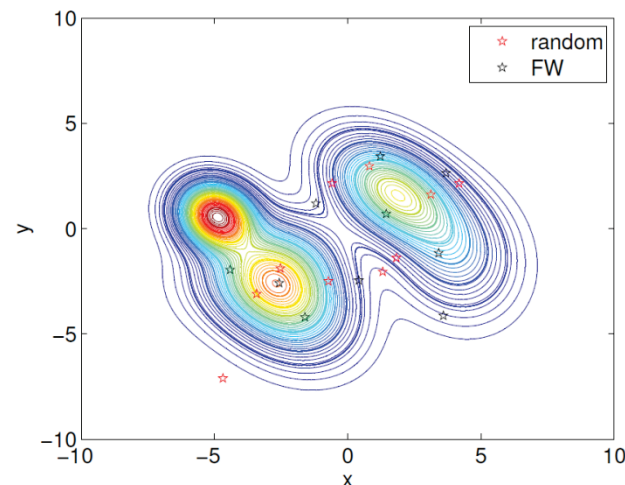
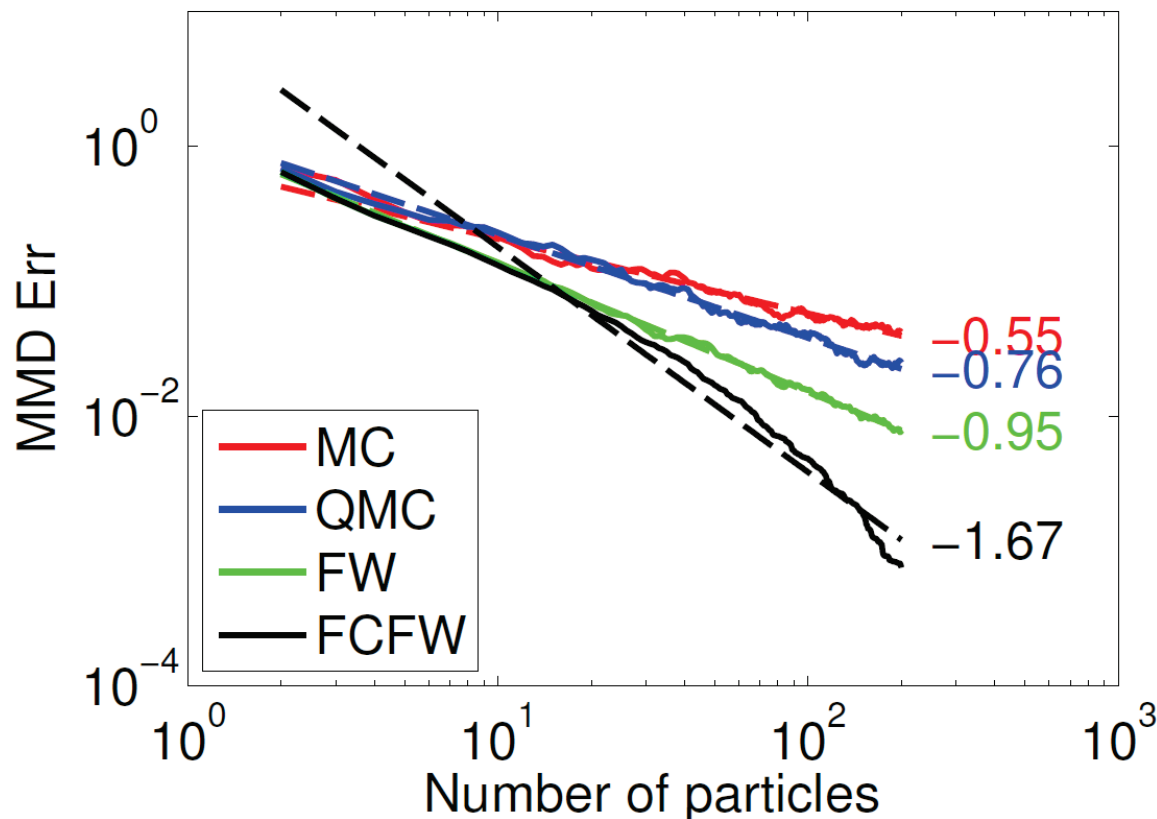
[L.-J., Lindsten, Bach AISTATS 2015]



FW quadrature

- for a mixture of Gaussians:

$$d = 2, \quad K = 5, \quad \sigma^2 = 1$$



[L.-J., Lindsten, Bach, AISTATS 15]

Frank-Wolfe Algorithms for Saddle Point Problems

[Gidel, Jebara & L.-J., AISTATS 2017]

with
Gauthier Gidel



Université 
de Montréal

Overview

Saddle point problem: solve $\min_{x \in \mathcal{X}} \max_{y \in \mathcal{Y}} \mathcal{L}(x, y)$

- want to solve using **only with LMOs**
- approach: extend FW to saddle point problems
- *straightforward* extension but **nontrivial** analysis
- related work:
 - [Lan arXiv 2013] -> use smoothing
 - [He & Harchaoui NIPS 2015] -> approximate projections
 - [Juditsky & Nemirovski MathProg 2016]
-> VIP transformations

Saddle point and link with variational inequalities

Let $\mathcal{L} : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$, where $\mathcal{X}$ and $\mathcal{Y}$ are convex and compact.

Saddle point problem: solve $\min_{x \in \mathcal{X}} \max_{y \in \mathcal{Y}} \mathcal{L}(x, y)$

A solution (x^*, y^*) is called a *Saddle Point*.

► *Necessary stationary conditions:*

$$\langle x - x^*, \nabla_x \mathcal{L}(x^*, y^*) \rangle \geq 0$$

$$\langle y - y^*, -\nabla_y \mathcal{L}(x^*, y^*) \rangle \geq 0$$

► *Variational inequality:*

$$\forall z \in \mathcal{X} \times \mathcal{Y} \quad \langle z - z^*, r(z^*) \rangle \geq 0$$

where $(x^*, y^*) = z^*$ and $r(z) = (\nabla_x \mathcal{L}(z), -\nabla_y \mathcal{L}(z))$

► *Sufficient condition: Global solution* if $\mathcal{L}$ *convex-concave*. $\forall (x, y) \in \mathcal{X} \times \mathcal{Y}$

$x' \mapsto \mathcal{L}(x', y)$ is convex and $y' \mapsto \mathcal{L}(x, y')$ is concave.

Motivations

- two-player games: $\min_{x \in \Delta(I)} \max_{y \in \Delta(J)} x^\top M y$
- structured SVM: $\min_{\omega \in \mathbb{R}^d} \lambda \Omega(\omega) + \frac{1}{n} \sum_{i=1}^n \underbrace{\max_{y \in \mathcal{Y}_i} (L_i(y) - \langle \omega, \phi_i(y) \rangle)}_{\text{structured hinge loss}}$

Regularization: penalized $\rightarrow$ constrained.

$$\min_{\substack{\Omega(\omega) \leq \beta}} \max_{\alpha \in \Delta(|\mathcal{Y}|)} b^T \alpha - \omega^T M \alpha$$

Hard to project when:

- ▶ *Structured sparsity* norm (group lasso norm).
- ▶ The output $\mathcal{Y}$ is structured: *exponential size*.

- $\rightarrow$ still looking for more: **call for applications!**

Standard methods to solve SP

projected gradient

Simple algorithm to solve Saddle point optimization:

$$\mathbf{x}^+ = P_{\mathcal{X}}(\mathbf{x} - \gamma \nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}, \mathbf{y}))$$

$$\mathbf{y}^+ = P_{\mathcal{Y}}(\mathbf{y} + \gamma \nabla_{\mathbf{y}} \mathcal{L}(\mathbf{x}, \mathbf{y}))$$

Can oscillate so use averaging:

$$\frac{1}{T} \sum_{t=1}^T (\mathbf{x}^{(t)}, \mathbf{y}^{(t)}) \xrightarrow{T \rightarrow \infty} (\mathbf{x}^*, \mathbf{y}^*)$$

projected extra-gradient

$$\bar{\mathbf{x}} = P_{\mathcal{X}}(\mathbf{x} - \gamma \nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}, \mathbf{y}))$$

$$\bar{\mathbf{y}} = P_{\mathcal{Y}}(\mathbf{y} + \gamma \nabla_{\mathbf{y}} \mathcal{L}(\mathbf{x}, \mathbf{y}))$$

$$\mathbf{x}^+ = P_{\mathcal{X}}(\mathbf{x} - \gamma \nabla_{\mathbf{x}} \mathcal{L}(\bar{\mathbf{x}}, \bar{\mathbf{y}}))$$

$$\mathbf{y}^+ = P_{\mathcal{Y}}(\mathbf{y} + \gamma \nabla_{\mathbf{y}} \mathcal{L}(\bar{\mathbf{x}}, \bar{\mathbf{y}}))$$

More stable and usually **faster**:

$$(\mathbf{x}^{(T)}, \mathbf{y}^{(T)}) \xrightarrow{T \rightarrow \infty} (\mathbf{x}^*, \mathbf{y}^*)$$

SP-FW

Algorithm Saddle point FW algorithm

- 1: Let $\mathbf{z}^{(0)} = (\mathbf{x}^{(0)}, \mathbf{y}^{(0)}) \in \mathcal{X} \times \mathcal{Y}$
 - 2: **for** $t = 0 \dots T$ **do**
 - 3: Compute $\mathbf{r}^{(t)} := \begin{pmatrix} \nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}^{(t)}, \mathbf{y}^{(t)}) \\ -\nabla_{\mathbf{y}} \mathcal{L}(\mathbf{x}^{(t)}, \mathbf{y}^{(t)}) \end{pmatrix}$
 - 4: Compute $\mathbf{s}^{(t)} \in \operatorname{argmin}_{\mathbf{z} \in \mathcal{X} \times \mathcal{Y}} \langle \mathbf{z}, \mathbf{r}^{(t)} \rangle$
 - 5: Compute $g_t := \langle \mathbf{z}^{(t)} - \mathbf{s}^{(t)}, \mathbf{r}^{(t)} \rangle$
 - 6: **if** $g_t \leq \epsilon$ **then return** $\mathbf{z}^{(t)}$
 - 7: Let $\gamma = \min(1, \frac{\nu}{C} g_t)$ **or** $\gamma = \frac{2}{2+t}$
 - 8: Update $\mathbf{z}^{(t+1)} := (1 - \gamma)\mathbf{z}^{(t)} + \gamma\mathbf{s}^{(t)}$
 - 9: **end for**
-

proposed by
[Hammond 1984] with
 $O(1/t)$ step-size

30 years old
conjecture for
polytopes!

Convergence

SP extension of FW with
away step:

- **Linear** rate with **adaptive** step size $\gamma_t := \frac{\nu}{2LD^2}g_t$.

$$\min_{s \leq t} g_s \leq \left(1 - \nu^2 \frac{\delta^2 \mu}{D^2 2L}\right)^t$$

- **Sublinear** rate with **universal** step size $\gamma_t := \frac{2}{2+k(t)}$.

Assumptions

Similar assumptions as AFW:

- $\mathcal{L}$ is L -**smooth** and μ -**strongly convex**.
- $\mathcal{X}$ and $\mathcal{Y}$ **polytopes**.

Additional assumption on **bilinearity**:

$$\nu := \frac{1}{2} - \frac{\sqrt{2}\|M\|D}{\mu\delta} > 0$$

Details on the additional assumption

“Strong convexity μ big enough compared to the bilinear coupling $\|M\|$ ”

$$\mathcal{L}(\mathbf{x}, \mathbf{y}) = f(\mathbf{x}) + \mathbf{x}^\top M \mathbf{y} - g(\mathbf{y}).$$

$$D := \max\{\text{diam}(\mathcal{X}), \text{diam}(\mathcal{Y})\}, \quad \delta := \min\{PWidth(\mathcal{X}), PWidth(\mathcal{Y})\}$$

Difficulties for saddle point

Usual descent Lemma:

$$h_{t+1} \leq h_t - \underbrace{\gamma_t g_t}_{\geq 0} + \gamma_t^2 \frac{L \|d^{(t)}\|^2}{2}$$

With γ_t small enough the sequence decreases.

For saddle point problem the Lipschitz gradient property gives

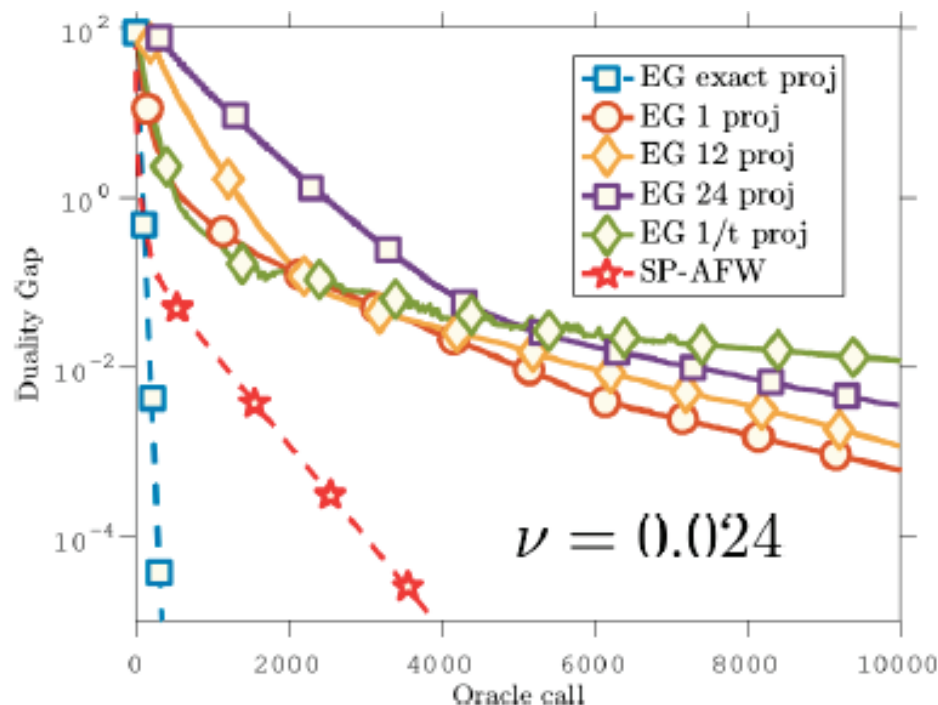
$$\mathcal{L}_{t+1} - \mathcal{L}^* \leq \mathcal{L}_t - \mathcal{L}^* - \underbrace{\gamma_t \left(g_t^{(x)} - g_t^{(y)} \right)}_{\text{arbitrary sign}} + \gamma_t^2 \frac{L \|d^{(t)}\|^2}{2}.$$

- ▶ Cannot control the oscillation of the sequence.
- ▶ Must introduce other quantities to establish convergence.

Toy experiments

■ SP-AFW vs. extragradient with approx. projection

[He & Harchaoui
NIPS 2015]



Theoretical:

$$\gamma_t = \frac{\nu}{C} g_t$$

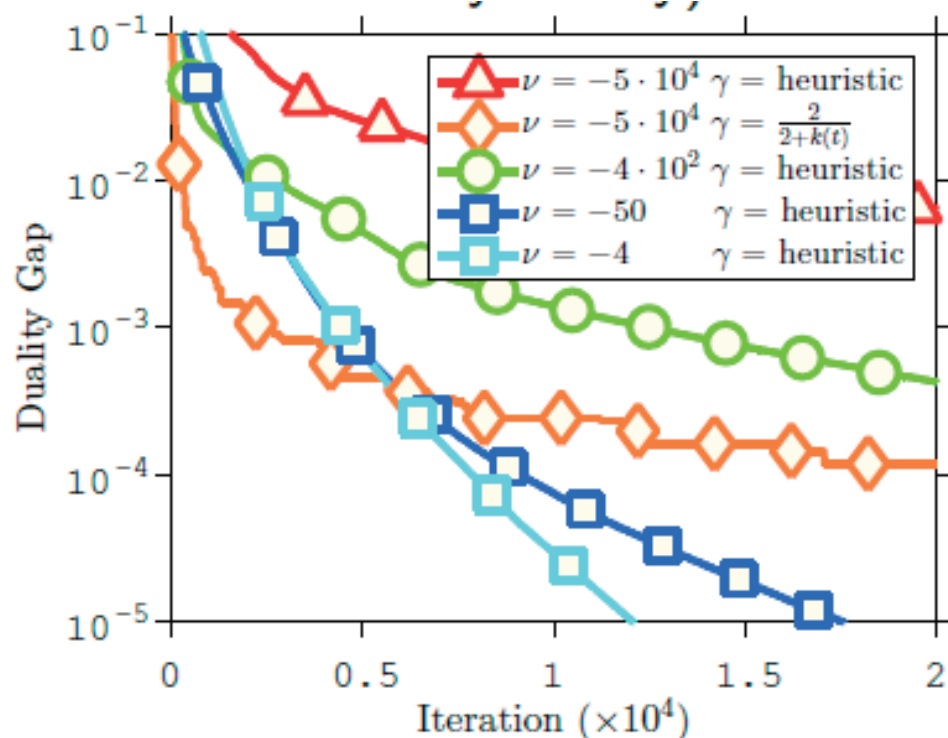
$$\mathcal{L}(x, y) := \frac{\mu}{2} \|x - x^*\|_2^2 + (x - x^*)^\top M(y - y^*) - \frac{\mu}{2} \|y - y^*\|_2^2$$

$$\mathcal{X} = \mathcal{Y} := [0, 1]^d \quad d = 30 \quad C := 2LD^2 \quad L = \mu$$

Toy experiments

- SP-AFW with heuristic step sizes when $\nu < 0$

(not covered by theory)



Heuristic:

$$\gamma_t = \frac{g_t}{C + 2 \frac{\|M\|^2 D^2}{\mu}}$$

$$\mathcal{L}(x, y) := \frac{\mu}{2} \|x - x^*\|_2^2 + (x - x^*)^\top M (y - y^*) - \frac{\mu}{2} \|y - y^*\|_2^2$$

$$\mathcal{X} = \mathcal{Y} := [0, 1]^d \quad d = 30 \quad C := 2LD^2 \quad L = \mu$$

Discussion

- also linear convergence of SP-FW on product of **strongly convex** sets
- for bilinear objective, Karlin's conjecture [1960] gives $O(1/\sqrt{t})$ rate (only empirical so far)
- more general convergence still open!

Thank you! Any question?

constants...

$$C_f := \sup_{\substack{\mathbf{x}, \mathbf{s} \in \mathcal{D}, \\ \gamma \in [-1, 1], \\ \mathbf{y} = \mathbf{x} + \gamma(\mathbf{s} - \mathbf{x})}} \frac{2}{\gamma^2} \left(f(\mathbf{y}) - f(\mathbf{x}) - \langle \nabla f(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle \right)$$

$$\frac{\mu}{2} \|\mathbf{y} - \mathbf{x}\|^2 \leq f(\mathbf{y}) - f(\mathbf{x}) - \langle \nabla f(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle \leq \frac{L}{2} \|\mathbf{y} - \mathbf{x}\|^2$$

$$\mu_f := \inf_{\substack{\mathbf{x}, \mathbf{y} \in \mathcal{D} \\ \text{s.t. } \langle \nabla f(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle < 0}} \frac{2}{\gamma^2(\mathbf{x}, \mathbf{y})} \left(f(\mathbf{y}) - f(\mathbf{x}) - \langle \nabla f(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle \right)$$

$$\gamma(\mathbf{x}, \mathbf{y}) := \frac{\langle \nabla f(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle}{\langle \nabla f(\mathbf{x}), \mathbf{s}_f(\mathbf{x}) - \mathbf{v}_f(\mathbf{x}) \rangle}$$

towards
vertex

away
vertex

Part I: Adaptive quadrature rule with Frank-Wolfe optimization

- Approximating integrals: $\int_{\mathcal{X}} f(x)p(x)dx \approx \frac{1}{N} \sum_{i=1}^N f(x^{(i)})$

for **fixed** p , and multiple f 's in a RKHS $\mathcal{H}$

- Random sampling $x^{(i)} \sim p(x)$ yields $O(1/\sqrt{N})$ error
- Kernel herding [Chen et al. 10] (can) yield $O(1/N)$ error!
(need finite dim. $\mathcal{H}$)  (like quasi-MC)
- -> generalized to FW optimization [Bach et al. 12] and could even get $O(e^{-cN})$ error

- Trick: run Frank-Wolfe optimization on dummy objective:

where $\mathcal{M} = \text{cl-conv}(\Phi(\mathcal{X}))$

is the marginal polytope

$$\min_{g \in \mathcal{M}} \frac{1}{2} \|g - \mu(p)\|_{\mathcal{H}}^2$$

and $\mu(p) = \mathbb{E}_{p(x)} \Phi(x)$ is the *mean map*

 representer: $k(x, \cdot) \in \mathcal{H}$

Approx. integrals in RKHS

- Why? Well, controlling **moment discrepancy** $\|\mu(\hat{p}) - \mu(p)\|_{\mathcal{H}}$ is enough to control **error of integrals** in RKHS $\mathcal{H}$:

- Reproducing property: $f \in \mathcal{H} \Rightarrow f(x) = \langle f, \Phi(x) \rangle$

- Define *mean map*: $\mu(p) = \mathbb{E}_{p(x)} \Phi(x)$

- Want to approximate integrals of the form:

$$\mathbb{E}_{p(x)} f(x) = \mathbb{E}_{p(x)} \langle f, \Phi(x) \rangle = \langle f, \mu(p) \rangle$$

- Use weighted sum to get approximated mean: $\hat{p} = \sum_{i=1}^N w_t^{(i)} \delta_{x^{(i)}}$

$$\mu(\hat{p}) = \mathbb{E}_{\hat{p}(x)} \Phi(x) = \sum_{i=1}^N w^{(i)} \Phi(x^{(i)}) \Rightarrow \mathbb{E}_{\hat{p}(x)} f(x) = \sum_{i=1}^N w^{(i)} f(x^{(i)})$$

- Approximation error is then bounded by:

$$|\mathbb{E}_{p(x)} f(x) - \mathbb{E}_{\hat{p}(x)} f(x)| \leq \|f\|_{\mathcal{H}} \|\mu(p) - \mu(\hat{p})\|_{\mathcal{H}}$$

FW quadrature

- Run Frank-Wolfe optimization on dummy objective:

$$\min_{g \in \mathcal{M}} \frac{1}{2} \|g - \mu(p)\|_{\mathcal{H}}^2$$

where $\mathcal{M} = \text{cl-conv}(\Phi(\mathcal{X}))$
is the marginal polytope

and $\mu(p) = \mathbb{E}_{p(x)} \Phi(x)$ is the *mean map*

FW-Quad

repeat:

input: p

1) FW search:

$$x^{(k+1)} = \arg \min_{x \in \mathcal{X}} g_k(x) - \mu(p)(x)$$

2) convex combo:

$$g_{k+1} = (1 - \gamma_k) g_k + \gamma_k \Phi(x^{(k+1)})$$

e.g. minimum of a difference of
mixture of Gaussian bumps!
(for a Gaussian kernel)

output:

$$\begin{aligned} g_N &= \sum_{i=1}^N w^{(i)} \Phi(x^{(i)}) \\ &= \mu(\hat{p}_N) \end{aligned}$$

- Requirements: can compute $\mu(p)$

+ approx. solve (1) -> use exhaustive search through M random samples from p
-> **super-samples** selection [Chen et al. 10]