

Low-Discrepancy Perturbation Design for Numerically Stable Gradient Estimation in Feedback Optimization

Carmen Jimenez Cortes^a, Isaac Monteiro^b Samuel Coogan^c, *Senior Member, IEEE*,
Marco Gamarra^d, *Member, IEEE*, and Makhin Thitsa^b

Abstract—We present a numerically stable gradient estimation algorithm for unknown objective functions in feedback-based optimization settings. Unlike classical extremum seeking control, which recovers gradient components sequentially via distinct demodulation frequencies, our approach estimates all components simultaneously, making it well-suited for high-dimensional and multi-agent systems. The core innovation lies in the use of carefully designed perturbations based on low-discrepancy sequences, which ensure structured and efficient probing. We provide theoretical guarantees of numerical stability, and numerical simulations validate the method's effectiveness in an integrated sensing and communication application.

Index Terms—Extremum Seeking, Optimization algorithms, Numerical Algorithms

I. INTRODUCTION

FEEDBACK optimization control is a strategy that embeds optimization directly within the feedback control loop, enabling a system to continuously move toward an optimal operating point [1]–[4]. The core objective is to adjust control inputs in real time based on system feedback, with the aim of optimizing a cost function subject to system constraints. This method is particularly useful in dynamic systems where the underlying relationship between control inputs and outputs may be unknown or difficult to model explicitly. Feedback optimization control is experiencing renewed interest because

modern systems must optimize while operating, under uncertainty, time variation, and limited modeling accuracy, conditions under which feedback-based approaches offer a clear advantage. Recent advances span online and real-time optimization frameworks, with increasing emphasis on robustness to uncertainty and model mismatch [5]–[8].

Gradient-free optimization methods such as covariance matrix adaptation evolution strategies (CMAES) [9], genetic algorithms, and neural-network-based approaches offer flexibility for nonconvex and poorly modeled problems, but they suffer from several important drawbacks. These methods typically require a large number of function evaluations, making them unsuitable for expensive simulations or real-time physical systems. Their performance also degrades rapidly with increasing problem dimension due to growing population sizes, covariance computations, and exploration requirements. Moreover, these approaches often lack strong theoretical guarantees on convergence and numerical stability. Classical extremum seeking control (ESC) [1], which relies on sequential, frequency-separated demodulation to recover gradient components, suffers from frequency interference and scales poorly with increasing dimensionality.

The work in [10] investigates stochastic perturbations and proves convergence to the extremum through a method based on averaging theory. Our approach aligns closely with the frameworks in [11], [12], which employ multi-perturbation excitation over a finite horizon or leverage historical data to recover the full gradient via regression in a single update. These regression-based schemes are increasingly viable due to high-speed sensing, which enables high-rate output sampling. However, convergence to the optimizer critically depends on the accuracy of the resulting gradient estimate. In particular, [11] derives bounded estimation error for a batch least-squares estimator using adaptive step size, while [12] establishes closed-loop stability via a Lyapunov–Razumikhin argument in a dither-free least-squares framework. In this work, we focus on reducing estimation error induced by ill-conditioned covariance matrices, with emphasis on excitation design that improves conditioning and thereby tightens gradient estimation error bounds. As such, we present results on static maps as opposed to the dynamic maps investigated in [1], [3], [5], [8].

To this end, we introduce a numerically stable algorithm for gradient estimation of unknown objective functions in feed-

*This research was supported in part by the Office of Naval Research under award N00014-21-1-2759. Makhin Thitsa was supported by the Air Force Research Laboratory, Information Directorate through the SFFP program.

¹Carmen (Thiemann) Jimenez Cortes is with the Department of Electrical and Computer Engineering, Kennesaw State University, Marietta, GA 30060, USA cthieman@kennesaw.edu

²Isaac Monteiro and Makhin Thitsa are with the School of Electrical and Computer Engineering, Mercer University, Macon, GA 31207, USA Isaac.Alexander.Monteiro@live.mercer.edu, thitsam@mercer.edu

³Samuel Coogan is with the School of Electrical and Computer Engineering, Georgia Institute of Technology, Atlanta, GA 30332, USA sam.coogan@gatech.edu

⁴Marco Gamarra is with the Air Force Research Laboratory, Rome, NY, USA marco.gamarra@us.af.mil

The views and opinions expressed by the authors herein do not necessarily reflect the official policies or positions of the United States Government or its agencies.

Distribution A. Approved for public release: distribution unlimited. Case Number: AFRL-2026-1350. Dated 19 Mar 2026.

back optimization. The method leverages carefully designed perturbations based on low-discrepancy sequences to ensure well-conditioned regression. The key contributions are:

- **Numerical stability and efficient probing:** The low-discrepancy perturbation design provides theoretical guarantees on a well-conditioned gradient estimation while enabling structured and efficient probing.
- **Application validation:** The proposed method is particularly well-suited for high-dimensional, multi-agent systems. More specifically, its application is shown in an Integrated Sensing and Communication (ISAC) framework, key to 6G wireless networks.

Section II reviews the necessary background material. Section III presents the mathematical analysis and the main theorem. In Section IV, the efficiency of the proposed method is demonstrated through numerical experiments in an integrated sensing and communication application. Finally, Section V concludes the paper.

II. BACKGROUND MATERIAL AND MATHEMATICAL PRELIMINARIES

A. Notation

In this letter, we let $\nabla V(x)$ represent the gradient operator applied to $V(x)$, while the subscript notation $\nabla_x V(x)$ specifically refers to the vector of partial derivatives with respect to the multi-dimensional state x . Lowercase boldface letters, such as $\mathbf{y}$, represent vectors, whereas italics denote scalars. For a matrix A , $[A]_{ij}$ refers to the ij th element, $\lambda_{\max}(A)$, $\lambda_{\min}(A)$, $\kappa(A)$, $\|A\|_2$ are maximum and minimum eigenvalues, the condition number and the spectral norm, respectively. $|\cdot|$ denotes the absolute value, and $\|\cdot\|$ denotes the 2-norm. The indicator function is noted as $\mathbf{1}(\cdot)$.

B. Mathematical Preliminaries

We consider a general system model represented as $\mathbf{y} = h(\mathbf{u})$, where $\mathbf{u} \in \mathbb{R}^n$ denotes the control input vector, $\mathbf{y} \in \mathbb{R}^p$ is the measured system output, and $h : \mathbb{R}^n \rightarrow \mathbb{R}^p$ is an unknown and possibly nonlinear mapping from inputs to outputs. The objective is to minimize a cost function defined as $J(\mathbf{u}) = f(\mathbf{y}) = f(h(\mathbf{u}))$, subject to constraints on $\mathbf{u}$ and (or) $\mathbf{y}$. This framework enables feedback-based adjustment of inputs to achieve system-level objectives in the absence of an explicit system model. Classical extremum seeking control (ESC) struggles in high-dimensional systems because it requires injecting distinct dither signals into each input to estimate gradient components. In an n -dimensional system, this entails n separate excitation frequencies, increasing bandwidth demands and the risk of frequency interference.

Demodulation-free ESC, as formulated in [11], [12] provide a more efficient gradient estimation framework. Instead of injecting periodic signals, this approach applies m perturbations at each iteration or collects m historical data points. These data points are then used to estimate the gradient of the cost function more efficiently. Let $\delta_i \in \mathbb{R}^n$ denote the i th perturbation vector. Assuming small perturbations and differentiability

of the objective function f , we can approximate the function value at the perturbed input as:

$$f(\mathbf{u} + \delta_i) = f(\mathbf{u}) + \nabla f(\mathbf{u}) \cdot \delta_i + r_i,$$

where r_i is the remainder error of the Taylor expansion. Defining $y_i = f(\mathbf{u} + \delta_i) - f(\mathbf{u})$ yields, $y_i = \delta_i^T \nabla f(\mathbf{u}) + r_i$. Stacking m such measurements yields the linear system of equations:

$$\mathbf{y} = D_m \nabla^T f(\mathbf{u}) + \mathbf{r}, \quad (1)$$

where $D_m \in \mathbb{R}^{m \times n}$ is the perturbation matrix with rows δ_i^T and $\mathbf{r} = [r_1, \dots, r_m]^T$. The gradient can then be estimated by solving a least-squares problem:

$$\hat{\mathbf{g}} = \arg \min_{\mathbf{g} \in \mathbb{R}^n} \|\mathbf{y} - D_m \mathbf{g}\|,$$

which has the closed-form solution: $\hat{\mathbf{g}} = (D_m^T D_m)^{-1} D_m^T \mathbf{y}$. Substituting (1) yields $\hat{\mathbf{g}} - \nabla^T f(\mathbf{u}) = (D_m^T D_m)^{-1} D_m^T \mathbf{r}$. If $\|\mathbf{r}\| \leq R$, then,

$$\|\hat{\mathbf{g}} - \nabla^T f(\mathbf{u})\| \leq \sqrt{\frac{\kappa(D_m^T D_m)}{\lambda_{\max}(D_m^T D_m)}} R. \quad (2)$$

[11] proposed designing an adaptive step size to reduce R , and hence the error bound. This letter focuses on tightening this error bound by ensuring that the matrix $D_m^T D_m$ is well-conditioned by carefully choosing the perturbation vectors.

III. METHODOLOGY

We propose using low-discrepancy sequences, specifically, Sobol sequences, to construct the perturbation directions. These sequences provide quasi-uniform coverage of the input space, avoiding clustering and ensuring even distribution. Perturbations are generated by evaluating orthonormal basis functions at Sobol sequence points [13], and scaling is applied independently in each dimension to control perturbation magnitude.

Let $\mathcal{P}_N = \{x_1, x_2, \dots, x_N\} \subset [0, 1]^d$ be a set of N points. The **discrepancy** of $\mathcal{P}_N$ with respect to a family $\mathcal{A}$ of subsets of $[0, 1]^d$ is defined as:

$$D_N(\mathcal{P}_N, \mathcal{A}) = \sup_{A \in \mathcal{A}} \left| \frac{1}{N} \sum_{n=1}^N \mathbf{1}_A(x_n) - \lambda_d(A) \right|,$$

where $\mathbf{1}_A(x)$ is the indicator function of A , and $\lambda_d(A)$ is the Lebesgue measure of A .

The **star discrepancy**, a specific form of discrepancy commonly used in numerical integration, is defined by:

$$D_N^*(\mathcal{P}_N) = \sup_{0 \leq \mathbf{u} \leq \mathbf{1}} \left| \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{[0, \mathbf{u}]}(x_n) - \prod_{i=1}^d u_i \right|.$$

Low star discrepancy ensures better convergence properties when using these points for function approximation.

Sobol sequences are low-discrepancy, quasi-random sequences that fill the input space more uniformly than purely random samples. Their main advantage lies in their star discrepancy: for d dimensional Sobol sequences, the star discrepancy is $D_N^*(\text{Sobol}) = \mathcal{O}((\log N)^d/N)$. When $d = 1$, the Sobol sequence reduces to Van der Corput sequence [14], whose convergence rate is asymptotically better than that of random samples, which have $D_N^*(\text{random}) = \mathcal{O}(1/\sqrt{N})$.

A. Construction of the Perturbation Matrix

To estimate gradients efficiently in high-dimensional spaces, it is essential to construct a well-conditioned perturbation matrix D_m . This matrix is formed using a set of smooth orthonormal functions and scaling constants that control the amplitude of perturbations in each dimension.

Let $\phi_j : \mathbb{R} \rightarrow \mathbb{R}$ for $1 \leq j \leq n$ be a set of smooth orthonormal functions defined on the interval $[0, 1] \subset \mathbb{R}$, and let $k_j > 0$ be dimension-dependent scaling factors. For a set of m sample points $\{x_i\}_{i=1}^m$ drawn from a one dimensional Sobol sequence, the perturbation matrix $D_m \in \mathbb{R}^{m \times n}$ is defined element-wise as: $[D_m]_{ij} = k_j \phi_j(x_i)$. Each column of D_m corresponds to one basis function scaled by a k_j , allowing control over the perturbation magnitude per input dimension. The full matrix takes the form:

$$D_m = \begin{bmatrix} k_1 \phi_1(x_1) & k_2 \phi_2(x_1) & \cdots & k_n \phi_n(x_1) \\ k_1 \phi_1(x_2) & k_2 \phi_2(x_2) & \cdots & k_n \phi_n(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ k_1 \phi_1(x_m) & k_2 \phi_2(x_m) & \cdots & k_n \phi_n(x_m) \end{bmatrix}.$$

B. Properties of $D_m^T D_m$

The condition number of $D_m^T D_m$ governs the stability of the least-squares gradient estimate. The matrix product $D_m^T D_m$ has entries:

$$[D_m^T D_m]_{pq} = \sum_{k=1}^m k_p k_q \phi_p(x_k) \phi_q(x_k), \quad \forall p, q \in \{1, \dots, n\}. \quad (3)$$

Since the functions $\{\phi_j\}$ are orthonormal on $[0, 1]$, the continuous inner product satisfies:

$$\int_0^1 k_p k_q \phi_p(x) \phi_q(x) dx = \begin{cases} k_p^2, & \text{if } p = q, \\ 0, & \text{if } p \neq q. \end{cases} \quad (4)$$

It is also worth noting that if we define matrix $G \in \mathbb{R}^{n \times n}$ as:

$$[G]_{pq} = \int_0^1 k_p k_q \phi_p(x) \phi_q(x) dx, \quad (5)$$

the condition number of G is:

$$\kappa(G) = \|G\|_2 \cdot \|G^{-1}\|_2 = \frac{k_{\max}^2}{k_{\min}^2}, \quad (6)$$

where $k_{\max} = \max_j k_j$ and $k_{\min} = \min_j k_j$.

C. Main Theorem on Matrix Conditioning

Our main theorem in this section provides an upper bound on the condition number of the matrix $D_m^T D_m$ based on the properties of the basis functions and sampling scheme. The following lemmas: Koksma-Hlawka and Weyl's inequalities are key to the proof of the main theorem.

Lemma 1: (Koksma-Hlawka [15]) Let $f : [0, 1]^d \rightarrow \mathbb{R}$ be of bounded variation, and let $(x_n)_{n=1}^N \subset [0, 1]^d$. Then the following inequality holds.

$$\left| \frac{1}{N} \sum_{n=1}^N f(x_n) - \int_{[0,1]^d} f(x) dx \right| \leq V_{\text{HK}}(f) \cdot D_N^*(x_1, \dots, x_N),$$

where $V_{\text{HK}}(f)$ is the variation of f in the sense of Hardy and Krause, and $D_N^*(x_1, \dots, x_N)$ is the star-discrepancy of the point set $\{x_1, \dots, x_N\}$.

Lemma 2: (Weyl's Inequality [16]) Let A, B be Hermitian matrices with dimension n on an inner product space with eigenvalues arranged in descending order, $\lambda_1 \geq \dots \geq \lambda_n$, then, $|\lambda_k(A+B) - \lambda_k(A)| \leq \|B\|_2$, $\forall 1 \leq k \leq n$.

Theorem 1: Let $\{\phi_j : \mathbb{R} \rightarrow \mathbb{R}\}_{j=1}^n$ be a set of smooth, orthonormal functions on $[0, 1] \subset \mathbb{R}$. Let $k_j > 0$ for all $1 \leq j \leq n$, and suppose $\{x_i\}_{i=1}^m \subset [0, 1]$ is a set of one dimensional Sobol sequence sample points. Define the matrix $D_m \in \mathbb{R}^{m \times n}$ by $[D_m]_{ij} = k_j \phi_j(x_i)$, $\forall 1 \leq i \leq m, 1 \leq j \leq n$. Then the condition number of $D_m^T D_m$ satisfies:

$$\kappa(D_m^T D_m) = \frac{k_{\max}^2}{k_{\min}^2} + \mathcal{O}\left(\frac{\log m}{m}\right),$$

where $k_{\max} = \max_{1 \leq j \leq n} k_j$, $k_{\min} = \min_{1 \leq j \leq n} k_j$.

Proof: We invoke the Koksma-Hlawka inequality, which bounds the integration error of a function f of bounded variation $V_{\text{HK}}(f)$ using the star discrepancy D_m^* of the point set:

$$\left| \frac{1}{m} \sum_{k=1}^m f(x_k) - \int_0^1 f(x) dx \right| \leq V_{\text{HK}}(f) \cdot D_m^*.$$

Let $f_{pq}(x) = k_p k_q \phi_p(x) \phi_q(x)$. Then,

$$\left| \frac{1}{m} \sum_{k=1}^m f_{pq}(x_k) - \int_0^1 f_{pq}(x) dx \right| \leq V_{\text{HK}}(f_{pq}) \cdot D_m^*.$$

Hence, from Equations (3),(4) and (5), we obtain the error bound: $(D_m^T D_m - mG)_{pq} = \mathcal{O}(\log m)$. Note that $D_m^T D_m$ and G respectively defined in (3) and (5) are both Hermitian and thus satisfy the following form of Weyl's inequality. Let $E = D_m^T D_m - mG$. Then,

$$|\lambda_{\max}(D_m^T D_m) - m\lambda_{\max}(G)| \leq \|E\|_2, \quad (7)$$

$$|\lambda_{\min}(D_m^T D_m) - m\lambda_{\min}(G)| \leq \|E\|_2. \quad (8)$$

Thus, we have:

$$\lambda_{\max}(D_m^T D_m) \leq m\lambda_{\max}(G) + \|E\|_2,$$

$$\lambda_{\min}(D_m^T D_m) \geq m\lambda_{\min}(G) - \|E\|_2.$$

The condition number then satisfies:

$$\begin{aligned} \kappa(D_m^T D_m) &= \frac{\lambda_{\max}(D_m^T D_m)}{\lambda_{\min}(D_m^T D_m)} \leq \frac{\lambda_{\max}(G) + \frac{\|E\|_2}{m}}{\lambda_{\min}(G) - \frac{\|E\|_2}{m}} \\ &= \kappa(G) + \mathcal{O}\left(\frac{\log m}{m}\right). \end{aligned}$$

Thus, by (6),

$$\kappa(D_m^T D_m) = \frac{k_{\max}^2}{k_{\min}^2} + \mathcal{O}\left(\frac{\log m}{m}\right).$$

Remark 1: Since $V_{\text{HK}}(f_{pq}) = \mathcal{O}(pq)$ and $p, q \leq n$, $|(E)_{pq}| = \mathcal{O}(n^2 \log m)$. Passing from an entry-wise bound to an operator norm bound introduces an additional factor of n , yielding the worst-case estimate for the condition number, $\kappa(D_m^T D_m) = \mathcal{O}\left(n^3 \frac{\log m}{m}\right)$. However, this estimate is highly

conservative, as it assumes fully coherent accumulation across all matrix entries. In practice, the oscillatory structure of the polynomial products leads to substantial cancellation, and the effective growth rate of V_{KH} is considerably smaller than the worst-case $O(pq)$ behavior. Consequently, the observed growth rate is closer to $\kappa(D_m^T D_m) = O\left(n^2 \frac{\log m}{m}\right)$. This scaling is consistent with the behavior observed in our numerical experiments as featured in Figure 1.

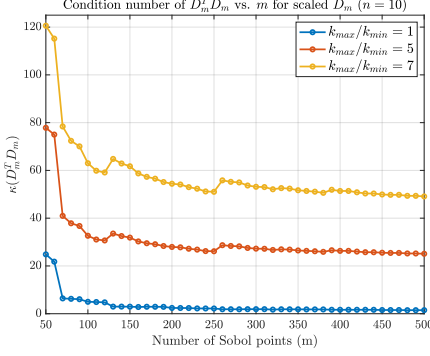


Fig. 1: Condition number test for a 10 dimensional problem. $\kappa(D_m^T D_m)$ converges to k_{max}^2/k_{min}^2 .

In systems where the nominal parameter magnitudes are comparable across dimensions, the ratio $k_{max}/k_{min} \approx 1$, indicating a well-conditioned problem. When the nominal values vary significantly across dimensions, this ratio can become large, resulting in a higher condition number, but the gradient can still be accurately approximated provided that the condition number remains a reasonably small finite value ($< 10^3$). In practice, the number of perturbations is typically chosen to exceed the system dimension, i.e., $m > n$, to introduce redundancy and reduce the risk of noninvertibility of $D_m^T D_m$. Our theorem shows that, with the proposed structured perturbation design, $D_m^T D_m$ converges to a well-conditioned matrix at a significantly faster rate than when random perturbations are used.

Once the gradient is estimated sufficiently accurately, the inexact gradient descent iteration $\mathbf{x}_{k+1} = \mathbf{x}_k - \alpha(\nabla^T f(\mathbf{x}_k) + \mathbf{e}_k)$, where α is the step size and $\mathbf{e}_k = \hat{\mathbf{g}}_k - \nabla^T f(\mathbf{x}_k)$ denotes the gradient estimation error at iteration k , is known to be input-to-state stable (ISS) under suitable regularity conditions on the objective function and sufficiently small α . In particular, sufficiently small gradient estimation errors ensure that the iterates converge to, and ultimately remain within, a small neighborhood of the optimum [17]. Figure 2 and 3 compare the scalability of the classical ESC and our approach for the minimization of a simple function $f(\mathbf{u}) = \sum_{i=1}^n (u_i - 1)^2$, with the constraint: $-5 \leq u_i \leq 5, \forall i \in \{1, \dots, n\}$. For the same problem, Figure 4 compares Low-Discrepancy ESC with regression-based ESC using random perturbations over 100 runs with $n = 20, m = 50$. Our proposed method consistently reduces costs more effectively, outperforming the best regression-based ESC run at every estimation interval.

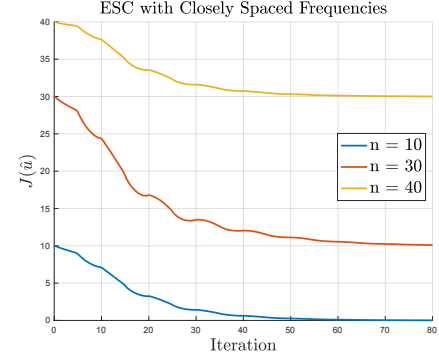


Fig. 2: Objective function over time using traditional ESC. The perturbation amplitude is $a = 0.2$, the base frequency is $\omega_0 = 0.2\pi$, the frequency spacing is $\Delta\omega = 0.4\pi$, and the filter time constant is $\tau = 1.0$. As the number of dimensions n increases, traditional methods fail to minimize $J(\hat{\mathbf{u}})$

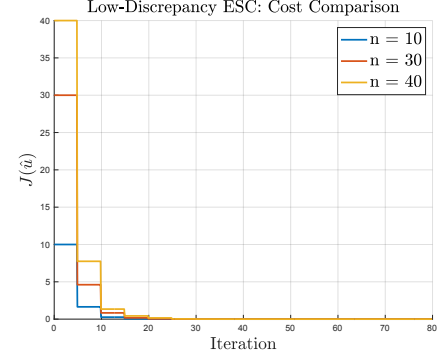


Fig. 3: Objective function over time using Low-Discrepancy ESC showing convergence for dimensions $n = 10, 30, 40$

IV. ISAC: INTEGRATED SENSING AND COMMUNICATION

Our proposed method is particularly well suited for high-dimensional and multiagent problems, with a direct application being Integrated Sensing and Communication (ISAC). ISAC is a framework in which the same signals, spectrum, hardware, and infrastructure are used simultaneously for both wireless communication and environmental sensing (radar-like functions). One of its key applications is 6G wireless networks, as it enables devices to communicate while simultaneously mapping and sensing their surroundings.

We model the ISAC framework as a multiagent problem, where agents must optimize both a desired formation as well as maximize the quality of their communication. We model the multi-agent system as a directed graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, where $\mathcal{V} = \{1, \dots, N\}$ is the set of all agents and $\mathcal{E} \subseteq \mathcal{V} \times \mathcal{V}$ is the set of active communication and sensing links. An edge $(i, j) \in \mathcal{E}$ exists if agent i can interact with agent j .

A. Formation Problem

In general, the goal is for each edge $(i, j) \in \mathcal{E}$, agents i and j need to hold a relative position $\mathbf{x}_i - \mathbf{x}_j \approx \Omega_{ij}$. This is achieved using the Lyapunov function $V(\mathbf{x}) = \frac{1}{2} \sum_{(i,j) \in \mathcal{E}} \|\mathbf{x}_i - \mathbf{x}_j - \Omega_{ij}\|^2$, where both $\mathbf{x}_i$ and $\mathbf{x}_j$ can be

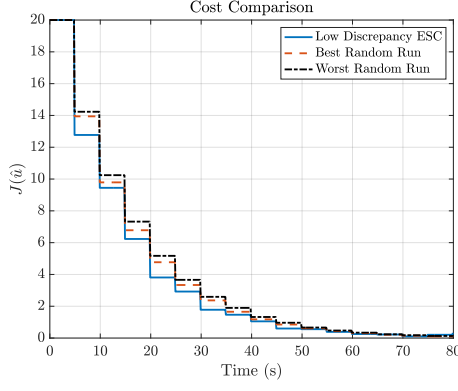


Fig. 4: Comparison of Low Discrepancy ESC and the best-case (lowest total cost) and worst-case (highest total cost) trajectories from 100 experiments of regression-based ESC using random perturbations

measured. $V(\mathbf{x})$ is known, and therefore its gradient can be obtained analytically.

B. Communication Problem

The communication among the agents in the network is modeled via the quality function $Q(\mathbf{x})$, which expresses how good the overall sensing is. In general, $Q(\mathbf{x})$ is not known, but can be measured. We propose the following $Q(\mathbf{x})$:

$$Q(\mathbf{x}) = \sum_{(i,j) \in \mathcal{E}} \frac{1}{\kappa} \log \left(1 + \exp \left[e^{-\beta \|\mathbf{x}_i - \mathbf{x}_j\|^2} - \sum_{k \in \mathcal{N}_i \setminus \{j\}} e^{-\beta \|\mathbf{x}_i - \mathbf{x}_k\|^2} \cos^2(\tilde{\theta}_{ijk}) g(\tilde{\theta}_{ijk}) \right] \right) \quad (9)$$

$$\tilde{\theta}_{ijk} = \arccos \left(\frac{\tanh \left(\alpha \frac{(\mathbf{x}_j - \mathbf{x}_i)(\mathbf{x}_k - \mathbf{x}_i)}{\|\mathbf{x}_j - \mathbf{x}_i\| \|\mathbf{x}_k - \mathbf{x}_i\| + \varepsilon} \right)}{\tanh(\alpha)} \right) \quad (10)$$

$$g(\tilde{\theta}_{ijk}) = \frac{1}{2} \left[1 + \tanh \left(s \left(\frac{\pi}{2} - \tilde{\theta}_{ijk} \right) \right) \right]. \quad (11)$$

$Q(\mathbf{x})$ can be interpreted as the addition of two different penalties. First, a *distance penalty*, which penalizes agents trying to communicate for being too far apart, and second, an *interference penalty*, which accounts for external agents k blocking the communication line of sight between two other agents i and j . In equations (9) and (10), $\beta > 0$ controls the decay of the distance-based interaction term, $\alpha > 0$ the smoothness of the angular clipping, $s > 0$ the steepness of the gating transition near $\pi/2$, $\kappa > 0$ the softness of the softplus activation, and $\varepsilon > 0$ is a small numerical constant to prevent division by zero.

C. Updated Lyapunov function

To simultaneously optimize the formation of the agents and their communication quality in the network, we combine functions $V(\mathbf{x})$ and $Q(\mathbf{x})$ in the new Lyapunov function:

$$\tilde{V}(\mathbf{x}) = V(\mathbf{x}) + \lambda Q(\mathbf{x})^{-1}. \quad (12)$$

Minimizing $\tilde{V}(\mathbf{x})$, encourages a larger $Q(\mathbf{x})$ and smaller $V(\mathbf{x})$. If the communication quality deteriorates and $Q \rightarrow 0^+$ then $\tilde{V}$ would cause the agents to push far away from their current states. Note that

$$\nabla_{\mathbf{x}} \tilde{V}(\mathbf{x}) = \nabla_{\mathbf{x}} V(\mathbf{x}) - \lambda Q(\mathbf{x})^{-2} \nabla_{\mathbf{x}} Q(\mathbf{x}). \quad (13)$$

Applying the gradient descent law $\dot{\mathbf{x}} = -\nabla_{\mathbf{x}} \tilde{V}(\mathbf{x})$, we obtain the stacked system dynamics:

$$\dot{\mathbf{x}} = -\nabla_{\mathbf{x}} V(\mathbf{x}) + \lambda Q(\mathbf{x})^{-2} \nabla_{\mathbf{x}} Q(\mathbf{x}). \quad (14)$$

To solve this dynamics, we compute $\nabla_{\mathbf{x}} V(\mathbf{x})$ analytically, we measure $\lambda Q(\mathbf{x})$, and we estimate $\nabla_{\mathbf{x}} Q(\mathbf{x})$ by proposed low-discrepancy ESC. $\tilde{V}(\mathbf{x})$ is robust to small values in $Q(\mathbf{x})^{-1}$ as $V(\mathbf{x})$ enforces formation constraints that prevent the edge cases where $Q(\mathbf{x}) \rightarrow 0^+$, specifically, divergence and overlap of neighboring agents.

D. Numerical Results

We simulate a network of six agents, each modeled with single integrator dynamics $[\dot{x}_1, \dot{x}_2]^T = [u_1, u_2]$, yielding a twelve-dimensional problem. The shifted Legendre polynomials orthonormal on $[0, 1]$ are selected as the basis functions, ϕ_j ¹. This choice was motivated by implementation convenience rather than a comparative performance analysis of various polynomial types. We run two different simulations. In both cases, we use $\lambda = 10$, a perturbation amplitude of 0.05, $\alpha_{ij}(\cdot) = 1$, and $\beta = 1, \alpha = 5 \cdot 10^{-4}, s = 10, \kappa = 1, \epsilon = 10^{-8}$.

In both simulations, the agents start evenly distributed around a circle, and the chosen desired formation is a 2 by 3 grid. Note that these simulations maintain a constant number of total gradient updates, which we denote as “batches”. The circular formation is ideal for communication, as it reduces the interferences among agents to a minimum. Once the agents start displacing, the communication deteriorates, causing the “drop” shown in Figure 5. At their final positions, both the formation and quality are such that $\tilde{V}(\mathbf{x})$ is minimized. The values of $\tilde{V}(\mathbf{x})$, $V(\mathbf{x})$ and $Q(\mathbf{x})$ over the simulation horizon are included in Figure 5. Starting with $\tilde{V}(\mathbf{x}) = 87.481$, $V(\mathbf{x}) = 87$, and $Q(\mathbf{x}) = 20.813$, after 500 batches, we yield the final values $\tilde{V}(\mathbf{x}) = 0.491$, $V(\mathbf{x}) = 3.81 \cdot 10^{-4}$, and $Q(\mathbf{x}) = 20.384$. With 60 perturbations, and the same initial values of $\tilde{V}(\mathbf{x})$, $V(\mathbf{x})$, and $Q(\mathbf{x})$ we yield the final values $\tilde{V}(\mathbf{x}) = 0.491$, $V(\mathbf{x}) = 3.79 \cdot 10^{-4}$, and $Q(\mathbf{x}) = 20.384$, which are almost identical to those obtained with only 24 perturbations per estimating, highlighting that a substantial reduction in number of perturbation does not degrade the performance significantly.

Lastly, we also include a comparison of the individual gradient components over the entire simulation horizon in Figure 6 for the case with only 24 perturbations, and Figure 7 shows the individual gradients for 60 perturbations. The solid lines represent the our proposed Low-Discrepancy ESC technique, while the dashed lines denote the finite difference benchmarks for each of the $2N$ dimensions. Despite using

¹Normalized Legendre polynomials are orthonormal on $[-1, 1]$, whereas the Sobol sequence is defined on $[0, 1]$. Therefore, the normalized shifted Legendre polynomials $P_n(2x - 1)$ are used.

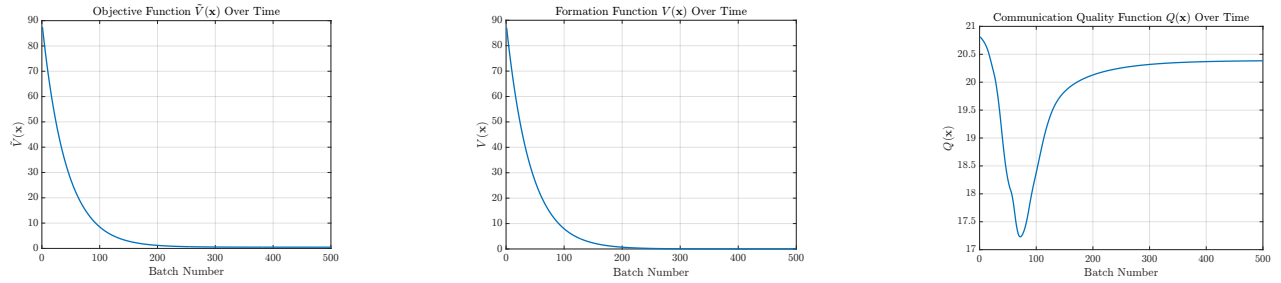


Fig. 5: Objective function $\hat{V}(\mathbf{x})$ and its components over 500 batches, with 24 perturbations per gradient estimate

only 24 perturbations in a twelve-dimensional setting, the gradient components are estimated reasonably well. When the number of perturbations is increased to 60, the estimates track respective gradient component with excellent accuracy.

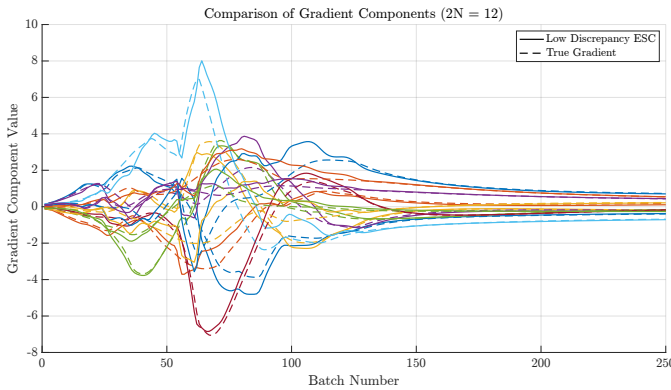


Fig. 6: Comparison of estimated gradient components and true gradient components over batch iterations using 24 perturbations in twelve-dimensional setting (first 250 batches).

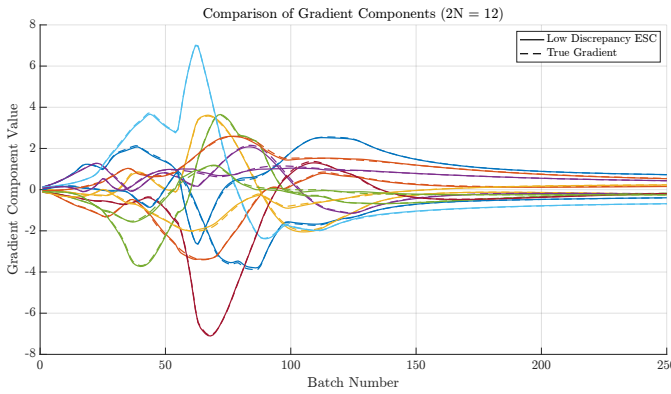


Fig. 7: Comparison of estimated gradient components and true gradient components over batch iterations using 60 perturbations in twelve-dimensional setting (first 250 batches).

V. CONCLUSION

We presented a method for estimating the gradient of an unknown function using low discrepancy perturbations. Perturbations are generated using low-discrepancy Sobol sequences and orthonormal basis functions, providing strong numerical stability. Theoretical analysis and simulations demonstrate the

numerical stability and efficiency of the method, particularly in high-dimensional settings. The effectiveness of our method is demonstrated in the Integrated Sensing and Communication application in the multi-agent setting.

REFERENCES

- [1] M. Krstić and H.-H. Wang, “Stability of extremum seeking feedback for general nonlinear dynamic systems,” *Automatica*, vol. 36, no. 4, pp. 595–601, 2000.
- [2] M. Krstić, “Performance improvement and limitations in extremum seeking control,” *Systems & Control Letters*, vol. 39, pp. 313–326, 2000.
- [3] H.-H. Wang, M. Krstić, and G. Bastin, “Optimizing bioreactors by extremum-seeking control,” *International Journal of Adaptive Control and Signal Processing*, vol. 13, no. 8, pp. 651–669, 1999.
- [4] K. B. Ariyur and M. Krstić, *Real-Time Optimization by Extremum Seeking Control*. John Wiley & Sons, 2003.
- [5] M. Colombino, E. Dall’Anese, and A. Bernstein, “Online optimization as a feedback controller: Stability and tracking,” *IEEE Transactions on Control of Network Systems*, vol. 7, no. 1, pp. 422–432, 2020.
- [6] D. Krishnamoorthy and S. Skogestad, “Real-time optimization as a feedback control problem – a review,” *Computers & Chemical Engineering*, vol. 165, p. 107951, 2022.
- [7] A. Hauswirth, S. Bolognani, G. Hug, and F. Dörfler, “Optimization algorithms as robust feedback controllers,” *Annual Reviews in Control*, vol. 57, p. 100902, 2024.
- [8] T. Liu, T. Liu, and Z.-P. Jiang, “Feedback optimization of nonlinear strict-feedback systems,” *Journal of Systems Science and Complexity*, 2025.
- [9] N. Hansen and A. Ostermeier, “Completely derandomized self-adaptation in evolution strategies,” *Evolutionary Computation*, vol. 9, no. 2, pp. 159–195, 2001.
- [10] C. Manzie and M. Krstić, “Extremum seeking with stochastic perturbations,” *IEEE Transactions on Automatic Control*, 2009, technical Note.
- [11] C. Danielson, S. A. Bortoff, and A. Chakrabarty, “Extremum seeking control with an adaptive gain based on gradient estimation error,” *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, vol. 53, no. 1, pp. 152–164, 2023.
- [12] B. Hunnekens, M. A. Haring, N. van de Wouw, and H. Nijmeijer, “A dither-free extremum-seeking control approach using 1st-order least-squares fits for gradient estimation,” in *53rd IEEE conference on decision and control*. IEEE, 2014, pp. 2679–2684.
- [13] I. M. Sobol, “On the distribution of points in a cube and the approximate evaluation of integrals,” *USSR Computational Mathematics and Mathematical Physics*, vol. 7, no. 4, pp. 86–112, 1967.
- [14] J. G. van der Corput, “Verteilungsfunktionen I,” *Proceedings of the Koninklijke Nederlandse Akademie van Wetenschappen*, vol. 38, pp. 813–821, 1935.
- [15] J. Dick and F. Pillichshammer, *Digital Nets and Sequences: Discrepancy Theory and Quasi-Monte Carlo Integration*, ser. Cambridge Monographs on Applied and Computational Mathematics. Cambridge University Press, 2010, vol. 17.
- [16] R. A. Horn and C. R. Johnson, *Matrix Analysis*, 2nd ed. Cambridge University Press, 2012.
- [17] L. Cui, Z.-P. Jiang, and E. D. Sontag, “Small-disturbance input-to-state stability of perturbed gradient flows: Applications to lqr problem,” *Systems & Control Letters*, vol. 188, p. 105804, 2024.