
A Unified Approach for Maximizing Continuous DR-submodular Functions

Mohammad Pedramfar
Purdue University
mpedramf@purdue.edu

Christopher John Quinn
Iowa State University
cjquinn@iastate.edu

Vaneet Aggarwal
Purdue University
vaneet@purdue.edu

Abstract

This paper presents a unified approach for maximizing continuous DR-submodular functions that encompasses a range of settings and oracle access types. Our approach includes a Frank-Wolfe type offline algorithm for both monotone and non-monotone functions, with different restrictions on the general convex set. We consider settings where the oracle provides access to either the gradient of the function or only the function value, and where the oracle access is either deterministic or stochastic. We determine the number of required oracle accesses in all cases. Our approach gives new/improved results for nine out of the sixteen considered cases, avoids computationally expensive projections in three cases, with the proposed framework matching performance of state-of-the-art approaches in the remaining four cases. Notably, our approach for the stochastic function value-based oracle enables the first regret bounds with bandit feedback for stochastic DR-submodular functions.

1 Introduction

The problem of optimizing DR-submodular functions over a convex set has attracted considerable interest in both the machine learning and theoretical computer science communities [2, 5, 16, 26]. This is due to its many practical applications in modeling real-world problems, such as influence/revenue maximization, facility location, and non-convex/non-concave quadratic programming [3, 9, 18, 15, 20], as well as more recently identified applications like serving heterogeneous learners under networking constraints [20] and joint optimization of routing and caching in networks [21].

Numerous studies investigated developing approximation algorithms for constrained DR-submodular maximization, utilizing a variety of algorithms and proof analysis techniques. These studies have addressed both monotone and non-monotone functions and considered various types of constraints on the feasible region. The studies have also considered different types of oracles—gradient oracles and value oracles, where the oracles could be exact (deterministic) or stochastic. Lastly, for some of the aforementioned offline problem settings, some studies have also considered analogous online optimization problem settings as well, where performance is measured in regret over a horizon. This paper aims to unify the disparate offline problems under a single framework by providing a comprehensive algorithm and analysis approach that covers a broad range of setups. By providing a unified framework, this paper presents novel results for several cases where previous research was either limited or non-existent, both for offline optimization problems and extensions to related stochastic online optimization problems.

This paper presents a Frank-Wolfe based meta-algorithm for (offline) constrained DR-submodular maximization where we could only query within the constraint set, with sixteen variants for sixteen problem settings. The algorithm is designed to handle settings where (i) the function is monotone

This work was supported in part by the National Science Foundation under grants CCF-2149588 and CCF-2149617, and Cisco, Inc.

Table 1: Offline DR-submodular optimization results.

F	Set	Oracle	Setting	Reference	Appx.	Complexity
Monotone	$0 \in \mathcal{K}$	∇F	det.	[4], (*)	$1 - 1/e$	$O(1/\epsilon)$
			stoch.	[23], (*) [32] †	$1 - 1/e$ $1 - 1/e$	$O(1/\epsilon^3)$ $O(1/\epsilon^2)$
		F	det.	This paper	$1 - 1/e$	$O(1/\epsilon^3)$
			stoch.	This paper	$1 - 1/e$	$O(1/\epsilon^5)$
	general†	∇F	det.	[16] ‡ This paper	$1/2$ $1/2$	$O(1/\epsilon)$ $\tilde{O}(1/\epsilon)$
			stoch.	[16] ‡ This paper	$1/2$ $1/2$	$O(1/\epsilon^2)$ $\tilde{O}(1/\epsilon^3)$
		F	det.	This paper	$1/2$	$\tilde{O}(1/\epsilon^3)$
			stoch.	This paper	$1/2$	$\tilde{O}(1/\epsilon^5)$
	Non-Monotone	∇F	det.	[3], (*)	$1/e$	$O(1/\epsilon)$
			stoch.	[23], (*)	$1/e$	$O(1/\epsilon^3)$
		F	det.	This paper	$1/e$	$O(1/\epsilon^3)$
			stoch.	This paper	$1/e$	$O(1/\epsilon^5)$
		general	det.	[12] [11] [10], (*)	$\frac{1-h}{3\sqrt{3}}$ $\frac{1-h}{4}$ $\frac{1-h}{4}$	$O(e^{\sqrt{dL/\epsilon}})$ $O(e^{\sqrt{dL/\epsilon}})$ $O(1/\epsilon)$
					$\frac{1-h}{4}$	$O(1/\epsilon^3)$
			stoch.	This paper	$\frac{1-h}{4}$	$O(1/\epsilon^3)$
		F	det.	This paper	$\frac{1-h}{4}$	$O(1/\epsilon^3)$
			stoch.	This paper	$\frac{1-h}{4}$	$O(1/\epsilon^5)$

This table compares the different results for the number of oracle calls (complexity) *within the feasible set* for DR-submodular maximization. Shaded rows indicate problem settings for which our work has the **first guarantees** or **beats the SOTA**. The rows marked with a blue star (*) correspond to cases where Algorithm 2 **generalizes the corresponding algorithm** and therefore has the same performance. The different columns enumerate properties of the function, the convex feasible region (downward-closed, includes the origin, or general), and the oracle, as well as the approximation ratios and oracle complexity (the number of queries needed to achieve the stated approximation ratio with at most $\epsilon > 0$ additive error). We have $h := \min_{\mathbf{x} \in \mathcal{K}} \|\mathbf{x}\|_\infty$. (See Appendix B regarding [23] and [32]). † when the oracle can be queried for any points in $[0, 1]^d$ (even outside the feasible region $\mathcal{K}$), the problem of optimizing monotone DR-submodular functions over a general convex set simplifies — [4] and [23] achieve the same ratios and complexity bounds as listed above for $0 \in \mathcal{K}$; [8] can achieve an approximation ratio of $1 - 1/e$ with the $O(1/\epsilon^3)$ and $O(1/\epsilon^5)$ complexity for exact and stochastic value oracles respectively. ‡ [16] and [32] use gradient ascent, requiring potentially computationally expensive projections.

or non-monotone, (ii) the feasible region is a downward-closed (d.c.) set (extended to include 0 for monotone functions) or a general convex set, (iii) gradient or value oracle access is available, and (iv) the oracle is exact or stochastic. Table 1 enumerates the cases and corresponding results on oracle complexity (further details are provided in Appendix A). We derive the first oracle complexity guarantees for nine cases, derive the oracle complexity in three cases where previous result had a computationally expensive projection step [32, 16] (and we obtain matching complexity in one of these), and obtain matching guarantees in the remaining four cases. In addition to proving approximation ratios and oracle complexities for several (challenging) settings that are the first or improvements over the state of the art, the *technical novelties of our approach* include:

- (i) A new construction procedure of a shrunk constraint set that allows us to work with lower dimensional feasible sets when given a value oracle, resulting in the first results on general lower dimensional feasible sets given a value oracle.
- (ii) The first Frank-Wolfe type algorithm for analyzing monotone functions over a general convex set for any type of oracle, where only feasible points can be queried.
- (iii) Shedding light on a previously unexplained gap in approximation guarantees for monotone DR-submodular maximization. Specifically, by considering the notion of query sets and assuming that the oracles can only be queried within the constraint set, we divide the class of

monotone submodular maximization into monotone submodular maximization over convex sets containing the origin and monotone submodular maximization over general convex sets. Moreover, we conjecture that the $1/2$ approximation coefficient, which has been considered sub-optimal in the literature, is optimal when oracle queries can only be made within the constraint set. (See Appendix B for more details.)

Furthermore, we also consider online stochastic DR-submodular optimization with bandit feedback, where an agent sequentially picks actions (from a convex feasible region), receives stochastic rewards (in expectation a DR-submodular function) but no additional information, and seeks to maximize the expected cumulative reward. Performance is measured against the best action in expectation (or a near-optimal baseline when the offline problem is NP-hard but can be approximated to within α in polynomial time), the difference denoted as expected α -regret. For such problems, when only bandit feedback is available (it is typically a strong assumption that semi-bandit or full-information feedback is available), the agent must be able to learn from stochastic value oracle queries over the feasible actions action. By designing offline algorithms that only query feasible points, we made it possible to convert those offline algorithms into online algorithms. In fact, because of how we designed the offline algorithms, we are able to access them in a black-box fashion for online problems when only bandit feedback is available. Note that previous works on DR-submodular maximization with bandit feedback in monotone settings (e.g. [31], [25] and [30]) explicitly assume that the convex set contains the origin.

For each of the offline setups, we extend the offline algorithm (the respective variants for stochastic value oracle) and oracle query guarantees to provide algorithms and α -regret bounds in the bandit feedback scenario. Table 2 enumerates the problem settings and expected regret bounds with bandit and semi-bandit feedback. The key contributions of this work can be summarized as follows:

1. This paper proposes a unified approach for maximizing continuous DR-submodular functions in a range of settings with different oracle access types, feasible region properties, and function properties. A Frank-Wolfe based algorithm is introduced, which compared to SOTA methods for each of the sixteen settings, achieves the best-known approximation coefficients for each case while providing (i) the first guarantees in nine cases, (ii) reduced computational complexity by avoiding projections in three cases, and (iii) matching guarantees in remaining four cases.

2. In particular, this paper gives the first results on offline DR-submodular maximization (for both monotone and non-monotone functions) over general convex sets and even for downward-closed convex sets, when only a value oracle is available over the feasible set. Most prior works on offline DR-submodular maximization require access to a gradient oracle.

3. The results, summarized in Table 2, are presented with two feedback models—bandit feedback where only the (stochastic) reward value is available and semi-bandit feedback where a single stochastic sample of the gradient at the location is provided. This paper presents the first regret analysis with bandit feedback for stochastic DR-submodular maximization for both monotone and non-monotone functions. For semi-bandit feedback case, we provide the first result in one case, improve the state of the art result in two cases, and gives the result without computationally intensive projections in one case.

Table 2: Online stochastic DR-submodular optimization.

F	Set	Feedback	Reference	Coef. α	α -Regret
Monotone	$0 \in \mathcal{K}$	∇F	[6]†, This paper	$1/e$	$O(T^{2/3})$
		F	This paper	$1 - 1/e$	$O(T^{3/4})$
	general	∇F	[16] ‡ This paper	$1/2$	$O(T^{1/2})$
		F	This paper	$1/2$	$\dot{O}(T^{3/4})$
		F	This paper	$1/2$	$\tilde{O}(T^{5/6})$
	Non-mono.	d.c.	∇F	This paper	$1/e$
F			This paper	$1/e$	$O(T^{5/6})$
general		∇F	This paper	$\frac{1-h}{4}$	$O(T^{3/4})$
		F	This paper	$\frac{1-h}{4}$	$O(T^{5/6})$

This table compares the different results for the expected α -regret for online stochastic DR-submodular maximization for the under bandit and semi-bandit feedback. Shaded rows indicate problem settings for which our work has the **first guarantees** or **beats the SOTA**. We have $h := \min_{x \in \mathcal{K}} \|x\|_\infty$. † the analysis in [6] has an error (see the supplementary material for details). ‡ [16] uses gradient ascent, requiring potentially computationally expensive projections.

Related Work: The key related works are summarized in Tables 1 and 2. We briefly discuss some key works here; see the supplementary materials for more discussion. For online DR-submodular optimization with bandit feedback, there has been some prior works in the adversarial setup [31, 33, 25, 30] which are not included in Table 2 as we consider the stochastic setup. [31] considered monotone DR-submodular functions over downward-closed convex sets and achieved $(1 - 1/e)$ -regret of $O(T^{8/9})$ in adversarial setting. This was later improved by [25] to $O(T^{5/6})$. [30] further improved the regret bound to $O(T^{3/4})$. However, it should be noted that they use a convex optimization subroutine at each iteration which could be even more computationally expensive than projection. [33] considered non-monotone DR-submodular functions over downward-closed convex sets and achieved $1/e$ -regret of $O(T^{8/9})$ in adversarial setting.

2 Background and Notation

We introduce some basic notions, concepts and assumptions which will be used throughout the paper. For any vector $\mathbf{x} \in \mathbb{R}^d$, $[\mathbf{x}]_i$ is the i -th entry of $\mathbf{x}$. We consider the partial order on $\mathbb{R}^d$ where $\mathbf{x} \leq \mathbf{y}$ if and only if $[\mathbf{x}]_i \leq [\mathbf{y}]_i$ for all $1 \leq i \leq d$. For two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$, the *join* of $\mathbf{x}$ and $\mathbf{y}$, denoted by $\mathbf{x} \vee \mathbf{y}$ and the *meet* of $\mathbf{x}$ and $\mathbf{y}$, denoted by $\mathbf{x} \wedge \mathbf{y}$, are defined by

$$\mathbf{x} \vee \mathbf{y} := (\max\{[\mathbf{x}]_i, [\mathbf{y}]_i\})_{i=1}^d \quad \text{and} \quad \mathbf{x} \wedge \mathbf{y} := (\min\{[\mathbf{x}]_i, [\mathbf{y}]_i\})_{i=1}^d, \quad (1)$$

respectively. Clearly, we have $\mathbf{x} \wedge \mathbf{y} \leq \mathbf{x} \leq \mathbf{x} \vee \mathbf{y}$. We use $\|\cdot\|$ to denote the Euclidean norm, and $\|\cdot\|_\infty$ to denote the supremum norm. In the paper, we consider a bounded convex domain $\mathcal{K}$ and w.l.o.g. assume that $\mathcal{K} \subseteq [0, 1]^d$. We say that $\mathcal{K}$ is *down-closed* (d.c.) if there is a point $\mathbf{u} \in \mathcal{K}$ such that for all $\mathbf{z} \in \mathcal{K}$, we have $\{\mathbf{x} \mid \mathbf{u} \leq \mathbf{x} \leq \mathbf{z}\} \subseteq \mathcal{K}$. The *diameter* D of the convex domain $\mathcal{K}$ is defined as $D := \sup_{\mathbf{x}, \mathbf{y} \in \mathcal{K}} \|\mathbf{x} - \mathbf{y}\|$. We use $\mathbb{B}_r(\mathbf{x})$ to denote the open ball of radius r centered at $\mathbf{x}$. More generally, for a subset $X \subseteq \mathbb{R}^d$, we define $\mathbb{B}_r(X) := \bigcup_{\mathbf{x} \in X} \mathbb{B}_r(\mathbf{x})$. If A is an affine subspace of $\mathbb{R}^d$, then we define $\mathbb{B}_r^A(X) := A \cap \mathbb{B}_r(X)$. For a function $F : \mathcal{D} \rightarrow \mathbb{R}$ and a set $\mathcal{L}$, we use $F|_{\mathcal{L}}$ to denote the restriction of F to the set $\mathcal{D} \cap \mathcal{L}$. For a linear space $\mathcal{L}_0 \subseteq \mathbb{R}^d$, we use $P_{\mathcal{L}_0} : \mathbb{R}^d \rightarrow \mathcal{L}_0$ to denote the projection onto $\mathcal{L}_0$. We will use $\mathbb{R}_+^d$ to denote the set $\{\mathbf{x} \in \mathbb{R}^d \mid \mathbf{x} \geq \mathbf{0}\}$. For any set $X \subseteq \mathbb{R}^d$, the affine hull of X , denoted by $\text{aff}(X)$, is defined to be the intersection of all affine subsets of $\mathbb{R}^d$ that contain X . The *relative interior* of a set X is defined by

$$\text{relint}(X) := \{\mathbf{x} \in X \mid \exists \varepsilon > 0, \mathbb{B}_\varepsilon^{\text{aff}(X)}(\mathbf{x}) \subseteq X\}.$$

It is well known that for any non-empty convex set $\mathcal{K}$, the set $\text{relint}(\mathcal{K})$ is always non-empty. We will always assume that the feasible set contains at least two points and therefore $\dim(\text{aff}(\mathcal{K})) \geq 1$, otherwise the optimization problem is trivial and there is nothing to solve.

A set function $f : \{0, 1\}^d \rightarrow \mathbb{R}^+$ is called *submodular* if for all $\mathbf{x}, \mathbf{y} \in \{0, 1\}^d$ with $\mathbf{x} \geq \mathbf{y}$, we have

$$f(\mathbf{x} \vee \mathbf{a}) - f(\mathbf{x}) \leq f(\mathbf{y} \vee \mathbf{a}) - f(\mathbf{y}), \quad \forall \mathbf{a} \in \{0, 1\}^d. \quad (2)$$

Submodular functions can be generalized over continuous domains. A function $F : [0, 1]^d \rightarrow \mathbb{R}^+$ is called *DR-submodular* if for all vectors $\mathbf{x}, \mathbf{y} \in [0, 1]^d$ with $\mathbf{x} \leq \mathbf{y}$, any basis vector $\mathbf{e}_i = (0, \dots, 0, 1, 0, \dots, 0)$ and any constant $c > 0$ such that $\mathbf{x} + c\mathbf{e}_i \in [0, 1]^d$ and $\mathbf{y} + c\mathbf{e}_i \in [0, 1]^d$, it holds that

$$F(\mathbf{x} + c\mathbf{e}_i) - F(\mathbf{x}) \geq F(\mathbf{y} + c\mathbf{e}_i) - F(\mathbf{y}). \quad (3)$$

Note that if function F is differentiable then the diminishing-return (DR) property (3) is equivalent to $\nabla F(\mathbf{x}) \geq \nabla F(\mathbf{y})$ for $\mathbf{x} \leq \mathbf{y}$ with $\mathbf{x}, \mathbf{y} \in [0, 1]^d$. A function $F : \mathcal{D} \rightarrow \mathbb{R}^+$ is *G-Lipschitz continuous* if for all $\mathbf{x}, \mathbf{y} \in \mathcal{D}$, $\|F(\mathbf{x}) - F(\mathbf{y})\| \leq G\|\mathbf{x} - \mathbf{y}\|$. A differentiable function $F : \mathcal{D} \rightarrow \mathbb{R}^+$ is *L-smooth* if for all $\mathbf{x}, \mathbf{y} \in \mathcal{D}$, $\|\nabla F(\mathbf{x}) - \nabla F(\mathbf{y})\| \leq L\|\mathbf{x} - \mathbf{y}\|$.

A (possibly randomized) offline algorithm is said to be an α -approximation algorithm (for constant $\alpha \in (0, 1]$) with $\epsilon \geq 0$ additive error for a class of maximization problems over non-negative functions if, for any problem instance $\max_{\mathbf{z} \in \mathcal{K}} F(\mathbf{z})$, the algorithm output $\mathbf{x}$ that satisfies the following relation with the optimal solution $\mathbf{z}^*$

$$\alpha F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{x})] \leq \epsilon, \quad (4)$$

where the expectation is with respect to the (possible) randomness of the algorithm. Further, we assume an oracle that can query the value $F(\mathbf{x})$ or the gradient $\nabla F(\mathbf{x})$. The number of calls to the oracle to achieve the error in (4) is called the *evaluation complexity*.

3 Offline Algorithms and Guarantees

In this section, we consider the problem of maximizing a DR-submodular function over a general convex set in sixteen different cases, enumerated in Table 1. After setting up the problem in Section 3.1, we then explain two key elements of our proposed algorithm when we only have access to a value oracle, (i) the Black Box Gradient Estimate (BBGE) procedure (Algorithm 1) to balance bias and variance in estimating gradients (Section 3.2) and (ii) the construction of a shrunk feasible region to avoid infeasible value oracle queries during the BBGE procedure (Section 3.3). Our main algorithm is proposed in Section 3.4 and analyzed in Section 3.5.

3.1 Problem Setup

We consider a general *non-oblivious* constrained stochastic optimization problem

$$\max_{\mathbf{z} \in \mathcal{K}} F(\mathbf{z}) := \max_{\mathbf{z} \in \mathcal{K}} \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x}; \mathbf{z})} [\hat{F}(\mathbf{z}, \mathbf{x})], \quad (5)$$

where F is a DR-submodular function, and $\hat{F} : \mathcal{K} \times \mathfrak{X} \rightarrow \mathbb{R}$ is determined by $\mathbf{z}$ and the random variable $\mathbf{x}$ which is independently sampled according to $\mathbf{x} \sim p(\mathbf{x}; \mathbf{z})$. We say the oracle has variance σ^2 if $\sup_{\mathbf{z} \in \mathcal{K}} \text{var}_{\mathbf{x} \sim p(\mathbf{x}; \mathbf{z})} [\hat{F}(\mathbf{z}, \mathbf{x})] = \sigma^2$. In particular, when $\sigma = 0$, then we say we have access to an exact (deterministic) value oracle. We will use $\hat{F}(\mathbf{z})$ to denote the random variables $\hat{F}(\mathbf{z}, \mathbf{x})$ where $\mathbf{x}$ is a random variable with distribution $p(\cdot; \mathbf{z})$. Similarly, we say we have access to a stochastic gradient oracle if we can sample from function $\hat{G} : \mathcal{K} \times \mathfrak{Y} \rightarrow \mathbb{R}$ such that $\nabla F(\mathbf{z}) = \mathbb{E}_{\mathbf{y} \sim q(\mathbf{y}; \mathbf{z})} [\hat{G}(\mathbf{z}, \mathbf{y})]$, and $\hat{G}$ is determined by $\mathbf{z}$ and the random variable $\mathbf{y}$ which is sampled according to $\mathbf{y} \sim q(\mathbf{y}; \mathbf{z})$. Note that oracles are only defined on the feasible set. We will use $\hat{G}(\mathbf{z})$ to denote the random variables $\hat{G}(\mathbf{z}, \mathbf{y})$ where $\mathbf{y}$ is a random variable with distribution $q(\cdot; \mathbf{z})$.

Assumption 1. We assume that $F : [0, 1]^d \rightarrow \mathbb{R}$ is DR-submodular, first-order differentiable, non-negative, G -Lipschitz for some $G < \infty$, and L -smooth for some $L < \infty$. We also assume the feasible region $\mathcal{K}$ is a closed convex domain in $[0, 1]^d$ with at least two points. Moreover, we also assume that we either have access to a value oracle with variance $\sigma_0^2 \geq 0$ or a gradient oracle with variance $\sigma_1^2 \geq 0$.

Remark 1. The proposed algorithm does not need to know the values of L , G , σ_0 or σ_1 . However, these constants appear in the final expressions of the number of oracle calls and the regret bounds.

3.2 Black Box Gradient Estimate

Without access to a gradient oracle (i.e., first-order information), we estimate gradient information using samples from a value oracle. We will use a variation of the “smoothing trick” technique [14, 17, 1, 27, 31, 8, 33], which involves averaging through spherical sampling around a given point.

Definition 1 (Smoothing Trick). For a function $F : \mathcal{D} \rightarrow \mathbb{R}$ defined on $\mathcal{D} \subseteq \mathbb{R}^d$, its δ -smoothed version $\tilde{F}_\delta$ is given as

$$\tilde{F}_\delta(\mathbf{x}) := \mathbb{E}_{\mathbf{z} \sim \mathbb{B}_\delta^{\text{aff}(\mathcal{D})}(\mathbf{x})} [F(\mathbf{z})] = \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{\text{aff}(\mathcal{D})-\mathbf{x}}(0)} [F(\mathbf{x} + \delta \mathbf{v})], \quad (6)$$

where $\mathbf{v}$ is chosen uniformly at random from the $\dim(\text{aff}(\mathcal{D}))$ -dimensional ball $\mathbb{B}_1^{\text{aff}(\mathcal{D})-\mathbf{x}}(0)$. Thus, the function value $\tilde{F}_\delta(\mathbf{x})$ is obtained by “averaging” F over a sliced ball of radius δ around $\mathbf{x}$.

When the value of δ is clear from the context, we may drop the subscript and simply use $\tilde{F}$ to denote the smoothed version of F . It can be easily seen that if F is DR-submodular, G -Lipschitz continuous, and L -smooth, then so is $\tilde{F}$ and $\|\tilde{F}(\mathbf{x}) - F(\mathbf{x})\| \leq \delta G$, for any point in the domain of both functions. Moreover, if F is monotone, then so is $\tilde{F}$ (Lemma 3). Therefore $\tilde{F}_\delta$ is an approximation of the function F . A maximizer of $\tilde{F}_\delta$ also maximizes F approximately.

Our definition of smoothing trick differs from the standard usage by accounting for the affine hull containing $\mathcal{D}$. This will be of particular importance when the feasible region is of (affine) dimension less than d , such as when there are equality constraints. When $\text{aff}(\mathcal{D}) = \mathbb{R}^d$, our definition reduces to the standard definition of the smoothing trick. In this case, it is well-known that the gradient of the smoothed function $\tilde{F}_\delta$ admits an unbiased one-point estimator [14, 17]. Using a two-point estimator instead of the one-point estimator results in smaller variance [1, 27]. In Algorithm 1, we adapt the two-point estimator to the general setting.

3.3 Construction of $\mathcal{K}_\delta$

We want to run Algorithm 1 as a subroutine within the main algorithm to estimate the gradient. However, in order to run Algorithm 1, we need to be able to query the oracle within the set $\mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\mathbf{x})$. Since the oracle can only be queried at points within the feasible set, we need to restrict our attention to a set $\mathcal{K}_\delta$ such that $\mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\mathcal{K}_\delta) \subseteq \mathcal{K}$. On the other hand, we want the optimal point of F within $\mathcal{K}_\delta$ to be close to the optimal point of F within $\mathcal{K}$. One way to ensure

that is to have $\mathcal{K}_\delta$ not be too small. More formally, we want that $\mathbb{B}_{\delta'}^{\text{aff}(\mathcal{K})}(\mathcal{K}_\delta) \supseteq \mathcal{K}$, for some value of $\delta' \geq \delta$ that is not too large. The constraint boundary could have a complex geometry, and simply maintaining a δ sized margin away from the boundary can result in big gaps between the boundary of $\mathcal{K}$ and $\mathcal{K}_\delta$. For example, in two dimensions, if $\mathcal{K}$ is polyhedral and has an acute angle, maintaining a δ margin away from both edges adjacent to the acute angle means the closest point in the $\mathcal{K}_\delta$ to the corner may be much more than δ . For this construction, we choose a $\mathbf{c} \in \text{relint}(\mathcal{K})$ and a real number $r > 0$ such that $\mathbb{B}_r^{\text{aff}(\mathcal{K})}(\mathbf{c}) \subseteq \mathcal{K}$. For any $\delta < r$, we define

$$\mathcal{K}_\delta^{\mathbf{c},r} := (1 - \frac{\delta}{r})\mathcal{K} + \frac{\delta}{r}\mathbf{c}. \quad (7)$$

Clearly if $\mathcal{K}$ is downward-closed, then so is $\mathcal{K}_\delta^{\mathbf{c},r}$. Lemma 7 shows that for any such choice of $\mathbf{c}$ and $r > 0$, we have $\frac{\delta'}{\delta} \leq \frac{D}{r}$. See Appendix G for more details about the choice of $\mathbf{c}$ and r . We will drop the superscripts in the rest of the paper when there is no ambiguity.

Remark 2. This construction is similar to the one carried out in [31] which was for d -dimensional downward-closed sets. Here we impose no restrictions on $\mathcal{K}$ beyond Assumption 1. A simpler construction of shrunken constraint set was proposed in [8]. However, as we discuss in Appendix D, they require to be able to query outside of the constraint set.

3.4 Generalized DR-Submodular Frank-Wolfe

The pseudocode of our proposed offline algorithm, Generalized DR-Submodular Frank-Wolfe, is shown in Algorithm 2. At a high-level, it follows the basic template of Frank-Wolfe type methods, where over the course of a pre-specified number of iterations, the gradient (or a surrogate thereof) is calculated, an optimization sub-routine with a linear objective is solved to find a feasible point whose difference (with respect to the current solution) has the largest inner product with respect to the gradient, and then the current solution is updated to move in the direction of that feasible point.

However, there are a number of important modifications to handle properties of the objective function, constraint set, and oracle type. For the oracle type, for instance, standard Frank-Wolfe methods assume access to a deterministic gradient oracle. Frank-Wolfe methods are known to be sensitive to errors in estimates of the gradient (e.g., see [16]). Thus, when only a stochastic gradient oracle or even more challenging, only a stochastic value oracle is available, the gradient estimators must be carefully designed to balance query complexity on the one hand and output error on the other. The Black Box Gradient Estimate (BBGE) sub-routine, presented in Algorithm 1, utilizes spherical sampling to produce an unbiased gradient estimate. This estimate is then combined with past estimates using momentum, as seen in [23], to control and reduce variance.

Algorithm 1 Black Box Gradient Estimate (BBGE)

- 1: **Input:** Point $\mathbf{z}$, sampling radius δ , constraint linear space $\mathcal{L}_0$, $k = \dim(\mathcal{L}_0)$, batch size B
 - 2: Sample $\mathbf{u}_1, \dots, \mathbf{u}_B$ i.i.d. from $S^{d-1} \cap \mathcal{L}_0$
 - 3: For $i = 1$ to B , let $\mathbf{y}_i^+ \leftarrow \mathbf{z} + \delta \mathbf{u}_i$, $\mathbf{y}_i^- \leftarrow \mathbf{z} - \delta \mathbf{u}_i$, and evaluate $\hat{F}(\mathbf{y}_i^+)$, $\hat{F}(\mathbf{y}_i^-)$
 - 4: $\mathbf{g} \leftarrow \frac{1}{B} \sum_{i=1}^B \frac{k}{2\delta} [\hat{F}(\mathbf{y}_i^+) - \hat{F}(\mathbf{y}_i^-)] \mathbf{u}_i$
 - 5: Output $\mathbf{g}$
-

Algorithm 2 Generalized DR-Submodular Frank-Wolfe

- 1: **Input:** Constraint set $\mathcal{K}$, iteration limit $N \geq 4$, sampling radius δ , gradient step-size $\{\rho_n\}_{n=1}^N$
 - 2: Construct $\mathcal{K}_\delta$
 - 3: Pick any $\mathbf{z}_1 \in \arg\min_{\mathbf{z} \in \mathcal{K}_\delta} \|\mathbf{z}\|_\infty$
 - 4: $\bar{\mathbf{g}}_0 \leftarrow \mathbf{0}$
 - 5: **for** $n = 1$ **to** N **do**
 - 6: $\mathbf{g}_n \leftarrow \text{estimate-grad}(\mathbf{z}_n, \delta, \mathcal{L}_0 = \text{aff}(\mathcal{K}) - \mathbf{z}_1)$
 - 7: $\bar{\mathbf{g}}_n \leftarrow (1 - \rho_n)\bar{\mathbf{g}}_{n-1} + \rho_n \mathbf{g}_n$
 - 8: $\mathbf{v}_n \leftarrow \text{optimal-direction}(\bar{\mathbf{g}}_n, \mathbf{z}_n)$
 - 9: $\mathbf{z}_{n+1} \leftarrow \text{update}(\mathbf{z}_n, \mathbf{v}_n, \varepsilon)$
 - 10: **end for**
 - 11: Output $\mathbf{z}_{N+1}$
-

Our algorithm design is influenced by state-of-the-art methods that have been developed for specific settings. One of the most closely related works is [8], which also dealt with using value oracle access for optimizing monotone functions. They used momentum and spherical sampling techniques that are similar to the ones we used in our Algorithm 1. However, we modified the sampling procedure and the solution update step. In their work, [8] also considered a shrunken feasible region to avoid sampling close to the boundary. However, they assumed that the value oracle could be queried outside the feasible set (see Appendix D for details).

In Algorithm 3, we consider the following cases for the function and the feasible set.

- (A) If F is monotone DR-submodular and $\mathbf{0} \in \mathcal{K}$, we choose

$$\text{optimal-direction}(\bar{\mathbf{g}}_n, \mathbf{z}_n) = \operatorname{argmax}_{\mathbf{v} \in \mathcal{K}_\delta - \mathbf{z}_1} \langle \mathbf{v}, \bar{\mathbf{g}}_n \rangle, \text{ update}(\mathbf{z}_n, \mathbf{v}_n, \varepsilon) = \mathbf{z}_n + \varepsilon \mathbf{v}_n,$$

and $\varepsilon = 1/N$. We start at a point near the origin and always move to points that are bigger with respect to the partial order on $\mathbb{R}^d$. In this case, since the function is monotone, the optimal direction is a maximal point with respect to the partial order. The choice of $\varepsilon = 1/N$ guarantees that after N steps, we arrive at a convex combination of points in the feasible set and therefore the final point is also in the feasible set. The fact that the origin is also in the feasible set shows that the intermediate points also belong to the feasible set.

- (B) If F is non-monotone DR-submodular and $\mathcal{K}$ is a downward closed set containing $\mathbf{0}$, we choose

$$\text{optimal-direction}(\bar{\mathbf{g}}_n, \mathbf{z}_n) = \operatorname{argmax}_{\substack{\mathbf{v} \in \mathcal{K}_\delta - \mathbf{z}_1 \\ \mathbf{v} \leq \mathbf{1} - \mathbf{z}_n}} \langle \mathbf{v}, \bar{\mathbf{g}}_n \rangle, \text{ update}(\mathbf{z}_n, \mathbf{v}_n, \varepsilon) = \mathbf{z}_n + \varepsilon \mathbf{v}_n,$$

and $\varepsilon = 1/N$. This case is similar to (A). However, since F is not monotone, we need to choose the optimal direction more conservatively.

- (C) If F is monotone DR-submodular and $\mathcal{K}$ is a general convex set, we choose

$$\text{optimal-direction}(\bar{\mathbf{g}}_n, \mathbf{z}_n) = \operatorname{argmax}_{\mathbf{v} \in \mathcal{K}_\delta} \langle \mathbf{v}, \bar{\mathbf{g}}_n \rangle, \text{ update}(\mathbf{z}_n, \mathbf{v}_n, \varepsilon) = (1 - \varepsilon)\mathbf{z}_n + \varepsilon \mathbf{v}_n,$$

and $\varepsilon = \log(N)/2N$. In this case, if we update like in cases (A) and (B), we do not have any guarantees of ending up in the feasible set, so we choose the update function to be a convex combination. Unlike (B), we do not need to limit ourselves in choosing the optimal direction and we simply choose ε to obtain the best approximation coefficient.

- (D) If F is non-monotone DR-submodular and $\mathcal{K}$ is a general convex set, we choose

$$\text{optimal-direction}(\bar{\mathbf{g}}_n, \mathbf{z}_n) = \operatorname{argmax}_{\mathbf{v} \in \mathcal{K}_\delta} \langle \mathbf{v}, \bar{\mathbf{g}}_n \rangle, \text{ update}(\mathbf{z}_n, \mathbf{v}_n, \varepsilon) = (1 - \varepsilon)\mathbf{z}_n + \varepsilon \mathbf{v}_n,$$

and $\varepsilon = \log(2)/N$. This case is similar to (C) and we choose ε to obtain the best approximation coefficient.

The choice of subroutine estimate-grad and ρ_n depend on the oracle. If we have access to a gradient oracle $\hat{G}$, we set estimate-grad($\mathbf{z}, \delta, \mathcal{L}_0$) to be the average of B evaluations of $P_{\mathcal{L}_0}(\hat{G}(\mathbf{z}))$. Otherwise, we run Algorithm 1 with input $\mathbf{z}, \delta, \mathcal{L}_0$. If we have access to a deterministic gradient oracle, then there is no need to use any momentum and we set $\rho_n = 1$. In other cases, we choose $\rho_n = \frac{2}{(n+3)^{2/3}}$.

3.5 Approximation Guarantees for the Proposed Offline Algorithm

Theorem 1. Suppose Assumption 1 holds. Let $N \geq 4$, $B \geq 1$ and choose $\mathbf{c} \in \mathcal{K}$ and $r > 0$ according to Section 3.3. If we have access to a gradient oracle, we choose $\delta = 0$, otherwise we choose $\delta \in (0, r/2)$. Then the following results hold for the output $\mathbf{z}_{N+1}$ of Algorithm 2.

- (A) If F is monotone DR-submodular and $\mathbf{0} \in \mathcal{K}$, then

$$(1 - e^{-1})F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \delta G(2 + \frac{\sqrt{d} + D}{r}). \quad (8)$$

- (B) If F is non-monotone DR-submodular and $\mathcal{K}$ is a downward closed set containing $\mathbf{0}$, then

$$e^{-1}F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \delta G(2 + \frac{\sqrt{d} + 2D}{r}). \quad (9)$$

(C) If F is monotone DR-submodular and $\mathcal{K}$ is a general convex set, then

$$\frac{1}{2}F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}\log(N)}{2N^{1/3}} + \frac{4DG + LD^2\log(N)^2}{8N} + \delta G(2 + \frac{D}{r}). \quad (10)$$

(D) If F is non-monotone DR-submodular and $\mathcal{K}$ is a general convex set, then

$$\frac{1}{4}(1 - \|\mathbf{z}_1\|_\infty)F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{DG + 2LD^2}{4N} + \delta G(2 + \frac{D}{r}). \quad (11)$$

In all these cases, we have

$$Q = \begin{cases} 0 & \text{det. grad. oracle,} \\ \max\{4^{2/3}G^2, 6L^2D^2 + \frac{4\sigma_1^2}{B}\} & \text{stoch. grad. oracle with variance } \sigma_1^2 > 0, \\ \max\{4^{2/3}G^2, 6L^2D^2 + \frac{4CkG^2 + 2k^2\sigma_0^2/\delta^2}{B}\} & \text{value oracle with variance } \sigma_0^2 \geq 0, \end{cases}$$

C is a constant, $k = \dim(\mathcal{K})$, $D = \text{diam}(\mathcal{K})$, and $\mathbf{z}^*$ is the global maximizer of F on $\mathcal{K}$.

Theorem 1 characterizes the worst-case approximation ratio α and additive error bounds for different properties of the function and feasible region, where the additive error bounds depend on selected parameters N for the number of iterations, batch size B , and sampling radius δ .

The proof of Parts (A)-(D) is provided in Appendix I-L, respectively.

The proof of parts (A), (B) and (D), when we have access to an exact gradient oracle is similar to the proofs presented in [4, 3, 24], respectively. Part (C) is the first analysis of a Frank-Wolfe type algorithm over general convex sets when the oracle can only be queried within the feasible set. When we have access to a stochastic gradient oracle, directly using a gradient sample can result in arbitrary bad performance as shown in Appendix B of [16]. The momentum technique, first used in continuous submodular maximization in [23], is used when we have access to a stochastic gradient oracle. The control on the estimate of the gradient is deferred to Lemma 9. Since the momentum technique is robust to noise in the gradient, when we only have access to a value oracle, we can use Algorithm 1, similar to [8], to obtain an unbiased estimate of the gradient and complete the proof.

Theorem 2 converts those bounds to characterize the oracle complexity for a user-specified additive error tolerance ϵ based on oracle properties (deterministic/stochastic gradient/value). The 16 combinations of the problem settings listed in Table 1 are enumerated by four cases (A)–(D) in Theorem 1 of function and feasible region properties (resulting in different approximation ratios) and the four cases 1–4 enumerated in Theorem 2 below of oracle properties. For the oracle properties, we consider the four cases as (Case 1): deterministic gradient oracle, (Case 2): stochastic gradient oracle, (Case 3): deterministic value oracle, and (Case 4): stochastic value oracle.

Theorem 2. *The number of oracle calls for different oracles to achieve an α -approximation error of smaller than ϵ using Algorithm 1 is*

$$\text{Case 1: } \tilde{O}(1/\epsilon), \quad \text{Cases 2, 3: } \tilde{O}(1/\epsilon^3), \quad \text{Case 4: } \tilde{O}(1/\epsilon^5).$$

Moreover, in all of the cases above, if F is non-monotone or $0 \in \mathcal{K}$, we may replace $\tilde{O}$ with O .

See Appendix M for proof.

4 Online DR-submodular optimization under bandit or semi-bandit feedback

In this section, we first describe the Black-box Explore-Then-Commit algorithm that uses the offline algorithm for exploration, and uses the solution of the offline algorithm for exploitation. This is followed by the regret analysis of the proposed algorithm. This is the first algorithm for stochastic continuous DR-submodular maximization under bandit feedback and obtains state-of-the-art for semi-bandit feedback.

4.1 Problem Setup

There are typically two settings considered in online optimization with bandit feedback. The first is the adversarial setting, where the environment chooses a sequence of functions $F_1, \dots, F_N$ and in each iteration n , the agent chooses a point $\mathbf{z}_n$ in the feasible set $\mathcal{K}$, observes $F_n(\mathbf{z}_n)$ and receives the reward $F_n(\mathbf{z}_n)$. The goal is to choose the sequence of actions that minimize the following notion of expected α -regret.

$$\mathcal{R}_{\text{adv}} := \alpha \max_{\mathbf{z} \in \mathcal{K}} \sum_{n=1}^N F_n(\mathbf{z}) - \mathbb{E} \left[\sum_{n=1}^N F_n(\mathbf{z}_n) \right]. \quad (12)$$

In other words, the agent's cumulative reward is being compared to α times the reward of the best *constant* action in hindsight. Note that, in this case, the randomness is over the actions of the policy.

The second is the stochastic setting, where the environment chooses a function $F : \mathcal{K} \rightarrow \mathbb{R}$ and a stochastic value oracle $\hat{F}$. In each iteration n , the agent chooses a point $\mathbf{z}_n$ in the feasible set $\mathcal{K}$, receives the reward $(\hat{F}(\mathbf{z}_n))_n$ by querying the oracle at $\mathbf{z}_n$ and observes this reward. Here the outer subscript n indicates that the result of querying the oracle at time n , since the oracle is stochastic. The goal is to choose the sequence of actions that minimize the following notion of expected α -regret.

$$\mathcal{R}_{\text{stoch}} := \alpha N \max_{\mathbf{z} \in \mathcal{K}} F(\mathbf{z}) - \mathbb{E} \left[\sum_{n=1}^N (\hat{F}(\mathbf{z}_n))_n \right] = \alpha N \max_{\mathbf{z} \in \mathcal{K}} F(\mathbf{z}) - \mathbb{E} \left[\sum_{n=1}^N F(\mathbf{z}_n) \right] \quad (13)$$

Further, two feedback models are considered – bandit and semi-bandit feedback. In the bandit feedback setting, the agent only observes the value of the function F_n at the point $\mathbf{z}_n$. In the semi-bandit setting, the agent has access to a gradient oracle instead of a value oracle and observes $\hat{G}(\mathbf{z}_n)$ at the point $\mathbf{z}_n$, where $\hat{G}$ is an unbiased estimator of ∇F .

In unstructured multi-armed bandit problems, any regret bound for the adversarial setup could be translated into bounds for the stochastic setup. However, having a non-trivial correlation between the actions of different arms complicates the relation between the stochastic and adversarial settings. Even in linear bandits, the relation between adversarial linear bandits and stochastic linear bandits is not trivial. (e.g. see Section 29 in [19]) While it is intuitively reasonable to assume that the optimal regret bounds for the stochastic case are better than that of the adversarial case, such a result is not yet proven for DR-submodular functions. Thus, while the cases of bandit feedback has been studied in the adversarial setup, the results do not reduce to stochastic setup. We also note that in the cases where there are adversarial setup results, this paper finds that the results in the stochastic setup achieve improved regret bounds (See Table 3 in Supplementary for the comparison).

4.2 Algorithm for DR-submodular maximization with Bandit Feedback

The proposed algorithm is described in Algorithm 3. In Algorithm 3, if there is semi-bandit feedback in the form of a stochastic gradient sample for each action $\mathbf{z}_n$, we run the offline algorithm (Algorithm 2) with parameters from the proof of case 2 of Theorem 2 for $T_0 = \lceil T^{3/4} \rceil$ total queries. If only the stochastic reward for each action $\mathbf{z}_n$ is available (bandit feedback), we run the offline algorithm (Algorithm 2) with parameters from the proof of case 4 of Theorem 2 for $T_0 = \lceil T^{5/6} \rceil$ total queries. Then, for the remaining time (exploitation phase), we run the last action in the exploration phase.

Algorithm 3 DR-Submodular Explore-Then-Commit

- 1: **Input:** Horizon T , inner time horizon T_0
 - 2: Run Algorithm 2 for T_0 , with according to parameters described in Theorem 2.
 - 3: **for** remaining time **do**
 - 4: Repeat the last action of Algorithm 2.
 - 5: **end for**
-

4.3 Regret Analysis for DR-submodular maximization with Bandit Feedback

In this section, we provide the regret analysis for the proposed algorithm. We note that by Theorem 2, Algorithm 2 requires a sample complexity of $\tilde{O}(1/\epsilon^5)$ with a stochastic value oracle for

offline problems (any of (A)–(D) in Theorem 1). Thus, the parameters and the results with bandit feedback are the same for all the four setups (A)–(D). Likewise, when a stochastic gradient oracle is available, Algorithm 2 requires a sample complexity of $\tilde{O}(1/\epsilon^3)$. Based on these sample complexities, the overall regret of online DR-submodular maximization problem is given as follows.

Theorem 3. *For an online constrained DR-submodular maximization problem over a horizon T , where the expected reward function F , feasible region type $\mathcal{K}$, and approximation ratio α correspond to any of the four cases (A)–(D) in Theorem 1, Algorithm 3 achieves α -regret (13) that is upper-bounded as:*

$$\text{Semi-bandit Feedback (Case 2): } \tilde{O}(T^{3/4}), \quad \text{Bandit Feedback (Case 4): } \tilde{O}(T^{5/6}).$$

Moreover, in either type of feedback, if F is non-monotone or $\mathbf{0} \in \mathcal{K}$, we may replace $\tilde{O}$ with O .

See Appendix N for the proof.

5 Conclusion

This work provides a novel and unified approach for maximizing continuous DR-submodular functions across various assumptions on function, constraint set, and oracle access types. The proposed Frank-Wolfe based algorithm improves upon existing results for nine out of the sixteen cases considered, and presents new results for offline DR-submodular maximization with only a value oracle. Moreover, this work presents the first regret analysis with bandit feedback for stochastic DR-submodular maximization, covering both monotone and non-monotone functions. These contributions significantly advance the field of DR-submodular optimization (with multiple applications) and open up new avenues for future research in this area.

Limitations: While the number of function evaluations in the different setups considered in the paper are state of the art, lower bounds have not been investigated.

References

- [1] Alekh Agarwal, Ofer Dekel, and Lin Xiao. Optimal algorithms for online convex optimization with multi-point bandit feedback. In *Proceedings of the 23rd Annual Conference on Learning Theory*, 2010.
- [2] Francis Bach. Submodular functions: from discrete to continuous domains. *Mathematical Programming*, 175:419–459, 2019.
- [3] An Bian, Kfir Levy, Andreas Krause, and Joachim M Buhmann. Continuous DR-submodular maximization: Structure and algorithms. In *Advances in Neural Information Processing Systems*, 2017.
- [4] Andrew An Bian, Baharan Mirzasoleiman, Joachim Buhmann, and Andreas Krause. Guaranteed Non-convex Optimization: Submodular Maximization over Continuous Domains. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, April 2017.
- [5] Yatao Bian, Joachim Buhmann, and Andreas Krause. Optimal continuous DR-submodular maximization and applications to provable mean field inference. In *Proceedings of the 36th International Conference on Machine Learning*, June 2019.
- [6] Lin Chen, Christopher Harshaw, Hamed Hassani, and Amin Karbasi. Projection-free online optimization with stochastic gradient: From convexity to submodularity. In *Proceedings of the 35th International Conference on Machine Learning*, July 2018.
- [7] Lin Chen, Hamed Hassani, and Amin Karbasi. Online continuous submodular maximization. In *Proceedings of the Twenty-First International Conference on Artificial Intelligence and Statistics*, April 2018.
- [8] Lin Chen, Mingrui Zhang, Hamed Hassani, and Amin Karbasi. Black box submodular maximization: Discrete and continuous settings. In *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, August 2020.

- [9] Josip Djolonga and Andreas Krause. From map to marginals: Variational inference in Bayesian submodular models. *Advances in Neural Information Processing Systems*, 2014.
- [10] Donglei Du. Lyapunov function approach for approximation algorithm design and analysis: with applications in submodular maximization. *arXiv preprint arXiv:2205.12442*, 2022.
- [11] Donglei Du, Zhicheng Liu, Chencheng Wu, Dachuan Xu, and Yang Zhou. An improved approximation algorithm for maximizing a DR-submodular function over a convex set. *arXiv preprint arXiv:2203.14740*, 2022.
- [12] Christoph Dürr, Nguyen Kim Thang, Abhinav Srivastav, and Léo Tible. Non-monotone DR-submodular maximization: Approximation and regret guarantees. *arXiv preprint arXiv:1905.09595*, 2019.
- [13] Maryam Fazel and Omid Sadeghi. *Fast First-Order Methods for Monotone Strongly DR-Submodular Maximization*, pages 169–179. SIAM Conference on Applied and Computational Discrete Algorithms, 2023.
- [14] Abraham D Flaxman, Adam Tauman Kalai, and H Brendan McMahan. Online convex optimization in the bandit setting: gradient descent without a gradient. In *Proceedings of the sixteenth annual ACM-SIAM symposium on Discrete algorithms*, pages 385–394, 2005.
- [15] Shuyang Gu, Chuangen Gao, Jun Huang, and Weili Wu. Profit maximization in social networks and non-monotone DR-submodular maximization. *Theoretical Computer Science*, 957:113847, 2023.
- [16] Hamed Hassani, Mahdi Soltanolkotabi, and Amin Karbasi. Gradient methods for submodular maximization. In *Advances in Neural Information Processing Systems*, 2017.
- [17] Elad Hazan et al. Introduction to online convex optimization. *Foundations and Trends® in Optimization*, 2(3-4):157–325, 2016.
- [18] Shinji Ito and Ryohei Fujimaki. Large-scale price optimization via network flow. *Advances in Neural Information Processing Systems*, 2016.
- [19] Tor Lattimore and Csaba Szepesvári. *Bandit algorithms*. Cambridge University Press, 2020.
- [20] Yuanyuan Li, Yuezhou Liu, Lili Su, Edmund Yeh, and Stratis Ioannidis. Experimental design networks: A paradigm for serving heterogeneous learners under networking constraints. *IEEE/ACM Transactions on Networking*, 2023.
- [21] Yuanyuan Li, Yuchao Zhang, Stratis Ioannidis, and Jon Crowcroft. Jointly optimal routing and caching with bounded link capacities. In *ICC 2023 - IEEE International Conference on Communications*, pages 1130–1136, 2023.
- [22] Aryan Mokhtari, Hamed Hassani, and Amin Karbasi. Conditional gradient method for stochastic submodular maximization: Closing the gap. In *International Conference on Artificial Intelligence and Statistics*, pages 1886–1895. PMLR, 2018.
- [23] Aryan Mokhtari, Hamed Hassani, and Amin Karbasi. Stochastic conditional gradient methods: From convex minimization to submodular maximization. *The Journal of Machine Learning Research*, 21(1):4232–4280, 2020.
- [24] Loay Mualem and Moran Feldman. Resolving the approximability of offline and online non-monotone DR-submodular maximization over general convex sets. In *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, April 2023.
- [25] Rad Niazadeh, Negin Golrezaei, Joshua Wang, Fransisca Susan, and Ashwinkumar Badanidiyuru. Online learning via offline greedy algorithms: Applications in market design and optimization. *Management Science*, 69(7):3797–3817, July 2023.
- [26] Rad Niazadeh, Tim Roughgarden, and Joshua R Wang. Optimal algorithms for continuous non-monotone submodular and DR-submodular maximization. *The Journal of Machine Learning Research*, 21(1):4937–4967, 2020.

- [27] Ohad Shamir. An optimal algorithm for bandit and zero-order convex optimization with two-point feedback. *Journal of Machine Learning Research*, 18(52):1–11, 2017.
- [28] Nguyen Kim Thang and Abhinav Srivastav. Online non-monotone DR-submodular maximization. *Proceedings of the AAAI Conference on Artificial Intelligence*, May 2021.
- [29] Jan Vondrák. Symmetry and approximability of submodular maximization problems. *SIAM Journal on Computing*, 42(1):265–304, 2013.
- [30] Zongqi Wan, Jialin Zhang, Wei Chen, Xiaoming Sun, and Zhijie Zhang. Bandit multi-linear dr-submodular maximization and its applications on adversarial submodular bandits. In *International Conference on Machine Learning*, 2023.
- [31] Mingrui Zhang, Lin Chen, Hamed Hassani, and Amin Karbasi. Online continuous submodular maximization: From full-information to bandit feedback. In *Advances in Neural Information Processing Systems*, volume 32, 2019.
- [32] Qixin Zhang, Zengde Deng, Zaiyi Chen, Haoyuan Hu, and Yu Yang. Stochastic continuous submodular maximization: Boosting via non-oblivious function. In *Proceedings of the 39th International Conference on Machine Learning*, 2022.
- [33] Qixin Zhang, Zengde Deng, Zaiyi Chen, Kuangqi Zhou, Haoyuan Hu, and Yu Yang. Online learning for non-monotone DR-submodular maximization: From full information to bandit feedback. In *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, April 2023.

A Details of Related Works

A.1 Offline DR-submodular maximization

The authors of [4] considered the problem of maximizing a monotone DR-submodular function over a downward-closed convex set given a deterministic gradient oracle. They showed that a variant of the Frank-Wolfe algorithm guarantees an optimal $(1 - \frac{1}{e})$ -approximation for this problem. While they only claimed their result for downward-closed convex sets, their result holds under a more general setting where the convex set contains the origin. In [3], a non-monotone variant of the algorithm for downward-closed convex sets with $\frac{1}{e}$ -approximation was proposed.

The authors of [16] used gradient ascent to obtain $\frac{1}{2}$ -guarantees for the maximization of a monotone DR-submodular function over a general convex set given a gradient oracle which could be stochastic. They proved that gradient ascent cannot guarantee better than a $\frac{1}{2}$ -approximation by constructing a convex set $\mathcal{K}$ and a function $F : \mathcal{K} \rightarrow \mathbb{R}$ such that F has a local maximum that is a $\frac{1}{2}$ -approximation of its optimal value on $\mathcal{K}$. They also showed that a Frank-Wolfe type algorithm similar to [4] cannot be directly used when we only have access to a stochastic gradient oracle. [32] extended projected gradient ascent using a line integral method, referred to as boosting, to obtain $(1 - 1/e)$ -approximation for convex sets containing the origin. Later, [23] resolved the issue of stochastic gradient oracles with a momentum technique and obtained $(1 - \frac{1}{e})$ -approximation in the case of monotone functions over sets that contain the origin, and $\frac{1}{e}$ -approximation in the case of non-monotone functions over downward closed sets. In [32] and the first case in [23], while they consider monotone DR-submodular functions over general convex sets $\mathcal{K}$, they query the oracle over the convex hull of $\mathcal{K} \cup \{0\}$ (See Appendix B).

For non-monotone maps over general convex sets, no constant approximation ratio can be guaranteed in sub-exponential time due to a hardness result by [29]. However, [12] bypassed this issue by finding an approximation guarantee that depends on the geometry of the convex set. Specifically, they showed that given a deterministic gradient oracle for a non-monotone function over a general convex set $\mathcal{K} \subseteq [0, 1]^d$, their proposed algorithm obtains $\frac{1}{3\sqrt{3}}(1 - h)$ -approximation of the optimal value where $h := \min_{\mathbf{z} \in \mathcal{K}} \|\mathbf{z}\|_\infty$. An improved sub-exponential algorithm was proposed by [11] that obtained a $\frac{1}{4}(1 - h)$ -approximation guarantees, which is optimal. Later, [10] provided the first polynomial time algorithm for this setting with the same approximation coefficient.

Remark 3. In the special case of maximizing a non-monotone continuous DR-submodular over a box, i.e. $[0, 1]^d$, one could discretize the problem and use discrete algorithms to solve the continuous version. The technique has been employed in [3] to obtain a $\frac{1}{3}$ -approximation and in [5, 26] to obtain $\frac{1}{2}$ -approximations for the optimal value. We have not included these results in Table 1 since using discretization has only been successfully applied to the case where the convex set is a box and can not be directly used in more general settings.

A.2 Online DR-submodular maximization with bandit feedback

There has been growing interest in online DR-submodular maximization in the recent years [7], [6], [31], [28], [25], [33], [13], [24]. Most of these results are focused on adversarial online full-information feedback. In the adversarial setting, the environment chooses a sequence of functions $F_1, \dots, F_N$ and in each iteration n , the agent chooses a point z_n in the feasible set $\mathcal{K}$, observes F_n and receives the reward $F_n(z_n)$. For the regret bound, the agents reward is being compared to α times the reward of the best *constant* action in hindsight. With full-information feedback, if at each iteration when the agent observes F_n , it may be allowed to query the value of ∇F_n or maybe F_n at any number of arbitrary points within the feasible set. Further, we consider stochastic setting, where the environment chooses a function $F : \mathcal{K} \rightarrow \mathbb{R}$ and a sequence of independent noise functions $\eta_n : \mathcal{K} \rightarrow \mathbb{R}$ with zero mean. In each iteration n , the agent chooses a point $\mathbf{z}_n$ in the feasible set $\mathcal{K}$, receives the reward $(F + \eta_n)(\mathbf{z}_n)$ and observes the reward. For the regret bound, the agents reward is being compared to α times the reward of the best action. Detailed formulation of adversarial and stochastic setups and why adversarial results cannot be reduced to stochastic results is given in Section 4.1. In this paper, we consider two feedback models – bandit feedback where only the (stochastic) reward value is available and semi-bandit feedback where a single stochastic sample of the gradient at the location is provided.

Function	Set	Setting	Reference	Appx.	Regret
Monotone	$0 \in \mathcal{K}$	stoch.	This paper	$1 - 1/e$	$O(T^{5/6})$
		adv.	[31]	$1 - 1/e$	$O(T^{8/9})$
			[25]	$1 - 1/e$	$O(T^{5/6})$
			[30] [†]	$1 - 1/e$	$O(T^{3/4})$
	general	stoch.	This paper	$1/2$	$\tilde{O}(T^{5/6})$
		adv.	-		
Non-monotone	d.c.	stoch.	This paper	$1/e$	$O(T^{5/6})$
		adv.	[33]	$1/e$	$O(T^{8/9})$
	general	stoch.	This paper	$\frac{1-h}{4}$	$O(T^{5/6})$
		adv.	-		

Table 3: This table presents the different results for the regret for DR-submodular maximization under bandit feedback, and gives the related works and regret bounds in the adversarial case. Note that the result marked by [†] uses a convex optimization subroutine at each iteration which could be even more computationally expensive than projection. As before, we have $h := \min_{\mathbf{x} \in \mathcal{K}} \|\mathbf{x}\|_\infty$.

Bandit Feedback: We note that this paper is the first work for bandit feedback for stochastic online DR-submodular maximization. The prior works on this topic has been in the adversarial setup [31, 33, 25, 30], and the results in this work is compared with their results in Table 3. In [31], the adversarial online setting with bandit feedback has been studied for monotone DR-submodular functions over downward-closed convex sets. Later [33] extended this framework to the setting with non-monotone DR-submodular functions over downward-closed convex sets. [25] described a framework for converting certain greedy-type offline algorithms with robustness guarantees into adversarial online algorithms for both full-information and bandit feedback. They apply their framework to obtain algorithms for non-monotone functions over a box, with $\frac{1}{2}$ -regret of $\tilde{O}(T^{4/5})$, and monotone function over downward-closed convex sets. The offline algorithm they use for downward-closed convex sets is the one described in [4] which only requires the convex set to contain the origin. They also use the construction of the shrunk constraint set described in [31]. By replacing that construction with ours, the result of [25] could be extended to monotone functions over all convex sets containing the origin. [30] improved the regret bound for monotone functions over convex sets containing the origin to $O(T^{3/4})$. However, they use a convex optimization subroutine at each iteration which could be even more computationally expensive than projection.

Semi-bandit Feedback: In semi-bandit feedback, a single stochastic sample of the gradient is available. The problem has been considered in [6], while the results have an error (See Appendix D). Further, they only obtain $\frac{1}{e}$ -regret for the monotone case. One could consider a generalization of the adversarial and stochastic setting in the following manner. The environment chooses a sequence of functions F_n and a sequence of value oracles $\hat{F}_n$ such that $\hat{F}_n$ estimates F_n . In each iteration n , the agent chooses a point $\mathbf{z}_n$ in the feasible set $\mathcal{K}$, receives the reward $(\hat{F}_n(\mathbf{z}_n))_n$ by querying the oracle at $\mathbf{z}_n$ and observes this reward. The goal is to choose the sequence of actions that minimize the following notion of expected α -regret.

$$\begin{aligned}
\mathcal{R}_{\text{stoch-adv}} &:= \alpha \max_{\mathbf{z} \in \mathcal{K}} \sum_{n=1}^N F_n(\mathbf{z}) - \mathbb{E} \left[\sum_{n=1}^N (\hat{F}_n(\mathbf{z}_n))_n \right] \\
&= \alpha \max_{\mathbf{z} \in \mathcal{K}} \sum_{n=1}^N F_n(\mathbf{z}) - \mathbb{E} \left[\sum_{n=1}^N F_n(\mathbf{z}_n) \right]
\end{aligned} \tag{14}$$

Algorithm 3 in [7] solves this problem in semi-bandit feedback setting with a deterministic value oracle and stochastic gradient oracles. Any bound for a problem in this setting implies bounds for stochastic semi-bandit and adversarial semi-bandit settings. The same is true for Mono-Frank-Wolfe Algorithms in [31, 33]. We have included these results in Table 2 as benchmark to compare with results in stochastic setting.

B Constraint Set and Query Set

In this work, we made the assumption that the query set is identical to the constraint set, i.e. oracles can only be queried within the constraint set. To the best of our knowledge, except in the context of online optimization with (semi-)bandit feedback, this is the first work on DR-submodular maximization that explicitly considers this assumption. Previous works assumed that we may query the oracle at any point within the unit box $[0, 1]^d$. Algorithms designed for non-monotone functions in prior works already satisfied the assumption we consider, so no changes in algorithms, proofs, or results are needed. However, the situation is different when the function is monotone. This assumption allows us to explain a previously unexplained gap in approximation guarantees for monotone DR-submodular maximization. Specifically, some prior works (enumerated below) studying monotone DR-submodular maximization over general convex sets obtained approximation guarantees of $1/2$ while others obtained $1 - 1/e$.

First we describe how some of previous results in literature with no apparent restriction on the query set may be reformulated as problems where the query set is equal to the constraint set. Let $\mathcal{K} \subseteq [0, 1]^d$ be a convex set, and define $\mathcal{K}^*$ as the convex hull of $\mathcal{K} \cup \{0\}$. For a problem in the setting of monotone functions over a general set $\mathcal{K}$, we can consider the same problem on $\mathcal{K}^*$. Since the function is monotone, the optimal solution in $\mathcal{K}^*$ is the same as the optimal solution in $\mathcal{K}$. However, solving this problem in $\mathcal{K}^*$ may require evaluating the function in the larger set $\mathcal{K}^*$, which may not always be possible. In fact, the result of [23] and [32] mentioned in Table 1 are for monotone functions over general convex sets $\mathcal{K}$, but their algorithms require evaluating the function on $\mathcal{K}^*$. This is why we have classified their results as algorithms for convex sets that contain the origin. The problem of offline DR-submodular maximization with only a value oracle was first considered by [8] for monotone maps over convex sets that contain the origin. However, their result requires querying in a neighborhood of $\mathcal{K}^*$ which violates our requirement to only query the oracle within the feasible set (see Appendix D).

In [16], a $1/2$ approximation guarantee was obtained by a projected gradient ascent method and this was shown by proving that the algorithm tends to a stationary point and proving that any stationary point is at least $1/2$ as good as the optimal point. Moreover, they construct examples with stationary points that are no better than $1/2$ of the optimal point.

The $1 - 1/e$ approximation guarantee was first reported for Frank-Wolfe methods, which (superficially) suggests that the gap may be due to algorithm or analysis differences. Later, [32] developed a projected gradient ascent based method that obtains a $1 - 1/e$ approximation guarantee where they consider general constraint set but their query set contains the origin.

However, the gap is not attributable to algorithm or analysis differences, but instead due to the fact that the query sets are different. In other words, the results that obtain a $1 - 1/e$ approximation guarantee are solving a different problem than the ones obtaining a $1/2$ approximation guarantee. A key ingredient to obtain $1 - 1/e$ is the ability to query the (gradient) oracle within the convex hull of $\mathcal{K} \cup \{0\}$. For monotone submodular maximization over general convex sets (not necessarily containing the origin), we can only guarantee a coefficient of $1/2$, both for Frank-Wolfe type methods (our work) and projection based methods (i.e. [16]). Therefore, the $1/2$ approximation could very well be optimal in its own setting.

To the best of our knowledge, in every paper where the $1/2$ approximation coefficient and $1 - 1/e$ approximation coefficient in the monotone setting are compared, the comparison was (unwittingly) between problems that are inherently mathematically different: [16] and [7] in experiments and main text; [6] and [8] in experiments; [33, 24], and [12] in related work section, [23] in the introduction and Table 2, [32] and [13] in the main claims.

Conjecture *The problem of maximizing a monotone DR-submodular continuous function subject to a general convex constraint, where oracle queries are limited to the feasible region, is NP-hard. For any $\epsilon > 0$, it cannot be approximated in polynomial time to within a ratio of $1/2 + \epsilon$ (up to low-order terms), unless $RP = NP$.*

C Brief discussion on oracle models in applications

For many problems, the ability to evaluate gradients directly requires strong assumptions about problem-specific parameters. Influence maximization and profit maximization form a family of problems that model choosing advertising resource allocations to maximize the expected number of customers, where there is an underlying diffusion model for how advertising resources spent (stochastically) activate customers over a social network. For common diffusion models, the objective function is known to be DR-submodular (see for instance [3] or [15]). The revenue (expected number of activated customers) is a monotone objective function; total profit (revenue from activated customers minus advertising costs) is a non-monotone objective. One significant challenge with these problems is that the objective function (and the gradients) cannot be analytically evaluated for general (non-bipartite) networks, even if all the underlying diffusion model parameters are known exactly. The mildest assumptions on knowledge/observability of the network diffusions for offline variants (respectively actions for online variants), especially fitting for user privacy and/or third-party access, leads to instantiations of queries as the agent selecting an advertising allocation within the budget (i.e., feasible point) and observing a (stochastic) count of activated customers. This corresponds to stochastic value oracle queries over the feasible region (respectively bandit feedback for online variants).

D Comments on previous results in literature

Construction of $\mathcal{K}'$ and error estimate in [8] In [8], the set $\mathcal{K}' + \delta \mathbf{1}$ plays a role similar to the set $\mathcal{K}_\delta$ defined in this paper. Algorithm 2, in the case with access to value oracle for monotone DR-submodular function with the constraint set $\mathcal{K}$, such that $\text{aff}(\mathcal{K}) = \mathbb{R}^d$ and $\mathbf{0} \in \mathcal{K}$, reduced to BBCG algorithm in [8] if we replace $\mathcal{K}_\delta$ with their construction of $\mathcal{K}' + \delta \mathbf{1}$. In their paper, $\mathcal{K}'$ is defined by

$$\mathcal{K}' := (\mathcal{K} - \delta \mathbf{1}) \cap [0, 1 - 2\delta]^d. \quad (15)$$

There are a few issues with this construction and the subsequent analysis that requires more care.

1. *The BBCG algorithm almost always needs to be able to query the value oracle outside the feasible set.*

We have

$$\mathcal{K}' + \delta \mathbf{1} = \mathcal{K} \cap [\delta, 1 - \delta]^d.$$

The BBCG algorithm starts at $\delta \mathbf{1}$ and behaves similar to Algorithm 2 in the monotone $\mathbf{0} \in \mathcal{K}$ case. It follows that the set of points that BBCG requires to be able to query is

$$Q_\delta := \mathbb{B}_\delta(\text{convex-hull}((\mathcal{K}' + \delta \mathbf{1}) \cup \{\delta \mathbf{1}\})) = \mathbb{B}_\delta(\text{convex-hull}(\mathcal{K} \cup \{\delta \mathbf{1}\}) \cap [\delta, 1 - \delta]^d).$$

If $\mathbf{1} \in \mathcal{K}$, then the problem becomes trivial since F is monotone. If $\mathcal{K}$ is contained in the boundary of $[0, 1]^d$, then we need to restrict ourselves to the affine subspace containing $\mathcal{K}$ and solve the problem in a lower dimension in order to be able to use BBCG algorithm as $\mathcal{K}'$ will be empty otherwise. We want to show that in all other cases, $Q_\delta \setminus \mathcal{K} \neq \emptyset$. If $\mathcal{K}'$ is non-empty and $\mathbf{1} \notin \mathcal{K}$, then let $\mathbf{x}_\delta$ be a maximizer of $\|\cdot\|_\infty$ over $\mathcal{K}' + \delta \mathbf{1}$. If $\mathbf{x}_\delta \neq (1 - \delta)\mathbf{1}$, then there is a point $\mathbf{y} \in \mathbb{B}_\delta(\mathbf{x}_\delta) \cap [\delta, 1 - \delta]^d \subseteq Q_\delta$ such that $\mathbf{y} > \mathbf{x}$ which implies that $\mathbf{y} \notin \mathcal{K}$. Therefore, we only need to prove the statement when $(1 - \delta)\mathbf{1} \in \mathcal{K} \cap [\delta, 1 - \delta]^d$ for all small δ . In this case, since $\mathcal{K}$ is closed, we see that $(1 - \delta)\mathbf{1} \rightarrow \mathbf{1} \in \mathcal{K}$. In other words, except in trivial cases, BBCG always requires being able to query outside the feasible set.

2. *The exact error bound could be arbitrarily far away from the correct error bound depending on the geometry of the constraint set.*

In Equation (69) in the appendix of [8], it is mentioned that

$$\tilde{F}(\mathbf{x}_\delta^*) \geq \tilde{F}(\mathbf{x}^*) - \delta G \sqrt{d}, \quad (16)$$

where $\mathbf{x}^*$ is the optimal solution and $\mathbf{x}_\delta^*$ is the optimal solution within $\mathcal{K}' + \delta \mathbf{1}$ and G is the Lipschitz constant. Next we construct an example where this inequality does not hold.

Consider the set $\mathcal{K} = \{(x, y) \in [0, 1]^2 \mid x + \lambda y \leq 1\}$ for some value of λ to be specified and let $F((x, y)) = Gx$. Clearly we have $\mathbf{x}^* = (1, 0)$. Thus, for any $\delta > 0$, we have

$$\mathcal{K}' + \delta \mathbf{1} = \{(x, y) \in [\delta, 1 - \delta]^2 \mid x + \lambda y \leq 1\}.$$

It follows that when $\lambda \leq \frac{1}{\delta} - 1$, then $\mathcal{K}'$ is non-empty and $\mathbf{x}_\delta^* = (1 - \lambda\delta, \delta)$. Then we have

$$\tilde{F}(\mathbf{x}_\delta^*) - \tilde{F}(\mathbf{x}^*) = -\lambda\delta G.$$

Therefore, (16) is correct if and only if $\lambda \leq \sqrt{d} = \sqrt{2}$. Since this does not hold in general as λ depends on the geometry of the convex set, this equation is not true in general making the overall proof incorrect. The issue here is that λ , which depends on the geometry of the convex set $\mathcal{K}$, should appear in (16). Without restricting ourselves to convex sets with “controlled” geometry and without including a term, such as $\frac{1}{r}$ in Theorem 1, we would not be able to use this method to obtain an error bound. We note that while their analysis has an issue, the algorithm is still fine. Using a proof technique similar to ours, their proof can be fixed, more precisely, we can modify (16) in a manner similar to (24) and (31), depending on the case, and that will help fix their proofs.

One-Shot Frank-Wolfe algorithm in [6] In [6], the authors claim their proposed algorithm, One-Shot Frank-Wolfe (OSFW), achieves a $(1 - \frac{1}{e})$ -regret for monotone DR-submodular maximization under semi-bandit feedback for general convex set with oracle access to the entire domain of F , i.e. $[0, 1]^d$. In their regret analysis in the last page of the supplementary material, the inequality $(1 - 1/T)^t \leq 1/e$ is used for all $0 \leq t \leq T - 1$. Such an inequality holds for $t = T$ but as t decreases, the value of $(1 - 1/T)^t$ becomes closer to 1 and the inequality fails. If we do not use this inequality and continue with the proof, we end up with the following approximation coefficient.

$$1 - \frac{1}{T} \sum_{t=0}^{T-1} (1 - 1/T)^t = 1 - \frac{1}{T} \cdot \frac{1 - (1 - 1/T)^T}{1 - (1 - 1/T)} = 1 - (1 - (1 - 1/T)^T) = (1 - 1/T)^T \sim \frac{1}{e}.$$

E Useful lemmas

Here we state some lemmas from the literature that we will need in our analysis of DR-submodular functions.

Lemma 1 (Lemma 2.2 of [24]). *For any two vectors $\mathbf{x}, \mathbf{y} \in [0, 1]^d$ and any continuously differentiable non-negative DR-submodular function F we have*

$$F(\mathbf{x} \vee \mathbf{y}) \geq (1 - \|\mathbf{x}\|_\infty)F(\mathbf{y}).$$

The following lemma can be traced back to [16] (see Inequality 7.5 in the arXiv version), and is also explicitly stated and proved in [12].

Lemma 2 (Lemma 1 of [12]). *For every two vectors $\mathbf{x}, \mathbf{y} \in [0, 1]^d$ and any continuously differentiable non-negative DR-submodular function F we have*

$$\langle \nabla F(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle \geq F(\mathbf{x} \vee \mathbf{y}) + F(\mathbf{x} \wedge \mathbf{y}) - 2F(\mathbf{x}).$$

F Smoothing trick

The following Lemma is well-known when $\text{aff}(\mathcal{D}) = \mathbb{R}^d$ (e.g., Lemma 1 in [8], Lemma 7 in [31]). The proof in the general case is similar to the special case $\text{aff}(\mathcal{D}) = \mathbb{R}^d$.

Lemma 3. *If $F : \mathcal{D} \rightarrow \mathbb{R}$ is DR-submodular, G -Lipschitz continuous, and L -smooth, then so is $\tilde{F}_\delta$ and for any $\mathbf{x} \in \mathcal{D}$ such that $\mathbb{B}_\delta^{\text{aff}(\mathcal{D})}(\mathbf{x}) \subseteq \mathcal{D}$, we have*

$$\|\tilde{F}_\delta(\mathbf{x}) - F(\mathbf{x})\| \leq \delta G.$$

Moreover, if F is monotone, then so is $\tilde{F}_\delta$.

Proof. Let $A := \text{aff}(\mathcal{D})$ and $A_0 := \text{aff}(\mathcal{D}) - \mathbf{x}$ for some $\mathbf{x} \in \mathcal{D}$. Using the assumption that F is G -Lipschitz continuous, we have

$$\begin{aligned} |\tilde{F}(\mathbf{x}) - \tilde{F}(\mathbf{y})| &= \left| \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})} [F(\mathbf{x} + \delta\mathbf{v}) - F(\mathbf{y} + \delta\mathbf{v})] \right| \\ &\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})} [|F(\mathbf{x} + \delta\mathbf{v}) - F(\mathbf{y} + \delta\mathbf{v})|] \\ &\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})} [G\|(\mathbf{x} + \delta\mathbf{v}) - (\mathbf{y} + \delta\mathbf{v})\|] \\ &= G\|\mathbf{x} - \mathbf{y}\|, \end{aligned}$$

and

$$\begin{aligned}
|\tilde{F}(\mathbf{x}) - F(\mathbf{x})| &= |\mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{x} + \delta \mathbf{v}) - F(\mathbf{x})]| \\
&\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[|F(\mathbf{x} + \delta \mathbf{v}) - F(\mathbf{x})|] \\
&\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[G\delta\|\mathbf{v}\|] \\
&\leq \delta G.
\end{aligned}$$

If F is G -Lipschitz continuous and continuous DR-submodular, then F is differentiable and we have $\nabla F(\mathbf{x}) \geq \nabla F(\mathbf{y})$ for $\forall \mathbf{x} \leq \mathbf{y}$. By definition of $\tilde{F}$, we see that $\tilde{F}$ is also differentiable and

$$\begin{aligned}
\nabla \tilde{F}(\mathbf{x}) - \nabla \tilde{F}(\mathbf{y}) &= \nabla \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{x} + \delta \mathbf{v})] - \nabla \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{y} + \delta \mathbf{v})] \\
&= \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[\nabla F(\mathbf{x} + \delta \mathbf{v}) - \nabla F(\mathbf{y} + \delta \mathbf{v})] \\
&\geq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[0] = 0,
\end{aligned}$$

for all $\mathbf{x} \leq \mathbf{y}$.

If F is L -smooth, then we have $\|\nabla F(\mathbf{x}) - \nabla F(\mathbf{y})\| \leq L\|\mathbf{x} - \mathbf{y}\|$, for all $\mathbf{x}, \mathbf{y} \in \mathcal{D}$. Therefore, we have

$$\begin{aligned}
\|\nabla \tilde{F}(\mathbf{x}) - \nabla \tilde{F}(\mathbf{y})\| &= \|\nabla \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{x} + \delta \mathbf{v})] - \nabla \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{y} + \delta \mathbf{v})]\| \\
&= \|\mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[\nabla F(\mathbf{x} + \delta \mathbf{v})] - \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[\nabla F(\mathbf{y} + \delta \mathbf{v})]\| \\
&\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[\|\nabla F(\mathbf{x} + \delta \mathbf{v}) - \nabla F(\mathbf{y} + \delta \mathbf{v})\|] \\
&\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[L\|\mathbf{x} - \mathbf{y}\|] = L\|\mathbf{x} - \mathbf{y}\|,
\end{aligned}$$

for all $\mathbf{x} \leq \mathbf{y}$.

If F is monotone, then we have $F(\mathbf{x}) \leq F(\mathbf{y})$ for all $\mathbf{x} \leq \mathbf{y}$. Therefore

$$\begin{aligned}
\tilde{F}(\mathbf{x}) - \tilde{F}(\mathbf{y}) &= \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{x} + \delta \mathbf{v})] - \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{y} + \delta \mathbf{v})] \\
&= \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[F(\mathbf{x} + \delta \mathbf{v}) - F(\mathbf{y} + \delta \mathbf{v})] \\
&\leq \mathbb{E}_{\mathbf{v} \sim \mathbb{B}_1^{A_0}(\mathbf{0})}[0] = 0,
\end{aligned}$$

for all $\mathbf{x} \leq \mathbf{y}$. Hence $\tilde{F}$ is also monotone. $\square$

Lemma 4 (Lemma 10 of [27]). *Let $\mathcal{D} \subseteq \mathbb{R}^d$ such that $\text{aff}(\mathcal{D}) = \mathbb{R}^d$. Assume $F : \mathcal{D} \rightarrow \mathbb{R}$ is a G -Lipschitz continuous function and let $\tilde{F}$ be its δ -smoothed version. For any $\mathbf{z} \in \mathcal{D}$ such that $\mathbb{B}_\delta(\mathbf{z}) \subseteq \mathcal{D}$, we have*

$$\begin{aligned}
\mathbb{E}_{\mathbf{u} \sim S^{d-1}} \left[\frac{d}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} \right] &= \nabla \tilde{F}(\mathbf{z}), \\
\mathbb{E}_{\mathbf{u} \sim S^{d-1}} \left[\left\| \frac{d}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} - \nabla \tilde{F}(\mathbf{z}) \right\|^2 \right] &\leq CdG^2,
\end{aligned}$$

where C is a constant.

When the convex feasible region $\mathcal{K}$ lies in an affine subspace, we cannot employ the standard spherical sampling method. We extend Lemma 4 to that case.

Lemma 5. *Let $\mathcal{D} \subseteq \mathbb{R}^d$ and $A := \text{aff}(\mathcal{D})$. Also let A_0 be the translation of A that contains 0 and let $k = \dim(A)$. Assume $F : \mathcal{D} \rightarrow \mathbb{R}$ is a G -Lipschitz continuous function and let $\tilde{F}$ be its δ -smoothed version. For any $\mathbf{z} \in \mathcal{D}$ such that $\mathbb{B}_\delta^A(\mathbf{z}) \subseteq \mathcal{D}$, we have*

$$\begin{aligned}
\mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap A_0} \left[\frac{k}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} \right] &= \nabla \tilde{F}(\mathbf{z}), \\
\mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap A_0} \left[\left\| \frac{k}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} - \nabla \tilde{F}(\mathbf{z}) \right\|^2 \right] &\leq CkG^2,
\end{aligned}$$

where C is the constant in Lemma 4.

Proof. First consider the case where $A = \mathbb{R}^k \times (0, \dots, 0)$. In this case, we restrict ourselves to first k coordinates and see that the problem reduces to Lemma 4.

For the general case, let O be an orthonormal transformation that maps $\mathbb{R}^k \times (0, \dots, 0)$ into A_0 . Now define $\mathcal{D}' = O^{-1}(\mathcal{D} - \mathbf{z})$ and $F' : \mathcal{D}' \rightarrow \mathbb{R} : x \mapsto F(O(x) + \mathbf{z})$. Let $\tilde{F}'$ be the δ -smoothed version of F' . Note that $O(\nabla \tilde{F}'(0)) = \nabla \tilde{F}(\mathbf{z})$. On the other hand, we have

$$\text{aff}(\mathcal{D}') = O^{-1}(A - \mathbf{z}) = O^{-1}(A_0) = \mathbb{R}^k \times (0, \dots, 0).$$

Therefore

$$\mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\frac{k}{2\delta} (F'(\delta \mathbf{u}) - F'(-\delta \mathbf{u})) \mathbf{u} \right] = \nabla \tilde{F}'(0),$$

and

$$\mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\left\| \frac{k}{2\delta} (F'(\delta \mathbf{u}) - F'(-\delta \mathbf{u})) \mathbf{u} - \nabla \tilde{F}'(0) \right\|^2 \right] \leq CkG^2.$$

Hence, if we set $\mathbf{v} = O^{-1}(\mathbf{u})$, we have

$$\begin{aligned} \mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap A_0} \left[\frac{k}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} \right] \\ &= \mathbb{E}_{\mathbf{v} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\frac{k}{2\delta} (F'(\delta \mathbf{v}) - F'(-\delta \mathbf{v})) O(\mathbf{v}) \right] \\ &= O \left(\mathbb{E}_{\mathbf{v} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\frac{k}{2\delta} (F'(\delta \mathbf{v}) - F'(-\delta \mathbf{v})) \mathbf{v} \right] \right) \\ &= O(\nabla \tilde{F}'(0)) \\ &= \nabla \tilde{F}(\mathbf{z}). \end{aligned}$$

Similarly

$$\begin{aligned} \mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap A_0} \left[\left\| \frac{k}{2\delta} (F(\mathbf{z} + \delta \mathbf{u}) - F(\mathbf{z} - \delta \mathbf{u})) \mathbf{u} - \nabla \tilde{F}(\mathbf{z}) \right\|^2 \right] \\ &= \mathbb{E}_{\mathbf{v} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\left\| \frac{k}{2\delta} (F'(\delta \mathbf{v}) - F'(-\delta \mathbf{v})) O(\mathbf{v}) - O(\nabla \tilde{F}'(0)) \right\|^2 \right] \\ &= \mathbb{E}_{\mathbf{v} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\left\| O \left(\frac{k}{2\delta} (F'(\delta \mathbf{v}) - F'(-\delta \mathbf{v})) \mathbf{v} - \nabla \tilde{F}'(0) \right) \right\|^2 \right] \\ &= \mathbb{E}_{\mathbf{v} \sim S^{d-1} \cap (\mathbb{R}^k \times (0, \dots, 0))} \left[\left\| \frac{k}{2\delta} (F'(\delta \mathbf{v}) - F'(-\delta \mathbf{v})) \mathbf{v} - \nabla \tilde{F}'(0) \right\|^2 \right] \\ &\leq CkG^2. \end{aligned} \quad \square$$

Remark 4. Note that the same argument may be applied to obtain the one-point gradient estimator:

$$\mathbb{E}_{\mathbf{u} \sim S^{d-1} \cap A_0} \left[\frac{k}{\delta} F(\mathbf{z} + \delta \mathbf{u}) \mathbf{u} \right] = \nabla \tilde{F}(\mathbf{z}).$$

G Construction of $\mathcal{K}_\delta$

Lemma 6. Let $\mathcal{K} \subseteq [0, 1]^d$ be a convex set containing the origin. Then for any choice of $\mathbf{c}$ and r with $\mathbb{B}_r^{\text{aff}(\mathcal{K})}(\mathbf{c}) \subseteq \mathcal{K}$, we have

$$\argmin_{\mathbf{z} \in \mathcal{K}_\delta} \|\mathbf{z}\|_\infty = \frac{\delta}{r} \mathbf{c} \quad \text{and} \quad \min_{\mathbf{z} \in \mathcal{K}_\delta} \|\mathbf{z}\|_\infty \leq \frac{\delta}{r}.$$

Proof. The claim follows immediately from the definition and the fact that $\|\mathbf{c}\|_\infty \leq 1$. $\square$

Lemma 7. Let $\mathcal{K}$ be an arbitrary convex set, $D := \text{Diam}(\mathcal{K})$ and $\delta' := \frac{\delta D}{r}$. We have

$$\mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\mathcal{K}_\delta) \subseteq \mathcal{K} \subseteq \mathbb{B}_{\delta'}^{\text{aff}(\mathcal{K})}(\mathcal{K}_\delta).$$

Proof. Define $\psi : \mathcal{K} \rightarrow \mathcal{K}_\delta := \mathbf{x} \mapsto (1 - \frac{\delta}{r})\mathbf{x} + \frac{\delta}{r}\mathbf{c}$. Let $\mathbf{y} \in \mathcal{K}_\delta$ and $\mathbf{x} = \psi^{-1}(\mathbf{y})$. Then

$$\begin{aligned} \mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\mathbf{y}) &= \mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\psi(\mathbf{x})) = \mathbb{B}_\delta^{\text{aff}(\mathcal{K})}((1 - \frac{\delta}{r})\mathbf{x} + \frac{\delta}{r}\mathbf{c}) \\ &= (1 - \frac{\delta}{r})\mathbf{x} + \mathbb{B}_\delta^{\text{aff}(\mathcal{K})}(\frac{\delta}{r}\mathbf{c}) = (1 - \frac{\delta}{r})\mathbf{x} + \frac{\delta}{r}\mathbb{B}_r^{\text{aff}(\mathcal{K})}(\mathbf{c}) \subseteq \mathcal{K}, \end{aligned}$$

where the last inclusion follows from the fact that $\mathcal{K}$ is convex and contains both $\mathbf{x}$ and $\mathbb{B}_r^{\text{aff}(\mathcal{K})}(\mathbf{c})$. On the other hand, for any $\mathbf{x} \in \mathcal{K} \subseteq \text{aff}(\mathcal{K})$, we have

$$\|\psi(\mathbf{x}) - \mathbf{x}\| = \frac{\delta}{r}\|\mathbf{x} - \mathbf{c}\| < \frac{\delta}{r}D = \delta'.$$

Therefore

$$\mathbf{x} \in \mathbb{B}_{\delta'}(\psi(\mathbf{x})) \cap \text{aff}(\mathcal{K}) = \mathbb{B}_{\delta'}^{\text{aff}(\mathcal{K})}(\psi(\mathbf{x})) \subseteq \mathbb{B}_{\delta'}^{\text{aff}(\mathcal{K})}(\mathcal{K}_\delta). \quad \square$$

Choice of $\mathbf{c}$ and r While the results hold for any choice of $\mathbf{c} \in \mathcal{K}$ and r with $\mathbb{B}_r^{\text{aff}}(\mathbf{c}) \subseteq \mathcal{K}$, as can be seen in Theorem 2, the approximation errors depends linearly on $1/r$. Therefore, it is natural to choose the point $\mathbf{c}$ that maximizes the value of r , the *Chebyshev center* of $\mathcal{K}$.

Analytic Constraint Model — Polytope When the feasible region $\mathcal{K}$ is characterized by a set of q linear constraints $\mathbf{A}\mathbf{x} \leq \mathbf{b}$ with a known coefficient matrix $\mathbf{A} \in \mathbb{R}^{q \times d}$ and vector $\mathbf{b} \in \mathbb{R}^q$, thus $\mathcal{K}$ is a polytope, by the linearity of the transformation (7), the shrunken feasible region $\mathcal{K}_\delta$ is similarly characterized by a (translated) set of q linear constraints $\mathbf{A}\mathbf{x} \leq (1 - \frac{\delta}{r})\mathbf{b} + \frac{\delta}{r}\mathbf{A}\mathbf{c}$.

H Variance reduction via momentum

In order to prove main regret bounds, we need the following variance reduction lemma, which is crucial in characterizing how much the variance of the gradient estimator can be reduced by using momentum. This lemma appears in [6] and it is a slight improvement of Lemma 2 in [22] and Lemma 5 in [23].

Lemma 8 (Theorem 3 of [6]). Let $\{\mathbf{a}_n\}_{n=0}^N$ be a sequence of points in $\mathbb{R}^d$ such that $\|\mathbf{a}_n - \mathbf{a}_{n-1}\| \leq G_0/(n+s)$ for all $1 \leq n \leq N$ with fixed constants $G_0 \geq 0$ and $s \geq 3$. Let $\{\tilde{\mathbf{a}}_n\}_{n=1}^N$ be a sequence of random variables such that $\mathbb{E}[\tilde{\mathbf{a}}_n | \mathcal{F}_{n-1}] = \mathbf{a}_n$ and $\mathbb{E}[\|\tilde{\mathbf{a}}_n - \mathbf{a}_n\|^2 | \mathcal{F}_{n-1}] \leq \sigma^2$ for every $n \geq 0$, where $\mathcal{F}_{n-1}$ is the σ -field generated by $\{\tilde{\mathbf{a}}_i\}_{i=1}^n$ and $\mathcal{F}_0 = \emptyset$. Let $\{\mathbf{d}_n\}_{n=0}^N$ be a sequence of random variables where $\mathbf{d}_0$ is fixed and subsequent $\mathbf{d}_n$ are obtained by the recurrence

$$\mathbf{d}_n = (1 - \rho_n)\mathbf{d}_{n-1} + \rho_n\tilde{\mathbf{a}}_n \quad (17)$$

with $\rho_n = \frac{2}{(n+s)^{2/3}}$. Then, we have

$$\mathbb{E}[\|\mathbf{a}_n - \mathbf{d}_n\|^2] \leq \frac{Q}{(n+s+1)^{2/3}}, \quad (18)$$

where $Q := \max\{\|\mathbf{a}_0 - \mathbf{d}_0\|^2(s+1)^{2/3}, 4\sigma^2 + 3G_0^2/2\}$.

We now analyze the variance of our gradient estimator, which, in the case when we only have access zeroth-order information, uses batched spherical sampling and momentum for gradient estimation. Calculations similar to the proof of the following Lemma, in the value oracle case, appear in the proof of Theorem 2 in [8]. The main difference is that here we consider a more general smoothing trick and therefore we estimate the gradient along the affine hull of $\mathcal{K}$.

Lemma 9. Under the assumptions of Theorem 1, in Algorithm 2, we have

$$\mathbb{E}[\|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|^2] \leq \frac{Q}{(n+4)^{2/3}},$$

for all $1 \leq n \leq N$ where $\mathcal{L} = \text{aff}(\mathcal{K})$,

$$Q = \begin{cases} 0 & \text{det. grad. oracle,} \\ \max\{4^{2/3}G^2, 6L^2D^2 + \frac{4\sigma_1^2}{B}\} & \text{stoch. grad. oracle with variance } \sigma_1^2 > 0, \\ \max\{4^{2/3}G^2, 6L^2D^2 + \frac{4CkG^2 + 2k^2\sigma_0^2/\delta^2}{B}\} & \text{value oracle with variance } \sigma_0^2 \geq 0, \end{cases}$$

C is a constant and $D = \text{diam}(\mathcal{K})$.

Remark 5. As we will see in the proof of Theorem 1, except for the case with deterministic gradient oracle, the dominating term in the approximation error is a constant multiple of

$$\frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|^2 \right].$$

Therefore, any improvement in Lemma 9 will result in direct improvement of the approximation error.

Proof. If we have access to a deterministic gradient oracle, then the claim is trivial. Let $\mathcal{F}_1 := \emptyset$ and $\mathcal{F}_n$ be the σ -field generated by $\{\bar{\mathbf{g}}_1, \dots, \bar{\mathbf{g}}_{n-1}\}$ and let

$$\sigma^2 = \begin{cases} \frac{\sigma_1^2}{B} & \text{stoch. grad. oracle with variance } \sigma_1^2 > 0, \\ \frac{CkG^2 + k^2\sigma_0^2/2\delta^2}{B} & \text{value oracle with variance } \sigma_0^2 \geq 0. \end{cases}$$

Let $\mathcal{L}_0$ denote the linear space $\mathcal{L} - \mathbf{x}$ for some $\mathbf{x} \in \mathcal{L}$. If we have access to a stochastic gradient oracle, then $\mathbf{g}_n$ is computed by taking the average of B gradient samples of $P_{\mathcal{L}_0}(\hat{G}(\mathbf{z}))$, i.e. the projection of $\hat{G}(\mathbf{z})$ onto the linear space $\mathcal{L}_0$. Since $P_{\mathcal{L}_0}$ is a 1-Lipschitz linear map, we see that

$$\mathbb{E}[P_{\mathcal{L}_0}(\hat{G}(\mathbf{z}))] = P_{\mathcal{L}_0}(\nabla \tilde{F}(\mathbf{z})) = \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z})$$

and

$$\begin{aligned} \mathbb{E} \left[\left\| P_{\mathcal{L}_0}(\hat{G}(\mathbf{z})) - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}) \right\|^2 \right] &= \mathbb{E} \left[\left\| P_{\mathcal{L}_0}(\hat{G}(\mathbf{z})) - P_{\mathcal{L}_0}(\nabla \tilde{F}(\mathbf{z})) \right\|^2 \right] \\ &\leq \mathbb{E} \left[\left\| \hat{G}(\mathbf{z}) - \nabla \tilde{F}(\mathbf{z}) \right\|^2 \right] \leq \sigma_1^2. \end{aligned}$$

Note that, in cases where we have access to a gradient oracle, we have $\delta = 0$ and $\tilde{F} = F$. Therefore

$$\mathbb{E}[\mathbf{g}_n | \mathcal{F}_{n-1}] = \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \quad \text{and} \quad \mathbb{E}[\|\mathbf{g}_n - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n)\|^2 | \mathcal{F}_{n-1}] \leq \frac{\sigma_1^2}{B} = \sigma^2.$$

Next we assume that we have access to a value oracle. By the unbiasedness of $\hat{F}$ and Lemma 5, we have

$$\begin{aligned} \mathbb{E} \left[\frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - \hat{F}(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} | \mathcal{F}_{n-1} \right] &= \mathbb{E} \left[\mathbb{E} \left[\frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - \hat{F}(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} | \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] | \mathcal{F}_{n-1} \right] \\ &= \mathbb{E} \left[\frac{k}{2\delta} (F(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} | \mathcal{F}_{n-1} \right] \\ &= \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \end{aligned}$$

and

$$\begin{aligned}
& \mathbb{E} \left[\left\| \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - \hat{F}(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right\|^2 \middle| \mathcal{F}_{n-1} \right] \\
&= \mathbb{E} \left[\mathbb{E} \left[\left\| \frac{k}{2\delta} (F(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right. \right. \right. \\
&\quad \left. \left. + \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^+)) \mathbf{u}_{n,i} \right. \right. \\
&\quad \left. \left. - \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^-) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} \right\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\leq \mathbb{E} \left[\mathbb{E} \left[\left\| \frac{k}{2\delta} (F(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\quad + \mathbb{E} \left[\mathbb{E} \left[\left\| \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^+)) \mathbf{u}_{n,i} \right\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\quad + \mathbb{E} \left[\mathbb{E} \left[\left\| \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^-) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} \right\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\leq \mathbb{E} \left[\left\| \frac{k}{2\delta} (F(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right\|^2 \middle| \mathcal{F}_{n-1} \right] \\
&\quad + \frac{k^2}{4\delta^2} \mathbb{E} \left[\mathbb{E} \left[|\hat{F}(\mathbf{y}_{n,i}^+) - F(\mathbf{y}_{n,i}^+)|^2 \cdot \|\mathbf{u}_{n,i}\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\quad + \frac{k^2}{4\delta^2} \mathbb{E} \left[\mathbb{E} \left[|\hat{F}(\mathbf{y}_{n,i}^-) - F(\mathbf{y}_{n,i}^-)|^2 \cdot \|\mathbf{u}_{n,i}\|^2 \middle| \mathcal{F}_{n-1}, \mathbf{u}_{n,i} \right] \middle| \mathcal{F}_{n-1} \right] \\
&\leq CkG^2 + \frac{k^2}{4\delta^2} \sigma_0^2 + \frac{k^2}{4\delta^2} \sigma_0^2 \\
&= CkG^2 + \frac{k^2}{2\delta^2} \sigma_0^2.
\end{aligned}$$

So we have

$$\mathbb{E} [\mathbf{g}_n | \mathcal{F}_{n-1}] = \mathbb{E} \left[\frac{1}{B} \sum_{i=1}^B \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - \hat{F}(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} \middle| \mathcal{F}_{n-1} \right] = \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n),$$

and

$$\begin{aligned}
& \mathbb{E} \left[\left\| \mathbf{g}_n - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right\|^2 \middle| \mathcal{F}_{n-1} \right] \\
&= \frac{1}{B^2} \sum_{i=1}^B \mathbb{E} \left[\left\| \frac{k}{2\delta} (\hat{F}(\mathbf{y}_{n,i}^+) - \hat{F}(\mathbf{y}_{n,i}^-)) \mathbf{u}_{n,i} - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) \right\|^2 \middle| \mathcal{F}_{n-1} \right] \\
&\leq \frac{CkG^2 + \frac{k^2}{2\delta^2} \sigma_0^2}{B} = \sigma^2.
\end{aligned}$$

Using Lemma 8 with $\mathbf{d}_n = \bar{\mathbf{g}}_n$, $\tilde{\mathbf{a}}_n = \mathbf{g}_n$, $\mathbf{a}_n = \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n)$ for all $n \geq 1$, $\mathbf{a}_0 = \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_1)$, $G_0 = 2LD$ and $s = 3$, we have

$$\mathbb{E} [\|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|^2] \leq \frac{Q'}{(n+4)^{2/3}}, \quad (19)$$

where $Q' = \max\{\|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_1)\|^2 4^{2/3}, 6L^2 D^2 + 4\sigma^2\}$. Note that by Lemma 3, we have $\|\nabla(\tilde{F}|_{\mathcal{L}})(x)\| \leq G$, thus we have $Q' \leq Q$. $\square$

I Proof of Theorem 1 for monotone maps over convex sets containing zero

Proof. By the definition of $\mathbf{z}_n$, we have $\mathbf{z}_n = \mathbf{z}_1 + \sum_{i=1}^{n-1} \frac{\mathbf{v}_i}{N}$. Therefore $\mathbf{z}_n - \mathbf{z}_1$ is a convex combination of $\mathbf{v}_n$'s and 0 which belong to $\mathcal{K}_\delta - \mathbf{z}_1$ and therefore $\mathbf{z}_n - \mathbf{z}_1 \in \mathcal{K}_\delta - \mathbf{z}_1$. Hence we have $\mathbf{z}_n \in \mathcal{K}_\delta \subseteq \mathcal{K}$ for all $1 \leq n \leq N + 1$.

Let $\mathcal{L} := \text{aff}(\mathcal{K})$. According to Lemma 3, the function $\tilde{F}$ is L -smooth. So we have

$$\begin{aligned} \tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) &\geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_{n+1} - \mathbf{z}_n \rangle - \frac{L}{2} \|\mathbf{z}_{n+1} - \mathbf{z}_n\|^2 \\ &= \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{v}_n \rangle - \frac{\varepsilon^2 L}{2} \|\mathbf{v}_n\|^2 \\ &\geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{v}_n \rangle - \frac{\varepsilon^2 L}{2} D^2 \\ &= \varepsilon \left(\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \right) - \frac{\varepsilon^2 L D^2}{2}. \end{aligned} \tag{20}$$

Let $\mathbf{z}_\delta^* := \operatorname{argmax}_{\mathbf{z} \in \mathcal{K}_\delta - \mathbf{z}_1} \tilde{F}(\mathbf{z})$. We have $\mathbf{z}_\delta^* \in \mathcal{K}_\delta - \mathbf{z}_1$, which implies that $\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \rangle$. Therefore

$$\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \rangle = \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle + \langle \bar{\mathbf{g}}_n - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle$$

Hence we obtain

$$\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle - \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* - \mathbf{v}_n \rangle$$

Using the Cauchy-Schwartz inequality, we have

$$\langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* - \mathbf{v}_n \rangle \leq \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| \|\mathbf{z}_\delta^* - \mathbf{v}_n\| \leq D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|$$

Therefore

$$\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle - D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|.$$

Plugging this into 20, we see that

$$\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 L D^2}{2}. \tag{21}$$

On the other hand, we have $\mathbf{z}_\delta^* \geq (\mathbf{z}_\delta^* - \mathbf{z}_n) \vee 0$. Since F is monotone continuous DR-submodular, by Lemma 3, so is $\tilde{F}$. Moreover monotonicity of $\tilde{F}$ implies that $\nabla(\tilde{F}|_{\mathcal{L}})$ is non-negative in positive directions. Therefore we have

$$\begin{aligned} \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \rangle &\geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), (\mathbf{z}_\delta^* - \mathbf{z}_n) \vee 0 \rangle && \text{(monotonicity)} \\ &\geq \tilde{F}(\mathbf{z}_n + ((\mathbf{z}_\delta^* - \mathbf{z}_n) \vee 0)) - \tilde{F}(\mathbf{z}_n) && \text{(DR-submodularity)} \\ &= \tilde{F}(\mathbf{z}_\delta^* \vee \mathbf{z}_n) - \tilde{F}(\mathbf{z}_n) \\ &\geq \tilde{F}(\mathbf{z}_\delta^*) - \tilde{F}(\mathbf{z}_n) \end{aligned}$$

After plugging this into (21) and re-arranging terms, we obtain

$$h_{n+1} \leq (1 - \varepsilon)h_n + \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| + \frac{\varepsilon^2 L D^2}{2}$$

where $h_n := \tilde{F}(\mathbf{z}_\delta^*) - \tilde{F}(\mathbf{z}_n)$. After taking the expectation and using Lemma 9, we see that

$$\mathbb{E}(h_{n+1}) \leq (1 - \varepsilon)\mathbb{E}(h_n) + \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} + \frac{\varepsilon^2 L D^2}{2}.$$

Using the above inequality recursively and $1 - \varepsilon \leq 1$, we have

$$\mathbb{E}[h_{N+1}] \leq (1 - \varepsilon)^N \mathbb{E}[h_1] + \sum_{n=1}^N \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} + \frac{N \varepsilon^2 L D^2}{2}.$$

Note that we have $\varepsilon = 1/N$. Using the fact that $(1 - \frac{1}{N})^N \leq e^{-1}$ and

$$\begin{aligned} \sum_{n=1}^N \frac{DQ^{1/2}}{(n+4)^{1/3}} &\leq DQ^{1/2} \int_0^N \frac{dx}{(x+4)^{1/3}} \leq DQ^{1/2} \left(\frac{3}{2} (N+4)^{2/3} \right) \\ &\leq DQ^{1/2} \left(\frac{3}{2} (2N)^{2/3} \right) \leq 3DQ^{1/2} N^{2/3}, \end{aligned} \quad (22)$$

we see that

$$\mathbb{E}[h_{N+1}] \leq e^{-1} \mathbb{E}[h_1] + \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N}.$$

By re-arranging the terms and using the fact that $\tilde{F}$ is non-negative, we conclude

$$\begin{aligned} (1 - e^{-1})\tilde{F}(\mathbf{z}_\delta^*) - \mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] &\leq -e^{-1}\tilde{F}(\mathbf{z}_1) + \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} \\ &\leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N}. \end{aligned} \quad (23)$$

According to Lemma 3, we have $\tilde{F}(\mathbf{z}_{N+1}) \leq F(\mathbf{z}_{N+1}) + \delta G$. Moreover, using Lemma 7, we see that $\mathbf{z}^* \in \mathbb{B}_{\delta'}(\mathcal{K}_\delta)$ where $\delta' = \delta D/r$. Therefore, there is a point $\mathbf{y}^* \in \mathcal{K}_\delta$ such that $\|\mathbf{y}^* - \mathbf{z}^*\| \leq \delta'$.

$$\begin{aligned} \tilde{F}(\mathbf{z}_\delta^*) &\geq \tilde{F}(\mathbf{y}^* - \mathbf{z}_1) \geq \tilde{F}(\mathbf{y}^*) - G\|\mathbf{z}_1\| \\ &\geq F(\mathbf{y}^*) - (\|\mathbf{z}_1\| + \delta)G \geq F(\mathbf{z}^*) - (\|\mathbf{z}_1\| + \delta + \frac{\delta D}{r})G. \end{aligned}$$

According to Lemma 6, we have $\|\mathbf{z}_1\| \leq \sqrt{d}\|\mathbf{z}_1\|_\infty \leq \delta\sqrt{d}/r$.

$$\tilde{F}(\mathbf{z}_\delta^*) \geq F(\mathbf{z}^*) - (1 + \frac{\sqrt{d} + D}{r})\delta G. \quad (24)$$

After plugging these into 23, we see that

$$\begin{aligned} (1 - e^{-1})F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] &\leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \delta G(2 + \frac{\sqrt{d} + D}{r}). \end{aligned} \quad \square$$

J Proof of Theorem 1 for non-monotone maps over downward-closed convex sets

Proof. Similar to Appendix I, we see that $\mathbf{z}_n \in \mathcal{K}_\delta$ for all $1 \leq n \leq N+1$ and

$$\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \geq \varepsilon \left(\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \right) - \frac{\varepsilon^2 LD^2}{2}. \quad (25)$$

Let $\mathbf{z}_\delta^* := \operatorname{argmax}_{\mathbf{z} \in \mathcal{K}_\delta - \mathbf{z}_1} \tilde{F}(\mathbf{z})$. We have $\mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n = (\mathbf{z}_\delta^* - \mathbf{z}_n) \vee 0 \leq \mathbf{z}_\delta^*$. Therefore, since $\mathcal{K}_\delta$ is downward-closed, we have $\mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \in \mathcal{K}_\delta - \mathbf{z}_1$. On the other hand, $\mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \leq \mathbf{1} - \mathbf{z}_n$. Therefore, we have $\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle$, which implies that

$$\begin{aligned} &\langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \\ &= \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle + \langle \bar{\mathbf{g}}_n - \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle \\ &\quad + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \\ &= \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle - \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, -\mathbf{v}_n + \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle \end{aligned}$$

Using the Cauchy-Shwarz inequality, we see that

$$\begin{aligned} \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, -\mathbf{v}_n + \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle &\leq \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| \|(\mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n) - \mathbf{v}_n\| \\ &\leq D\|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|. \end{aligned}$$

where the last inequality follows from the fact that both $\mathbf{v}_n$ and $\mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n$ belong to $\mathcal{K}_\delta$. Therefore

$$\begin{aligned} & \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \\ & \geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle - D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|. \end{aligned}$$

Plugging this into Equation (25), we get

$$\begin{aligned} & \tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \\ & \geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 LD^2}{2}. \end{aligned} \quad (26)$$

Next we show that

$$1 - \|\mathbf{z}_n\|_\infty \geq (1 - \varepsilon)^{n-1} \left(1 - \frac{\delta}{r}\right), \quad (27)$$

for all $1 \leq n \leq N + 1$. We use induction on n to show that for each coordinate $1 \leq i \leq d$, we have $1 - [\mathbf{z}_n]_i \geq (1 - \varepsilon)^{n-1}$. For $n = 1$, the claim follows from Lemma 6. Assuming that the inequality is true for n , using the fact that $\mathbf{v}_n \leq \mathbf{1} - \mathbf{z}_n$, we have

$$\begin{aligned} 1 - [\mathbf{z}_{n+1}]_i &= 1 - [\mathbf{z}_n]_i - \varepsilon [\mathbf{v}_n]_i \geq 1 - [\mathbf{z}_n]_i - \varepsilon (1 - [\mathbf{z}_n]_i) \\ &= (1 - \varepsilon)(1 - [\mathbf{z}_n]_i) \geq (1 - \varepsilon)^n \left(1 - \frac{\delta}{r}\right), \end{aligned}$$

which completes the proof by induction.

Since $\tilde{F}$ is DR-submodular, it is concave along non-negative directions. Therefore, using Lemma 1 and Equation (27), we have

$$\begin{aligned} \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* \vee \mathbf{z}_n - \mathbf{z}_n \rangle &\geq \tilde{F}(\mathbf{z}_\delta^* \vee \mathbf{z}_n) - \tilde{F}(\mathbf{z}_n) \\ &\geq (1 - \|\mathbf{z}_n\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) - \tilde{F}(\mathbf{z}_n) \\ &\geq (1 - \varepsilon)^{n-1} \left(1 - \frac{\delta}{r}\right) \tilde{F}(\mathbf{z}_\delta^*) - \tilde{F}(\mathbf{z}_n). \end{aligned}$$

Plugging this into Equation (26), we get

$$\begin{aligned} & \tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \\ & \geq \varepsilon \left((1 - \varepsilon)^{n-1} \left(1 - \frac{\delta}{r}\right) \tilde{F}(\mathbf{z}_\delta^*) - \tilde{F}(\mathbf{z}_n) \right) - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 LD^2}{2}. \end{aligned}$$

Taking expectations of both sides and using Lemma 9, we see that

$$\mathbb{E}(\tilde{F}(\mathbf{z}_{n+1})) \geq (1 - \varepsilon) \mathbb{E}(\tilde{F}(\mathbf{z}_n)) + \varepsilon (1 - \varepsilon)^{n-1} \left(1 - \frac{\delta}{r}\right) \tilde{F}(\mathbf{z}_\delta^*) - \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{\varepsilon^2 LD^2}{2}.$$

Using this inequality recursively and Equation (22), we get

$$\begin{aligned} \mathbb{E}(\tilde{F}(\mathbf{z}_{N+1})) &\geq (1 - \varepsilon)^N \mathbb{E}(\tilde{F}(\mathbf{z}_1)) + N \varepsilon (1 - \varepsilon)^{N-1} \left(1 - \frac{\delta}{r}\right) \tilde{F}(\mathbf{z}_\delta^*) \\ &\quad - \sum_{n=1}^N \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{N \varepsilon^2 LD^2}{2} \\ &\geq (1 - \varepsilon)^N \mathbb{E}(\tilde{F}(\mathbf{z}_1)) + N \varepsilon (1 - \varepsilon)^{N-1} \left(1 - \frac{\delta}{r}\right) \tilde{F}(\mathbf{z}_\delta^*) \\ &\quad - 3 \varepsilon D Q^{1/2} N^{2/3} - \frac{N \varepsilon^2 LD^2}{2}. \end{aligned}$$

Since $\delta < \frac{r}{2}$ and $\varepsilon = 1/N$, we have

$$(1 - \varepsilon)^N = \left(1 - \frac{1}{N}\right) (1 - \varepsilon)^{N-1} \geq \frac{1}{2} (1 - \varepsilon)^{N-1} \geq \frac{\delta}{r} (1 - \varepsilon)^{N-1} \geq \frac{\delta}{r} (1 - \varepsilon)^{N-1}.$$

Since $\tilde{F}$ is non-negative and G -Lipschitz, this implies that

$$\begin{aligned}
\mathbb{E}(\tilde{F}(\mathbf{z}_{N+1})) &\geq \frac{\delta}{r}(1-\varepsilon)^{N-1}\mathbb{E}(\tilde{F}(\mathbf{z}_1)) + (1-\varepsilon)^{N-1}(1-\frac{\delta}{r})\tilde{F}(\mathbf{z}_\delta^*) \\
&\quad - 3\varepsilon DQ^{1/2}N^{2/3} - \frac{N\varepsilon^2 LD^2}{2} \\
&= (1-\varepsilon)^{N-1}\tilde{F}(\mathbf{z}_\delta^*) + \frac{\delta}{r}(1-\varepsilon)^{N-1}(\mathbb{E}(\tilde{F}(\mathbf{z}_1)) - \tilde{F}(\mathbf{z}_\delta^*)) \\
&\quad - 3\varepsilon DQ^{1/2}N^{2/3} - \frac{N\varepsilon^2 LD^2}{2} \\
&\geq (1-\varepsilon)^{N-1}\tilde{F}(\mathbf{z}_\delta^*) - \frac{\delta}{r}(1-\varepsilon)^{N-1}GD - 3\varepsilon DQ^{1/2}N^{2/3} - \frac{N\varepsilon^2 LD^2}{2}.
\end{aligned}$$

After setting $\varepsilon = 1/N$ and using $(1 - 1/N)^{N-1} \geq e^{-1}$, we see that

$$\begin{aligned}
e^{-1}\tilde{F}(\mathbf{z}_\delta^*) - \mathbb{E}(\tilde{F}(\mathbf{z}_{N+1})) &\leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \frac{\delta}{r}(1-\varepsilon)^{N-1}GD \\
&\leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \frac{\delta}{r}GD.
\end{aligned}$$

Using the argument presented in Appendix I, i.e. Lemma 3 and Equation 24, we conclude that

$$e^{-1}F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{LD^2}{2N} + \delta G(2 + \frac{\sqrt{d} + 2D}{r}). \quad \square$$

K Proof of Theorem 1 for monotone maps over general convex sets

Proof. Using the fact that $\tilde{F}$ is L -smooth, we have

$$\begin{aligned}
\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) &\geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_{n+1} - \mathbf{z}_n \rangle - \frac{L}{2}\|\mathbf{z}_{n+1} - \mathbf{z}_n\|^2 \\
&= \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{v}_n - \mathbf{z}_n \rangle - \frac{\varepsilon^2 L}{2}\|\mathbf{v}_n - \mathbf{z}_n\|^2 \\
&\geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{v}_n - \mathbf{z}_n \rangle - \frac{\varepsilon^2 LD^2}{2} \\
&= \varepsilon \left(\langle \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) + \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle \right) - \frac{\varepsilon^2 LD^2}{2}.
\end{aligned} \tag{28}$$

Let $\mathbf{z}_\delta^* := \operatorname{argmax}_{\mathbf{z} \in \mathcal{K}_\delta} \tilde{F}(\mathbf{z})$. Using the fact that $\langle \bar{\mathbf{g}}_n, \mathbf{v}_n \rangle \geq \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* \rangle$, we have

$$\begin{aligned}
&\langle \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle \\
&\geq \langle \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle \\
&= \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle - \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* - \mathbf{v}_n \rangle.
\end{aligned}$$

Using the Cauchy-Schwarz inequality, we see that

$$\langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{z}_\delta^* - \mathbf{v}_n \rangle \leq \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| \|\mathbf{z}_\delta^* - \mathbf{v}_n\| \leq D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|.$$

Therefore

$$\begin{aligned}
&\langle \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle + \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n, \mathbf{v}_n - \mathbf{z}_n \rangle \\
&\geq \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle - D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\|.
\end{aligned}$$

Plugging this into 28, we get

$$\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 LD^2}{2}. \tag{29}$$

Using Lemma 2 and the fact that $\tilde{F}$ is monotone, we see that

$$\begin{aligned}\langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle &\geq \tilde{F}(\mathbf{z}_\delta^* \vee \mathbf{z}_n) + \tilde{F}(\mathbf{z}_\delta^* \wedge \mathbf{z}_n) - 2\tilde{F}(\mathbf{z}_n) \\ &\geq \tilde{F}(\mathbf{z}_\delta^*) + \tilde{F}(\mathbf{z}_\delta^* \wedge \mathbf{z}_n) - 2\tilde{F}(\mathbf{z}_n) \\ &\geq \tilde{F}(\mathbf{z}_\delta^*) - 2\tilde{F}(\mathbf{z}_n).\end{aligned}$$

After plugging this into (29), we get

$$\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \geq \varepsilon \tilde{F}(\mathbf{z}_\delta^*) - 2\varepsilon \tilde{F}(\mathbf{z}_n) - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 LD^2}{2}.$$

After taking the expectation, using Lemma 9 and re-arranging the terms, we see that

$$\mathbb{E}[\tilde{F}(\mathbf{z}_{n+1})] \geq (1 - 2\varepsilon)\mathbb{E}[\tilde{F}(\mathbf{z}_n)] + \varepsilon \tilde{F}(\mathbf{z}_\delta^*) - \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{\varepsilon^2 LD^2}{2}.$$

Using this inequality recursively together with Equation (22) and the fact that $\tilde{F}$ is non-negative, we get

$$\begin{aligned}\mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] &\geq (1 - 2\varepsilon)^N \mathbb{E}[\tilde{F}(\mathbf{z}_1)] + \varepsilon \tilde{F}(\mathbf{z}_\delta^*) \sum_{n=1}^N (1 - 2\varepsilon)^{N-n} \\ &\quad - \sum_{n=1}^N \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{N\varepsilon^2 LD^2}{2}. \\ &\geq \frac{1}{2}(1 - 2\varepsilon)^N \mathbb{E}[\tilde{F}(\mathbf{z}_1)] + \varepsilon \tilde{F}(\mathbf{z}_\delta^*) \sum_{n=1}^N (1 - 2\varepsilon)^{N-n} \\ &\quad - 3\varepsilon D Q^{1/2} N^{2/3} - \frac{N\varepsilon^2 LD^2}{2} \\ &= \frac{1}{2}(1 - 2\varepsilon)^N \mathbb{E}[\tilde{F}(\mathbf{z}_1)] + \frac{1}{2}(1 - (1 - 2\varepsilon)^N) \mathbb{E}[\tilde{F}(\mathbf{z}_\delta^*)] \\ &\quad - 3\varepsilon D Q^{1/2} N^{2/3} - \frac{N\varepsilon^2 LD^2}{2} \\ &= \frac{1}{2} \tilde{F}(\mathbf{z}_\delta^*) - \frac{1}{2}(1 - 2\varepsilon)^N (\tilde{F}(\mathbf{z}_\delta^*) - \mathbb{E}[\tilde{F}(\mathbf{z}_1)]) \\ &\quad - 3\varepsilon D Q^{1/2} N^{2/3} - \frac{N\varepsilon^2 LD^2}{2} \\ &\geq \frac{1}{2} \tilde{F}(\mathbf{z}_\delta^*) - \frac{1}{2}(1 - 2\varepsilon)^N DG - 3\varepsilon D Q^{1/2} N^{2/3} - \frac{N\varepsilon^2 LD^2}{2}.\end{aligned}$$

Note that $(1 - \log(N)/N)^N \leq e^{-\log(N)} = 1/N$. Therefore, since $\varepsilon = \log(N)/2N$, we have

$$\mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] \geq \frac{1}{2} \tilde{F}(\mathbf{z}_\delta^*) - \frac{DG}{2N} - \frac{3DQ^{1/2} \log(N)}{2N^{1/3}} - \frac{LD^2 \log(N)^2}{8N}. \quad (30)$$

According to Lemma 3, we have $\tilde{F}(\mathbf{z}_{N+1}) \leq F(\mathbf{z}_{N+1}) + \delta G$. Moreover, using Lemma 7, we see that $\mathbf{z}^* \in \mathbb{B}_{\delta'}(\mathcal{K}_\delta)$ where $\delta' = \delta D/r$. Therefore, there is a point $\mathbf{y}^* \in \mathcal{K}_\delta$ such that $\|\mathbf{y}^* - \mathbf{z}^*\| \leq \delta'$.

$$\tilde{F}(\mathbf{z}_\delta^*) \geq \tilde{F}(\mathbf{y}^*) \geq F(\mathbf{y}^*) \geq F(\mathbf{z}^*) - \delta G \geq F(\mathbf{z}^*) - (\delta + \frac{\delta D}{r})G. \quad (31)$$

After plugging these into (30), we see that

$$\begin{aligned}\frac{1}{2} \tilde{F}(\mathbf{z}^*) - \mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] &\leq \frac{3DQ^{1/2} \log(N)}{2N^{1/3}} + \frac{4DG + LD^2 \log(N)^2}{8N} + \delta G(2 + \frac{D}{r}).\end{aligned}$$

which completes the proof. $\square$

L Proof of Theorem 1 for non-monotone maps over general convex sets

Proof. First we show that

$$1 - \|\mathbf{z}_n\|_\infty \geq (1 - \varepsilon)^{n-1}(1 - \|\mathbf{z}_1\|_\infty), \quad (32)$$

for all $1 \leq n \leq N + 1$. We use induction on n to show that for each coordinate $1 \leq i \leq d$, we have $1 - [\mathbf{z}_n]_i \geq (1 - \varepsilon)^{n-1}(1 - [\mathbf{z}_1]_i)$. The claim is obvious for $n = 1$. Assuming that the inequality is true for n , we have

$$\begin{aligned} 1 - [\mathbf{z}_{n+1}]_i &= 1 - (1 - \varepsilon)[\mathbf{z}_n]_i - \varepsilon[\mathbf{v}_n]_i \geq 1 - (1 - \varepsilon)[\mathbf{z}_n]_i - \varepsilon \\ &= (1 - \varepsilon)(1 - [\mathbf{z}_n]_i) \geq (1 - \varepsilon)^n(1 - [\mathbf{z}_1]_i), \end{aligned}$$

which completes the proof by induction.

Let $\mathbf{z}_\delta^* := \operatorname{argmax}_{\mathbf{z} \in \mathcal{K}_\delta} \tilde{F}(\mathbf{z})$. Using the same arguments as in Appendix K, we see that

$$\tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) \geq \varepsilon \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 L D^2}{2}.$$

Using Lemmas 2 and 1 and Equation (32), we have

$$\begin{aligned} \langle \nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n), \mathbf{z}_\delta^* - \mathbf{z}_n \rangle &\geq \tilde{F}(\mathbf{z}_\delta^* \vee \mathbf{z}_n) + \tilde{F}(\mathbf{z}_\delta^* \wedge \mathbf{z}_n) - 2\tilde{F}(\mathbf{z}_n) \\ &\geq (1 - \|\mathbf{z}_n\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) + \tilde{F}(\mathbf{z}_\delta^* \wedge \mathbf{z}_n) - 2\tilde{F}(\mathbf{z}_n) \\ &\geq (1 - \varepsilon)^{n-1}(1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) + \tilde{F}(\mathbf{z}_\delta^* \wedge \mathbf{z}_n) - 2\tilde{F}(\mathbf{z}_n) \\ &\geq (1 - \varepsilon)^{n-1}(1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) - 2\tilde{F}(\mathbf{z}_n). \end{aligned}$$

Therefore

$$\begin{aligned} \tilde{F}(\mathbf{z}_{n+1}) - \tilde{F}(\mathbf{z}_n) &\geq \varepsilon(1 - \varepsilon)^{n-1}(1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) - 2\varepsilon \tilde{F}(\mathbf{z}_n) \\ &\quad - \varepsilon D \|\nabla(\tilde{F}|_{\mathcal{L}})(\mathbf{z}_n) - \bar{\mathbf{g}}_n\| - \frac{\varepsilon^2 L D^2}{2}. \end{aligned}$$

After taking the expectation, using Lemma 9 and re-arranging the terms, we see that

$$\begin{aligned} \mathbb{E}[\tilde{F}(\mathbf{z}_{n+1})] &\geq (1 - 2\varepsilon) \mathbb{E}[\tilde{F}(\mathbf{z}_n)] + \varepsilon(1 - \varepsilon)^{n-1}(1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) \\ &\quad - \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{\varepsilon^2 L D^2}{2}. \end{aligned} \quad (33)$$

Using this inequality recursively together with Equation (22), we see that

$$\begin{aligned} \mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] &\geq \varepsilon(1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) \sum_{n=1}^N (1 - \varepsilon)^{n-1} (1 - 2\varepsilon)^{N-n} \\ &\quad + (1 - 2\varepsilon)^N \mathbb{E}[\tilde{F}(\mathbf{z}_1)] - \sum_{n=1}^N \frac{\varepsilon D Q^{1/2}}{(n+4)^{1/3}} - \frac{N \varepsilon^2 L D^2}{2}. \end{aligned} \quad (34)$$

Elementary calculations show that $(1 - \frac{\varepsilon}{N})^{N-1} \geq e^{-c}$ for $0 \leq c \leq 2$ and $N \geq 4$.¹ Therefore, since $\varepsilon = \log(2)/N$, we have

$$(1 - 2\varepsilon)^N \geq e^{-2 \log(2)} (1 - 2\varepsilon) = \frac{1}{4} \left(1 - \frac{2 \log(2)}{N}\right) \geq \frac{1}{4N}. \quad (35)$$

On the other hand

$$\begin{aligned} \varepsilon \sum_{n=1}^N (1 - 2\varepsilon)^{N-n} (1 - \varepsilon)^{n-1} &= \varepsilon (1 - 2\varepsilon)^{N-1} \sum_{n=1}^N \left(\frac{1 - \varepsilon}{1 - 2\varepsilon}\right)^{n-1} \\ &\geq \varepsilon (1 - 2\varepsilon)^{N-1} \sum_{n=1}^N (1 + \varepsilon)^{n-1} \\ &= (1 - 2\varepsilon)^{N-1} ((1 + \varepsilon)^N - 1). \end{aligned}$$

¹For $0 \leq x \leq \frac{1}{2}$, we have $\log(1 - x) \geq -x - \frac{x^2}{2} - x^3$. Therefore, for $0 \leq c \leq 2$ and $N \geq 4$, we have $\log(1 - \frac{c}{N}) \geq -\frac{c}{N} - \frac{c^2}{2N^2} - \frac{c^3}{N^3} \geq -\frac{c}{N-1}$.

We have $(1 + \frac{c}{N})^N \geq e^c(1 - \frac{c^2}{2N})$ for $c \geq 0$ and $N \geq 1$.² Therefore

$$\begin{aligned}
\varepsilon \sum_{n=1}^N (1 - 2\varepsilon)^{N-n} (1 - \varepsilon)^{n-1} &= (1 - 2\varepsilon)^{N-1} \left((1 + \varepsilon)^N - 1 \right) \\
&\geq e^{-2 \log(2)} \left(\left(1 + \frac{\log(2)}{N} \right)^N - 1 \right) \\
&\geq e^{-2 \log(2)} \left(e^{\log 2} \left(1 - \frac{\log(2)^2}{2N} \right) - 1 \right) \\
&= \frac{1}{4} \left(1 - \frac{\log(2)^2}{N} \right) \\
&\geq \frac{1}{4} - \frac{1}{4N}.
\end{aligned}$$

Plugging this and 35 into 34 and using the fact that $\tilde{F}(z_1)$ is non-negative, we get

$$\begin{aligned}
\mathbb{E}[\tilde{F}(\mathbf{z}_{N+1})] &\geq \left(\frac{1}{4} - \frac{1}{4N} \right) (1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) + \frac{1}{4N} \mathbb{E}[\tilde{F}(\mathbf{z}_1)] - \frac{3DQ^{1/2}}{N^{1/3}} - \frac{LD^2}{2N} \\
&\geq \frac{1}{4} (1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) + \frac{1}{4N} \left(\mathbb{E}[\tilde{F}(\mathbf{z}_1)] - \tilde{F}(\mathbf{z}_\delta^*) \right) - \frac{3DQ^{1/2}}{N^{1/3}} - \frac{LD^2}{2N} \\
&\geq \frac{1}{4} (1 - \|\mathbf{z}_1\|_\infty) \tilde{F}(\mathbf{z}_\delta^*) - \frac{3DQ^{1/2}}{N^{1/3}} - \frac{DG + 2LD^2}{4N}.
\end{aligned}$$

Using the same argument as in Appendix K, we obtain

$$\frac{1}{4} (1 - \|\mathbf{z}_1\|_\infty) F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})] \leq \frac{3DQ^{1/2}}{N^{1/3}} + \frac{DG + 2LD^2}{4N} + \delta G(2 + \frac{D}{r}). \quad \square$$

M Proof of Theorem 2

Proof. Let $T = O(BN)$ denote the number of evaluations³ and let $\mathcal{E}_\alpha := \alpha F(\mathbf{z}^*) - \mathbb{E}[F(\mathbf{z}_{N+1})]$ denote the α -approximation error. We prove Cases 1-4 separately. Note that F being non-monotone or $\mathbf{0} \in \mathcal{K}$ correspond to cases (A), (B) and (D) of Theorem 1 where $\log(N)$ does not appear in the approximation error bound, which is why $\tilde{O}$ can be replaced with O .

Case 1 (deterministic gradient oracle): In this case, we have $Q = \delta = 0$. According to Theorem 1, in cases (A), (B) and (D), the approximation error is bounded by $\frac{DG + 2LD^2}{4N} = O(N^{-1})$, and thus we choose $T = N = \Theta(1/\epsilon)$ to get $\mathcal{E}_\alpha = O(\epsilon)$. Similarly, in case (C), we have

$$\mathcal{E}_\alpha \leq \frac{4DG + LD^2 \log(N)^2}{8N} = O(N^{-1} \log(N)^2).$$

We choose $T = N = \Theta(\log^2(\epsilon)/\epsilon)$ to bound α -approximation error by $O(\epsilon)$.

Case 2 (stochastic gradient oracle): In this case, we have $Q = \Theta(1)$ and $\delta = 0$. According to Theorem 1, in cases (A), (B) and (D), the approximation error is bounded by

$$\frac{3DQ^{1/2}}{N^{1/3}} + \frac{DG + 2LD^2}{4N} = O(N^{-1/3} + N^{-1}) = O(N^{-1/3}),$$

so we choose $N = \Theta(1/\epsilon^3)$, $B = 1$ and $T = \Theta(1/\epsilon^3)$ to get $\mathcal{E}_\alpha = O(\epsilon)$. Similarly, in case (C), we have

$$\mathcal{E}_\alpha \leq \frac{3DQ^{1/2} \log(N)}{2N^{1/3}} + \frac{4DG + LD^2 \log(N)^2}{8N} = O(N^{-1/3} \log(N) + N^{-1} \log(N)^2)$$

²For $x \geq 0$, we have $\log(1 + x) \geq x - \frac{x^2}{2}$ and $-x \geq \log(1 - x)$. Therefore $N \log(1 + \frac{c}{N}) \geq N(\frac{c}{N} - \frac{c^2}{2N^2}) = c - \frac{c^2}{2N} \geq c + \log(1 - \frac{c^2}{2N})$.

³We have $T = BN$ when we have access to a gradient oracle and $T = 2BN$ otherwise.

Since $\mathcal{E}_\alpha \leq O(N^{-1/3} \log(N)^2)$, we choose $N = \Theta(\log^6(\epsilon)/\epsilon^3)$, $B = 1$ and $T = \Theta(\log^6(\epsilon)/\epsilon^3)$ to bound α -approximation error by $O(\epsilon)$.

Case 3 (deterministic value oracle): In this case, we have $Q = \Theta(1)$ and $\delta \neq 0$. According to Theorem 1, in cases (A), (B) and (D), the approximation error is bounded by

$$\frac{3DQ^{1/2}}{N^{1/3}} + \frac{DG + 2LD^2}{4N} + O(\delta) = O(N^{-1/3} + \delta),$$

so we choose $\delta = \Theta(\epsilon)$, $N = \Theta(1/\epsilon^3)$, $B = 1$ and $T = \Theta(1/\epsilon^3)$ to get $\mathcal{E}_\alpha = O(\epsilon)$. Similarly, in case (C), we have

$$\mathcal{E}_\alpha \leq \frac{3DQ^{1/2} \log(N)}{2N^{1/3}} + \frac{4DG + LD^2 \log(N)^2}{8N} + O(\delta) = O(N^{-1/3} \log(N)^2 + \delta).$$

We choose $\delta = \Theta(\epsilon)$, $N = \Theta(\log^6(\epsilon)/\epsilon^3)$, $B = 1$ and $T = \Theta(\log^6(\epsilon)/\epsilon^3)$ to bound α -approximation error by $O(\epsilon)$.

Case 4 (stochastic value oracle): In this case, we have $Q = O(1) + O(\frac{1}{\delta^2 B})$ and $\delta \neq 0$. According to Theorem 1, in cases (A), (B) and (D), the approximation error is bounded by

$$\begin{aligned} \frac{3DQ^{1/2}}{N^{1/3}} + \frac{DG + 2LD^2}{4N} + O(\delta) &= O(Q^{1/2} N^{-1/3} + N^{-1} + \delta) \\ &= O(N^{-1/3} + \delta^{-1} B^{-1/2} N^{-1/3} + \delta), \end{aligned}$$

so we choose $\delta = \Theta(\epsilon)$, $N = \Theta(1/\epsilon^3)$, $B = \Theta(1/\epsilon^2)$ and $T = \Theta(1/\epsilon^5)$ to get $\mathcal{E}_\alpha = O(\epsilon)$. Similarly, in case (C), we have

$$\begin{aligned} \mathcal{E}_\alpha &\leq \frac{3DQ^{1/2} \log(N)}{2N^{1/3}} + \frac{4DG + LD^2 \log(N)^2}{8N} + O(\delta) \\ &= O(Q^{1/2} N^{-1/3} \log(N) + N^{-1} \log(N)^2 + \delta) \\ &= O(N^{-1/3} \log(N)^2 + \delta^{-1} B^{-1/2} N^{-1/3} \log(N)^2 + \delta). \end{aligned}$$

We choose $\delta = \Theta(\epsilon)$, $N = \Theta(\log^6(\epsilon)/\epsilon^3)$, $B = \Theta(1/\epsilon^2)$, and $T = \Theta(\log^6(\epsilon)/\epsilon^5)$ to bound α -approximation error by $O(\epsilon)$. $\square$

N Proof of Theorem 3

Proof. Since the parameters of Algorithm 2 are chosen according to Theorem 2, we see that the α -approximation error is bounded by $\tilde{O}(T_0^{-\beta})$ where $\beta = 1/3$ in case 2 (stochastic gradient oracle) and $\beta = 1/5$ in case 4 (stochastic value oracle).

Recall that F is G -Lipschitz and the feasible region $\mathcal{K}$ has diameter D . Thus, during the first T_0 time-steps, the α -regret can be bounded by

$$\sup_{z, z' \in \mathcal{K}} \alpha F(\mathbf{z}) - F(\mathbf{z}') \leq \sup_{z, z' \in \mathcal{K}} F(\mathbf{z}) - F(\mathbf{z}') \leq DG.$$

Therefore the total α -regret is bounded by

$$T_0 DG + (T - T_0) \tilde{O}(T_0^{-\beta}) \leq T_0 DG + T \tilde{O}(T_0^{-\beta}).$$

Since we have $T_0 = \Theta(T^{\frac{1}{\beta+1}})$, we see that

$$T_0 DG + T \tilde{O}(T_0^{-\beta}) = \tilde{O}(T^{\frac{1}{\beta+1}}) = \begin{cases} \tilde{O}(T^{\frac{3}{4}}) & \text{Case 2,} \\ \tilde{O}(T^{\frac{5}{6}}) & \text{Case 4.} \end{cases}$$

If F is non-monotone or $\mathbf{0} \in \mathcal{K}$, the exact same argument applies with $\tilde{O}$ replaced by O . $\square$