

LLOYD’S K -MEANS CLUSTERING ALGORITHM IS FRANK-WOLFE IN DISGUISE

Michael Pokojov^{*}, J. Marcus Jobe^{**} and Simon Lacoste-Julien[†]

^{*}Department of Mathematics & Statistics and School of Data Science, Old Dominion University, Norfolk, Virginia

^{**}Department of Information Systems and Analytics, Miami University, Oxford, Ohio

[†]Canada CIFAR AI Chair, Department of Computer Science and Operations Research & Mila, Université de Montréal, Montréal (QC) Canada

jobejm@miamioh.edu

Introduction

Lloyd’s K -means algorithm, also known as naïve K -means, is a widely used *ad hoc* optimization heuristic, designed to minimize the sum of squared errors (SSE) across all K -partitions of a dataset via iterative cluster refinement. We establish a novel connection between Lloyd’s algorithm and the Frank-Wolfe (FW) algorithm, a prominent first-order method for projection-free optimization. We demonstrate that Lloyd’s algorithm is a special case of FW. Leveraging recent advances in FW methods for concave objectives, we derive a non-asymptotic $\mathcal{O}(1/t)$ convergence rate to a local minimum of the SSE objective. To account for empty clusters, an outcome possible under Lloyd’s greedy assignment, we develop an FW variant for semismooth objectives while retaining the same convergence rate that is solely controlled by the initial SSE value. We illustrate our findings with a simulation study for spherical Gaussian mixtures and a real-world image segmentation data.

Contributions: Main contributions are summarized as follows:

- We show that the Lloyd’s algorithm with greedy cluster assignment is a special case of the FW algorithm with unit step-size for a semismooth concave objective over a polyhedral set. With empty clusters excluded from the feasible set, the objective becomes smooth. This observation allows us to leverage recent convergence results for concave FW to directly establish a non-asymptotic $\mathcal{O}(1/t)$ convergence rate (Yurtsever and Sra 2022).
- To accommodate for the practically relevant possibility of empty clusters, we develop an FW algorithm for general (semismooth) concave objectives with a new FW gap based on Clarke subdifferential. The same non-asymptotic $\mathcal{O}(1/t)$ convergence is recovered that solely depends on the initial global suboptimality (uniformly in sample size n , space dimension d and number of clusters K) and does not involve coresets considerations.

Lloyd’s K -means Clustering

Consider a dataset $\mathcal{D} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\} \subset \mathbb{R}^d$. We seek to minimize the sum of squares errors (SSE), a.k.a. within-cluster sum of squares (WCSS),

$$\text{SSE}(\mathcal{C}) \equiv \sum_{k: C_k \neq \emptyset} \sum_{\mathbf{x} \in C_k} \|\mathbf{x} - \bar{\mathbf{x}}_k\|^2 \quad (1)$$

with $\bar{\mathbf{x}}_k = \frac{1}{|C_k|} \sum_{\mathbf{x} \in C_k} \mathbf{x}$ over all K -partitions $\mathcal{C} = (C_1, C_2, \dots, C_K)$ of $\mathcal{D}$.

Algorithm 1 (Lloyd’s K -means algorithm).

- 1: Let the seeds $\boldsymbol{\mu}_1^{(-1)}, \dots, \boldsymbol{\mu}_K^{(-1)} \in \mathbb{R}^d$ be given.
- 2: Initialize:
- 3: $C_k^{(0)} := \{\mathbf{x}_i \mid \|\mathbf{x}_i - \boldsymbol{\mu}_k^{(-1)}\| \leq \|\mathbf{x}_i - \boldsymbol{\mu}_{k'}^{(-1)}\| \text{ for } 1 \leq k' \leq K\}$,
- 4: $\boldsymbol{\mu}_k^{(0)} := \begin{cases} \frac{1}{|C_k^{(0)}|} \sum_{\mathbf{x}_i \in C_k^{(0)}} \mathbf{x}_i, & \text{if } C_k^{(0)} \neq \emptyset, \\ \boldsymbol{\mu}_k^{(0)} := \boldsymbol{\mu}_k^{(-1)}, & \text{otherwise.} \end{cases}$
- 5: **for** $t = 0, \dots, T-1$ **do**
- 6: **for** $k = 1, \dots, K$ **do**
- 7: Let $C_k^{(t+1)} := \{\mathbf{x}_i \mid \|\mathbf{x}_i - \boldsymbol{\mu}_k^{(t)}\| \leq \|\mathbf{x}_i - \boldsymbol{\mu}_{k'}^{(t)}\| \text{ for } 1 \leq k' \leq K\}$.
- 8: Compute $\boldsymbol{\mu}_k^{(t+1)} := \begin{cases} \frac{1}{|C_k^{(t+1)}|} \sum_{\mathbf{x}_i \in C_k^{(t+1)}} \mathbf{x}_i, & \text{if } C_k^{(t+1)} \neq \emptyset, \\ \boldsymbol{\mu}_k^{(t)}, & \text{otherwise.} \end{cases}$
- 9: Let $\Delta \text{SSE}_t := \text{SSE}(\mathcal{C}^{(t)}) - \text{SSE}(\mathcal{C}^{(t+1)})$.
- 10: **if** $\Delta \text{SSE}_t \leq \epsilon$ **then**
- 11: **return** $C_1^{(t)}, \dots, C_K^{(t)}$ (and $\boldsymbol{\mu}_1^{(t)}, \dots, \boldsymbol{\mu}_K^{(t)}$)
- 12: **end if**
- 13: **end for**
- 14: **end for**
- 15: **return** $C_1^{(T)}, \dots, C_K^{(T)}$ (and $\boldsymbol{\mu}_1^{(T)}, \dots, \boldsymbol{\mu}_K^{(T)}$).

One-Hot Encoding and Extension

The K -means SSE objective reads as

$$f(\mathbf{w}) = \sum_{k=1}^K \sum_{i=1}^n w_{ik} \|\mathbf{x}_i - \mathbf{x}_k^w\|^2 \quad (2)$$

$$\text{with } \mathbf{x}_k^w := \frac{1}{n_k^w} \sum_{i=1}^n w_{ik} \mathbf{x}_i \quad \text{and} \quad n_k^w := \sum_{i=1}^n w_{ik}$$

using the usual ‘one-hot’ encoding $\mathbf{w} \in \{0, 1\}^{n \times K}$ such that $w_{ik} = 1$ if the i -th data point $\mathbf{x}_i \in \mathbb{R}^d$ belongs to the k -th cluster.

The SSE objective $f(\mathbf{w})$ can be continuously extended to $\mathcal{M} = \text{conv}(\mathcal{M}_0)$ with

$$\mathcal{M}_0 = \left\{ \mathbf{w} \in \{0, 1\}^{n \times K} \mid \sum_{k=1}^K w_{ik} = 1, i = 1, \dots, n \right\} \quad (3)$$

by letting $\sum_{i=1}^n w_{ik} \|\mathbf{x}_i - \bar{\mathbf{x}}_k^w\|^2 := 0$ for $n_k^w = 0$.

Frank-Wolfe Algorithm for Semismooth Concave Objectives

Despite being continuous, the Lloyd’s SSE objective fails to be smooth in $\mathcal{M}$. Therefore, we extend the recent results of Yurtsever and Sra 2022 to semismooth concave objectives.

Let $\mathcal{M} \subset \mathbb{R}^p$ be a non-empty convex compact set. We consider a general constrained optimization problem

$$\min_{\mathbf{x} \in \mathcal{M}} f(\mathbf{x}), \quad (4)$$

where f is still concave, but only semismooth. Let

$$\text{LMO}_{\mathcal{M}}(\mathbf{r}) := \arg \min_{\mathbf{s} \in \mathcal{M}} \langle \mathbf{s}, \mathbf{r} \rangle \quad (5)$$

denote the linear minimization oracle (LMO) at $\mathbf{r} \in \mathbb{R}^p$.

Suppose the function f possesses a (set-valued) Clarke subdifferential $\partial f(\mathbf{x})$ at any $\mathbf{x} \in \mathcal{M}$. Further, the set $\partial f(\mathbf{x}) \neq \emptyset$ is compact and convex for any $\mathbf{x} \in \mathcal{M}$. Algorithm 2 below protocols our semismooth FW with unit step-size. Note that for smooth objectives, it automatically reduces to the smooth version since $\partial f(\mathbf{x}) = \{\nabla f(\mathbf{x})\}$.

Algorithm 2 (Semismooth FW for concave functions).

- 1: Let $\mathbf{x}^{(0)} \in \mathcal{M}$.
- 2: **for** $t = 0, \dots, T-1$ **do**
- 3: Choose arbitrary $\mathbf{r}^{(t)} \in \partial f(\mathbf{x}^{(t)})$.
- 4: Compute $\mathbf{s}^{(t)} := \text{LMO}_{\mathcal{M}}(\mathbf{r}^{(t)})$.
- 5: Define the FW update direction $\mathbf{d}_t := \mathbf{s}^{(t)} - \mathbf{x}^{(t)}$.
- 6: Compute the FW gap $g_t := \langle \mathbf{d}_t, -\mathbf{r}^{(t)} \rangle$.
- 7: **if** $g_t \leq \epsilon$ **then**
- 8: **return** $\mathbf{x}^{(t)}$
- 9: **end if**
- 10: Select step-size $\gamma_t = 1$.
- 11: Update $\mathbf{x}^{(t+1)} := \mathbf{x}^{(t)} + \gamma_t \mathbf{d}_t \equiv \mathbf{s}^{(t)}$.
- 12: **end for**
- 13: **return** $\mathbf{x}^{(T)}$.

Before discussing non-asymptotic convergence of Algorithm 2, we introduce a suitable generalization of the FW gap to be used as a stationarity measure.

Cnt’d

Recall that $\mathbf{x} \in \mathcal{M}$ is referred to as Clarke-stationary if

$$0 = \max_{\mathbf{y} \in \mathcal{M}} \langle \mathbf{y} - \mathbf{x}, -\mathbf{r} \rangle = 0 \quad (6)$$

for at least one $\mathbf{r} \in \partial f(\mathbf{x})$. Similar to Liu et al. 2024, we introduce the Clarke version of FW gap

$$g_{\mathcal{C}}(\mathbf{x}) := \min_{\mathbf{r} \in \partial f(\mathbf{x})} \max_{\mathbf{y} \in \mathcal{M}} \langle \mathbf{y} - \mathbf{x}, -\mathbf{r} \rangle. \quad (7)$$

By compactness, the latter expression is well-defined and finite. Further, for any $\mathbf{r} \in \partial f(\mathbf{x})$,

$$\max_{\mathbf{y} \in \mathcal{M}} \langle \mathbf{y} - \mathbf{x}, -\mathbf{r} \rangle \geq \langle \mathbf{y} - \mathbf{x}, -\mathbf{r} \rangle|_{\mathbf{y}=\mathbf{x}} = 0.$$

Thus, we trivially have $g_{\mathcal{C}}(\mathbf{x}) \geq 0$. By Equation (6), $g_{\mathcal{C}}(\mathbf{x}) = 0$ if and only if $\mathbf{x}$ is Clarke stationary.

Similar to Yurtsever and Sra 2022, we can show:

Theorem 1 (Semismooth FW on concave objectives). Consider running the FW Algorithm 2 with unit step-size for the optimization problem in Equation (4). Then the minimal FW gap $\tilde{g}_t := \min_{0 \leq \tau \leq t} g_{\tau}$ encountered by the iterates after t iterations satisfies:

$$\tilde{g}_t \leq \frac{h_0}{t+1} \quad \text{for } 0 \leq t \leq T-1 \quad (8)$$

where $h_0 := f(\mathbf{x}^{(0)}) - \min_{\mathbf{x} \in \mathcal{M}} f(\mathbf{x})$ is the initial global suboptimality.

Lloyd’s K -Means as Frank-Wolfe for Concave Objective

Lemma 1 characterizes f as a concave function in the interior of $\mathcal{M} = \text{conv}(\mathcal{M}_0)$ (cf. Equation (3)). Since f is continuous in $\mathcal{M}$, due to Bauer’s maximum principle, f attains a global minimum at one of the extreme points $\mathcal{M}_0$. This also applies to any local minimum that either must be an extreme point itself or be ‘‘tied’’ with one.

Lemma 1. For any $\mathbf{w} \in \mathcal{M}$ with $n_k^w > 0$ for all k , ∇f and $\nabla^2 f$ can be expressed as

$$\begin{aligned} \partial_{w_{ik}} f(\mathbf{w}) &= \|\mathbf{x}_i - \bar{\mathbf{x}}_k^w\|^2 \quad \text{and} \\ \partial_{w_{ik}} \partial_{w_{jl}} f(\mathbf{w}) &= -\frac{2}{n_k^w} \cdot \mathbf{1}_{\{k=l\}} \cdot (\mathbf{x}_i - \bar{\mathbf{x}}_k^w) \cdot (\mathbf{x}_j - \bar{\mathbf{x}}_k^w). \end{aligned}$$

Moreover, the Hessian $\nabla^2 f$ is negative semidefinite.

Unfortunately, ∇f can not be continuously extended to $\mathcal{M}$ so that the smooth FW algorithm cannot be applied to f over $\mathcal{M}$. However, the following theorem holds for a semismooth extension of f .

Theorem 2. Let $\mathbf{w}^{(t)}$, $0 \leq t \leq T$, be produced by Lloyd’s K -means Algorithm 1. Then the membership vector $\mathbf{w}^{(t)}$ coincides with the t -th iterate of the FW Algorithm 2 applied to minimizing f over $\mathcal{M}^{\text{adm}}$ such that

$$\tilde{g}_t = \min_{0 \leq \tau \leq t} g_{\tau} = \min_{0 \leq \tau \leq t} \Delta \text{SSE}_{\tau} \leq \frac{h_0}{t+1}$$

for $t = 0, \dots, T-1$, where $h_0 = f(\mathbf{x}^{(0)}) - \min_{\mathbf{x} \in \mathcal{M}} f(\mathbf{x})$ and $\Delta \text{SSE}_t = f(\mathbf{w}^{(t)}) - f(\mathbf{w}^{(t+1)})$.

Simulation Study

We illustrate our theoretical results with a simulation study. Since the iteration becomes stationary in finite time, we focus on the *worst-case* performance over randomly selected seeds – for a fixed dataset or over multiple datasets sampled from a given population. Letting g_t^{WC} denote the worst-case (i.e., largest) FW gap (over all replications) at iteration t , the convergence rate of g_t^{WC} was estimated using OLS applied to the regression model $\log(g_t^{\text{WC}}) = \beta_0 + \beta_1 \log(t) + \varepsilon$ with i.i.d. normal errors ε . Preasymptotic regime was defined as bottom 2/3 of the $\log(t)$ range. Theoretical values were assumed $\beta_1 = -1$ and $\beta_0 = (n-1) \text{tr}(\Sigma)$, where Σ is estimated using sample covariance. Synthetic data were simulated from the

Gaussian mixture $\sum_{k=1}^K \pi_k \varphi_k(\mathbf{x} | \boldsymbol{\mu}_k, \sigma_k^2 \mathbf{I})$ with $\pi_k = \frac{1}{K}$ and $\sigma_k = 1$ (cf. `make_blobs()` of Python’s `scikit-learn`). For each (n, d, K) with $n = 500, 1000, 5000$, $d = 2, 5, 10$ and $K = 5, 10, 20, 50$, we performed 10,000 Monte Carlo (MC) simulations to estimate g_t^{WC} .

Figure 1 displays the worst-case FW gap plot vs t . Applying an upper-sided Student’s t -test of $H_0 : \beta_1 = -1$ vs $H_1 : \beta_1 > -1$ across all scenarios considered, the p -values never dropped below $1 - (1 - 0.05)^{1/36} \approx 0.0014$, which is the Šidák-adjusted value to maintain a family-wise 0.05 test size, the composite null hypothesis failed to be rejected, corroborating that the worst-case preasymptotic rate was never slower than $\mathcal{O}(1/t)$.

Example: Image Segmentation Dataset

Consider the image segmentation dataset from UCI Machine Learning Repository. Pooling training and test data, each of 2,310 instances is classified as one of the seven classes: BRICKFACE, CEMENT, FOLIAGE, GRASS, PATH, SKY and WINDOW. Keeping only continuous features (columns 6 through 19), we end up with 14 variables ($d = 14$). Pokojov and Jobe 2022 empirically demonstrated that the 7-component mixture is non-Gaussian. An MC simulation of size 50,000 was performed for $K = 7$. The results are displayed in Figure 2. The slope estimate is reported along with the standard error (in parentheses). Based on an upper-sided Student’s t -test with $\hat{\beta}_1 = -1.2451$ and $\text{se}(\hat{\beta}_1) = 0.1228$, the null hypothesis $H_0 : \beta_1 = -1$ failed to be rejected vs $H_1 : \beta_1 > -1$ at $\alpha = 0.05$, corroborating an $\mathcal{O}(1/t)$ worst-case rate.

Summary and Conclusions

We emphasize the importance of non-asymptotic convergence guarantees for Lloyd’s K -means – or any other machine learning algorithm for that matter – that hold uniformly with respect to the sample size, space dimension, hyperparameters, initial seeding or optimization ‘‘intricacies’’ like the curvature constant or geometric aspects of the domain. Building upon the work of (Yurtsever and Sra 2022) in the context of convex-concave procedure, we proved that Lloyd’s algorithm is a special case of FW. In doing so, we not only were able to perfectly carry over their results to Lloyd’s K -means clustering but developed an FW variant for general semismooth concave objectives and proved an $\mathcal{O}(1/t)$ convergence rate. Owing to concavity, we could employ the less technical Clarke’s instead of Goldstein’s subdifferential prevalent in the nonsmooth literature allowing for better transparency and broader accessibility.

References

- Liu, Z., C. Chen, L. Luo, and B.K.H. Low. 2024. ‘‘Zeroth-Order Methods for Constrained Nonconvex Nonsmooth Stochastic Optimization.’’ In *Proc of the 41st ICML*, 235:30842–30872. PMLR.
- Pokojov, M., and J. M. Jobe. 2022. ‘‘A robust deterministic affine-equivariant algorithm for multivariate location and scatter.’’ *Computational Statistics & Data Analysis* 172 (107475): 1–24.
- Yurtsever, Alp, and Suvrit Sra. 2022. ‘‘CCCP is Frank-Wolfe in Disguise.’’ In *Advances in Neural Information Processing Systems (NeurIPS)*, 35:29989–30001.

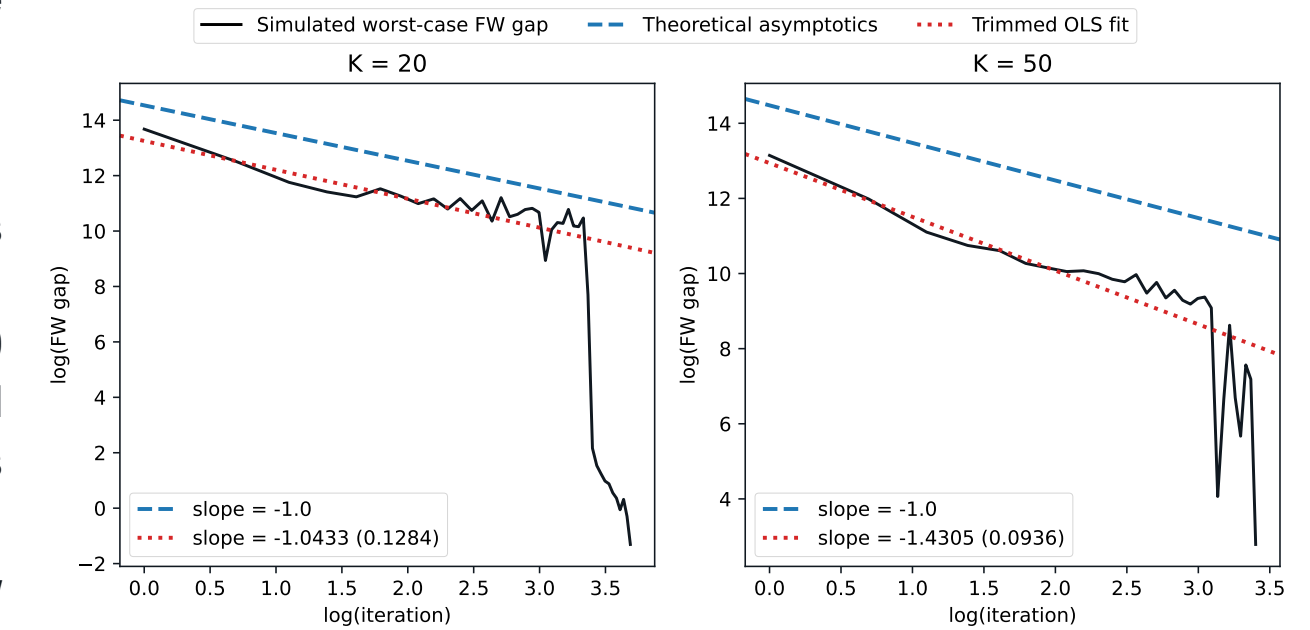


Fig. 1: Worst-case convergence rate on synthetic ‘‘blobs’’ data for $n = 5,000$ and $d = 10$.

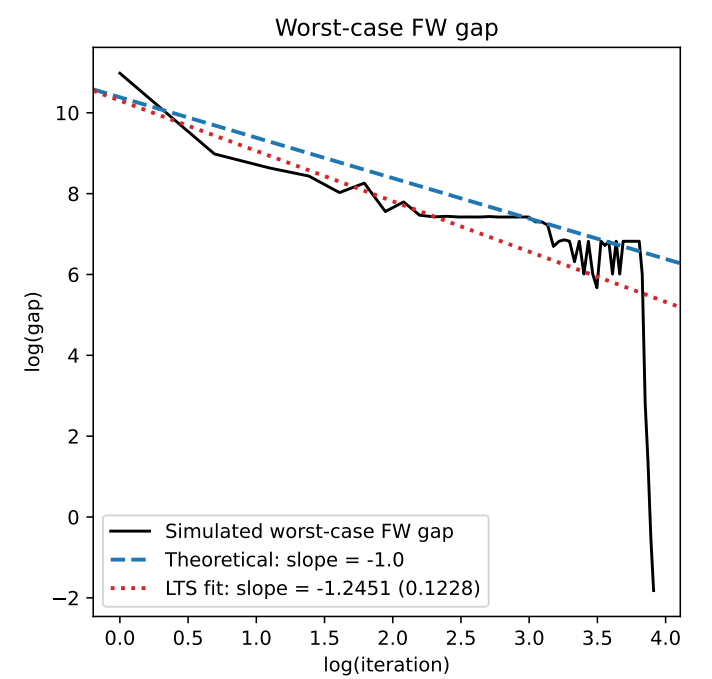


Fig. 2: Worst-case convergence rate.