

Model-Based Derivative-Free Optimization: Algorithms and Approximation Theory

Part 2: Fundamentals of Model-Based DFO

Lindon Roberts, University of Melbourne (lindon.roberts@unimelb.edu.au)

Supported by the Australian Research Council (DE240100006)

MoCaO Lectures

28 July 2026

Model-Based DFO

We want to solve

$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x})$$

where f is nonconvex and has L -Lipschitz continuous gradient (not accessible to algorithm!).

Model-Based DFO

We want to solve

$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x})$$

where f is nonconvex and has L -Lipschitz continuous gradient (not accessible to algorithm!).

Basic approach draws on standard/derivative-based trust-region methods:

- Build a local quadratic interpolation model to approximate f near $\mathbf{x}_k$
- Calculate a step $\|\mathbf{s}_k\| \leq \Delta_k$ that decreases the model
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$, and use this to accept/reject the step and choose Δ_{k+1}

e.g. [Conn et al., 2009]

1. **Main MBDFO Algorithm**
2. Convergence Theory
3. Interpolation Model Construction

Model Accuracy

Yesterday, we considered linear models based on finite differencing

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := f(\mathbf{x}_k) + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k), \quad \text{where} \quad [\mathbf{g}_k]_i = \frac{f(\mathbf{x}_k + h\mathbf{e}_i) - f(\mathbf{x}_k)}{h}$$

with error bounds

$$\|\mathbf{g}_k - \nabla f(\mathbf{x}_k)\| \leq \frac{\sqrt{n}Lh}{2} \quad \text{and} \quad |f(\mathbf{y}) - m_k(\mathbf{y})| \leq \frac{L}{2}\|\mathbf{y} - \mathbf{x}_k\|^2 + \frac{\sqrt{n}Lh}{2}\|\mathbf{y} - \mathbf{x}_k\|,$$

Model Accuracy

Yesterday, we considered linear models based on finite differencing

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := f(\mathbf{x}_k) + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k), \quad \text{where} \quad [\mathbf{g}_k]_i = \frac{f(\mathbf{x}_k + h\mathbf{e}_i) - f(\mathbf{x}_k)}{h}$$

with error bounds

$$\|\mathbf{g}_k - \nabla f(\mathbf{x}_k)\| \leq \frac{\sqrt{n}Lh}{2} \quad \text{and} \quad |f(\mathbf{y}) - m_k(\mathbf{y})| \leq \frac{L}{2}\|\mathbf{y} - \mathbf{x}_k\|^2 + \frac{\sqrt{n}Lh}{2}\|\mathbf{y} - \mathbf{x}_k\|,$$

For trust-region methods, our tentative step will be $\mathbf{x}_k + \mathbf{s}_k$, with $\|\mathbf{s}_k\| \leq \Delta_k$, and so

$$|f(\mathbf{x}_k + \mathbf{s}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)| \leq \frac{L}{2}\Delta_k^2 + \frac{\sqrt{n}Lh}{2}\Delta_k$$

Model Accuracy

Yesterday, we considered linear models based on finite differencing

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := f(\mathbf{x}_k) + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k), \quad \text{where} \quad [\mathbf{g}_k]_i = \frac{f(\mathbf{x}_k + h\mathbf{e}_i) - f(\mathbf{x}_k)}{h}$$

with error bounds

$$\|\mathbf{g}_k - \nabla f(\mathbf{x}_k)\| \leq \frac{\sqrt{n}Lh}{2} \quad \text{and} \quad |f(\mathbf{y}) - m_k(\mathbf{y})| \leq \frac{L}{2}\|\mathbf{y} - \mathbf{x}_k\|^2 + \frac{\sqrt{n}Lh}{2}\|\mathbf{y} - \mathbf{x}_k\|,$$

For trust-region methods, our tentative step will be $\mathbf{x}_k + \mathbf{s}_k$, with $\|\mathbf{s}_k\| \leq \Delta_k$, and so

$$|f(\mathbf{x}_k + \mathbf{s}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)| \leq \frac{L}{2}\Delta_k^2 + \frac{\sqrt{n}Lh}{2}\Delta_k$$

For minimal error, pick h no bigger than Δ_k ; e.g. if $h = \Delta_k$ we have

$$\|\mathbf{g}_k - \nabla f(\mathbf{x}_k)\| \leq \frac{\sqrt{n}L}{2}\Delta_k \quad \text{and} \quad |f(\mathbf{x}_k + \mathbf{s}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)| \leq \left(\frac{L}{2} + \frac{\sqrt{n}L}{2}\right)\Delta_k^2.$$

Model Accuracy

So far, we have two possible [local, linear models for the objective](#), with error bounds:

Model Accuracy

So far, we have two possible **local, linear models for the objective**, with error bounds:

Linear Taylor Series

- Model:

$$m_k(\mathbf{y}) = f(\mathbf{x}_k) + \nabla f(\mathbf{x}_k)^T (\mathbf{y} - \mathbf{x}_k)$$

- Model error:

$$|f(\mathbf{y}) - m_k(\mathbf{y})| \leq \frac{L}{2} \|\mathbf{y} - \mathbf{x}_k\|^2 = \mathcal{O}(\Delta_k^2)$$

- Gradient error:

$$\|\nabla f(\mathbf{x}_k) - \nabla m_k(\mathbf{x}_k)\| = 0$$

Finite Differencing, $h = \Delta_k$

- Model:

$$m_k(\mathbf{y}) = f(\mathbf{x}_k) + \mathbf{g}_k^T (\mathbf{y} - \mathbf{x}_k)$$

- Model error:

$$|f(\mathbf{y}) - m_k(\mathbf{y})| = \mathcal{O}(\Delta_k^2)$$

- Gradient error:

$$\|\nabla f(\mathbf{x}_k) - \nabla m_k(\mathbf{x}_k)\| \leq \mathcal{O}(\Delta_k)$$

for all $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$.

Fully Linear Models

Motivated by these, in our algorithm we assume we have a **local, quadratic model for the objective**

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = f(\mathbf{x}_k) + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k) + \frac{1}{2}(\mathbf{y} - \mathbf{x}_k)^T H_k(\mathbf{y} - \mathbf{x}_k)$$

Fully Linear Models

Motivated by these, in our algorithm we assume we have a **local, quadratic model for the objective**

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = f(\mathbf{x}_k) + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k) + \frac{1}{2}(\mathbf{y} - \mathbf{x}_k)^T H_k(\mathbf{y} - \mathbf{x}_k)$$

Our required notion of model accuracy is:

[Conn et al., 2009]

Fully Linear Model

A quadratic model m_k is a **fully linear** approximation to f in $B(\mathbf{x}_k, \Delta_k)$ if:

$$|f(\mathbf{y}) - m_k(\mathbf{y})| \leq \kappa_{\text{mf}} \Delta_k^2 \quad \text{and} \quad \|\nabla f(\mathbf{x}_k) - \nabla m_k(\mathbf{x}_k)\| \leq \kappa_{\text{mg}} \Delta_k$$

for all $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, with $\kappa_{\text{mf}}, \kappa_{\text{mg}} \geq 0$ independent of k .

Fully Linear Models

Motivated by these, in our algorithm we assume we have a **local, quadratic model for the objective**

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = f(\mathbf{x}_k) + \mathbf{g}_k^T (\mathbf{y} - \mathbf{x}_k) + \frac{1}{2} (\mathbf{y} - \mathbf{x}_k)^T H_k (\mathbf{y} - \mathbf{x}_k)$$

Our required notion of model accuracy is:

[Conn et al., 2009]

Fully Linear Model

A quadratic model m_k is a **fully linear** approximation to f in $B(\mathbf{x}_k, \Delta_k)$ if:

$$|f(\mathbf{y}) - m_k(\mathbf{y})| \leq \kappa_{\text{mf}} \Delta_k^2 \quad \text{and} \quad \|\nabla f(\mathbf{x}_k) - \nabla m_k(\mathbf{x}_k)\| \leq \kappa_{\text{mg}} \Delta_k$$

for all $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, with $\kappa_{\text{mf}}, \kappa_{\text{mg}} \geq 0$ independent of k .

For example, linear Taylor series or finite differencing gradients, with $H_k = 0$.

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Issue

In MBDFO, Δ_k is overloaded!

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Issue

In MBDFO, Δ_k is overloaded!

- Controls length of step, $\|\mathbf{s}_k\| \leq \Delta_k$

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Issue

In MBDFO, Δ_k is overloaded!

- Controls length of step, $\|\mathbf{s}_k\| \leq \Delta_k$
- Controls model accuracy, e.g. $|f(\mathbf{y}) - m_k(\mathbf{y})| = \mathcal{O}(\Delta_k^2)$

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Issue

In MBDFO, Δ_k is overloaded!

- Controls length of step, $\|\mathbf{s}_k\| \leq \Delta_k$
- Controls model accuracy, e.g. $|f(\mathbf{y}) - m_k(\mathbf{y})| = \mathcal{O}(\Delta_k^2)$
- Controls termination: stop when $\Delta_k \leq \Delta_{\min}$ (can't check $\|\nabla f(\mathbf{x}_k)\| \leq \epsilon$, and checking $\|\mathbf{g}_k\| \leq \epsilon$ only works if model error—i.e. Δ_k —is small!)

Algorithm Preliminaries

To build our MBDFO algorithm, we just need to replace a Taylor-type local model with a fully linear model (constructed however we wish) then proceed as before: minimize the model to get step $\mathbf{s}_k$, evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and accept/reject, ...

Issue

In MBDFO, Δ_k is overloaded!

- Controls length of step, $\|\mathbf{s}_k\| \leq \Delta_k$
- Controls model accuracy, e.g. $|f(\mathbf{y}) - m_k(\mathbf{y})| = \mathcal{O}(\Delta_k^2)$
- Controls termination: stop when $\Delta_k \leq \Delta_{\min}$ (can't check $\|\nabla f(\mathbf{x}_k)\| \leq \epsilon$, and checking $\|\mathbf{g}_k\| \leq \epsilon$ only works if model error—i.e. Δ_k —is small!)

This requires one small modification to the basic algorithm (*criticality check*)

With this, one iteration of the basic MBDFO algorithm is:

- Given $\mathbf{x}_k$ and Δ_k , **build fully linear model m_k**
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:

With this, one iteration of the basic MBDFO algorithm is:

- Given $\mathbf{x}_k$ and Δ_k , **build fully linear model m_k**
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:
 - (Very successful) If $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = 2\Delta_k$
 - (Successful) If $0.1 \leq \rho_k < 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = \Delta_k$
 - (Unsuccessful) If $\rho_k < 0.1$ or $\|\mathbf{g}_k\| < \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k$ and $\Delta_{k+1} = \frac{1}{2}\Delta_k$

With this, one iteration of the basic MBDFO algorithm is:

- Given $\mathbf{x}_k$ and Δ_k , **build fully linear model m_k**
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:
 - (Very successful) If $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = 2\Delta_k$
 - (Successful) If $0.1 \leq \rho_k < 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = \Delta_k$
 - (Unsuccessful) If $\rho_k < 0.1$ or $\|\mathbf{g}_k\| < \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k$ and $\Delta_{k+1} = \frac{1}{2}\Delta_k$

Criticality check: if $\|\mathbf{g}_k\| \ll \Delta_k$, shrink Δ_k (without calculating $\mathbf{s}_k$). Usually $\mu \ll 1$.

With this, one iteration of the basic MBDFO algorithm is:

- Given $\mathbf{x}_k$ and Δ_k , **build fully linear model m_k**
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:
 - (Very successful) If $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = 2\Delta_k$
 - (Successful) If $0.1 \leq \rho_k < 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = \Delta_k$
 - (Unsuccessful) If $\rho_k < 0.1$ or $\|\mathbf{g}_k\| < \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k$ and $\Delta_{k+1} = \frac{1}{2}\Delta_k$

Criticality check: if $\|\mathbf{g}_k\| \ll \Delta_k$, shrink Δ_k (without calculating $\mathbf{s}_k$). Usually $\mu \ll 1$.

i.e. If the algorithm thinks we're near a solution, shrink Δ_k to improve model accuracy, so we can be sure.

Step Calculation

Aside from model construction, the other key part of the algorithm is the [trust-region subproblem](#) (step calculation): e.g. [Conn et al., 2000]

$$\min_{\mathbf{s} \in \mathbb{R}^n} m(\mathbf{s}) := \mathbf{g}^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T H \mathbf{s} \quad \text{s.t. } \|\mathbf{s}\| \leq \Delta$$

Step Calculation

Aside from model construction, the other key part of the algorithm is the **trust-region subproblem** (step calculation): e.g. [Conn et al., 2000]

$$\min_{\mathbf{s} \in \mathbb{R}^n} m(\mathbf{s}) := \mathbf{g}^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T H \mathbf{s} \quad \text{s.t. } \|\mathbf{s}\| \leq \Delta$$

This problem is nonconvex (from H), but we can still characterize *global* minima.

Step Calculation

Aside from model construction, the other key part of the algorithm is the [trust-region subproblem](#) (step calculation): e.g. [Conn et al., 2000]

$$\min_{\mathbf{s} \in \mathbb{R}^n} m(\mathbf{s}) := \mathbf{g}^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T H \mathbf{s} \quad \text{s.t. } \|\mathbf{s}\| \leq \Delta$$

This problem is nonconvex (from H), but we can still characterize *global* minima.

Theorem

A point $\mathbf{s}$ is a global minimizer if and only if the KKT conditions (with Lagrange multiplier $\lambda \in \mathbb{R}$)

$$(H + \lambda I)\mathbf{s} = -\mathbf{g},$$

$$\|\mathbf{s}\| \leq \Delta, \quad \lambda \geq 0, \quad \lambda(\Delta - \|\mathbf{s}\|) = 0,$$

hold and $H + \lambda I$ is positive semidefinite (unique minimizer if positive definite).

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

- Interior solution: $\mathbf{s} = -H^{-1}\mathbf{g}$ if H positive definite and $\|\mathbf{s}\| \leq \Delta$ (i.e. $\lambda = 0$)

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

- Interior solution: $\mathbf{s} = -H^{-1}\mathbf{g}$ if H positive definite and $\|\mathbf{s}\| \leq \Delta$ (i.e. $\lambda = 0$)
- Boundary solution: $\lambda > 0$, so $\|\mathbf{s}\| = \Delta$ and $\mathbf{s} = -(H + \lambda I)^{-1}\mathbf{g}$ — 1D rootfinding problem in λ

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

- Interior solution: $\mathbf{s} = -H^{-1}\mathbf{g}$ if H positive definite and $\|\mathbf{s}\| \leq \Delta$ (i.e. $\lambda = 0$)
- Boundary solution: $\lambda > 0$, so $\|\mathbf{s}\| = \Delta$ and $\mathbf{s} = -(H + \lambda I)^{-1}\mathbf{g}$ — 1D rootfinding problem in λ
- Hard case: if H negative semidefinite and $\mathbf{g}$ orthogonal to eigenspace of $\lambda_{\min}(H)$, then $(H + \lambda I)\mathbf{s} = -\mathbf{g}$ for $\lambda = -\lambda_{\min}(H)$ is singular but consistent (pick boundary solution)

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

- Interior solution: $\mathbf{s} = -H^{-1}\mathbf{g}$ if H positive definite and $\|\mathbf{s}\| \leq \Delta$ (i.e. $\lambda = 0$)
- Boundary solution: $\lambda > 0$, so $\|\mathbf{s}\| = \Delta$ and $\mathbf{s} = -(H + \lambda I)^{-1}\mathbf{g}$ — 1D rootfinding problem in λ
- Hard case: if H negative semidefinite and $\mathbf{g}$ orthogonal to eigenspace of $\lambda_{\min}(H)$, then $(H + \lambda I)\mathbf{s} = -\mathbf{g}$ for $\lambda = -\lambda_{\min}(H)$ is singular but consistent (pick boundary solution)

For low-dimensional problems (e.g. $n = \mathcal{O}(100)$), finding a global minimizer is practical *and does improve the performance of MBDFO methods*.

Step Calculation

This gives us a way to find a global minimizer of the trust-region subproblem:

- Interior solution: $\mathbf{s} = -H^{-1}\mathbf{g}$ if H positive definite and $\|\mathbf{s}\| \leq \Delta$ (i.e. $\lambda = 0$)
- Boundary solution: $\lambda > 0$, so $\|\mathbf{s}\| = \Delta$ and $\mathbf{s} = -(H + \lambda I)^{-1}\mathbf{g}$ — 1D rootfinding problem in λ
- Hard case: if H negative semidefinite and $\mathbf{g}$ orthogonal to eigenspace of $\lambda_{\min}(H)$, then $(H + \lambda I)\mathbf{s} = -\mathbf{g}$ for $\lambda = -\lambda_{\min}(H)$ is singular but consistent (pick boundary solution)

For low-dimensional problems (e.g. $n = \mathcal{O}(100)$), finding a global minimizer is practical *and does improve the performance of MBDFO methods*.

However, for theory (and large-scale problems), approximate minimizers are acceptable...

Approximate Step Calculation

The simplest approximate minimizer is the **Cauchy point** — steepest descent with exact linesearch:

$$\mathbf{s}_C = -\alpha_C \mathbf{g}, \quad \alpha_C = \arg \min_{\alpha > 0} m(-\alpha \mathbf{g}) \quad \text{s.t. } \|\alpha \mathbf{g}\| \leq \Delta$$

Approximate Step Calculation

The simplest approximate minimizer is the **Cauchy point** — steepest descent with exact linesearch:

$$\mathbf{s}_C = -\alpha_C \mathbf{g}, \quad \alpha_C = \arg \min_{\alpha > 0} m(-\alpha \mathbf{g}) \quad \text{s.t. } \|\alpha \mathbf{g}\| \leq \Delta$$

If objective convex along this direction, get (constrained) minimizer, otherwise objective unbounded below (so must be at the boundary). Hence

$$\alpha_C = \begin{cases} \min \left(\frac{\|\mathbf{g}\|}{\mathbf{g}^T H \mathbf{g}}, \frac{\Delta}{\|\mathbf{g}\|} \right), & \mathbf{g}^T H \mathbf{g} > 0, \\ \frac{\Delta}{\|\mathbf{g}\|}, & \mathbf{g}^T H \mathbf{g} \leq 0 \end{cases}$$

Approximate Step Calculation

The simplest approximate minimizer is the **Cauchy point** — steepest descent with exact linesearch:

$$\mathbf{s}_C = -\alpha_C \mathbf{g}, \quad \alpha_C = \arg \min_{\alpha > 0} m(-\alpha \mathbf{g}) \quad \text{s.t. } \|\alpha \mathbf{g}\| \leq \Delta$$

If objective convex along this direction, get (constrained) minimizer, otherwise objective unbounded below (so must be at the boundary). Hence

$$\alpha_C = \begin{cases} \min \left(\frac{\|\mathbf{g}\|}{\mathbf{g}^T H \mathbf{g}}, \frac{\Delta}{\|\mathbf{g}\|} \right), & \mathbf{g}^T H \mathbf{g} > 0, \\ \frac{\Delta}{\|\mathbf{g}\|}, & \mathbf{g}^T H \mathbf{g} \leq 0 \end{cases}$$

which gives objective reduction of at least

$$m(\mathbf{0}) - m(\mathbf{s}_C) \geq \frac{1}{2} \|\mathbf{g}\| \min \left(\Delta, \frac{\|\mathbf{g}\|}{\|H\| + 1} \right)$$

Approximate Step Calculation

The simplest approximate minimizer is the **Cauchy point** — steepest descent with exact linesearch:

$$\mathbf{s}_C = -\alpha_C \mathbf{g}, \quad \alpha_C = \arg \min_{\alpha > 0} m(-\alpha \mathbf{g}) \quad \text{s.t. } \|\alpha \mathbf{g}\| \leq \Delta$$

If objective convex along this direction, get (constrained) minimizer, otherwise objective unbounded below (so must be at the boundary). Hence

$$\alpha_C = \begin{cases} \min \left(\frac{\|\mathbf{g}\|}{\mathbf{g}^T H \mathbf{g}}, \frac{\Delta}{\|\mathbf{g}\|} \right), & \mathbf{g}^T H \mathbf{g} > 0, \\ \frac{\Delta}{\|\mathbf{g}\|}, & \mathbf{g}^T H \mathbf{g} \leq 0 \end{cases}$$

which gives objective reduction of at least

$$m(\mathbf{0}) - m(\mathbf{s}_C) \geq \frac{1}{2} \|\mathbf{g}\| \min \left(\Delta, \frac{\|\mathbf{g}\|}{\|H\| + 1} \right)$$

This is the minimum required by theory, but is not enough to be practically useful.

An intermediate alternative is to use the [conjugate gradient method](#).

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $Hs = -g$ via minimizing

$$\min_s g^T s + \frac{1}{2} s^T H s$$

i.e. our subproblem without the constraint, and with H positive definite.

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $Hs = -g$ via minimizing

$$\min_s g^T s + \frac{1}{2} s^T H s$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $H\mathbf{s} = -\mathbf{g}$ via minimizing

$$\min_{\mathbf{s}} \mathbf{g}^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T H \mathbf{s}$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

- Solution $H\mathbf{s} = -\mathbf{g}$ reached (interior global minimizer)

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $H\mathbf{s} = -\mathbf{g}$ via minimizing

$$\min_{\mathbf{s}} \mathbf{g}^T \mathbf{s} + \frac{1}{2} \mathbf{s}^T H \mathbf{s}$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

- Solution $H\mathbf{s} = -\mathbf{g}$ reached (interior global minimizer)
- $\|\mathbf{s}\| > \Delta$ reached — shrink step to $\|\mathbf{s}\| = \Delta$ and terminate

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $Hs = -g$ via minimizing

$$\min_s g^T s + \frac{1}{2} s^T H s$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

- Solution $Hs = -g$ reached (interior global minimizer)
- $\|s\| > \Delta$ reached — shrink step to $\|s\| = \Delta$ and terminate
- Negative curvature detected — follow this direction until $\|s\| = \Delta$ and terminate

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $Hs = -g$ via minimizing

$$\min_s g^T s + \frac{1}{2} s^T H s$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

- Solution $Hs = -g$ reached (interior global minimizer)
- $\|s\| > \Delta$ reached — shrink step to $\|s\| = \Delta$ and terminate
- Negative curvature detected — follow this direction until $\|s\| = \Delta$ and terminate

This is cheap to implement for large-scale problems (Hessian-vector products only), and provides a reasonable model decrease in practice.

An intermediate alternative is to use the [conjugate gradient method](#). CG solves symmetric positive definite linear systems $Hs = -g$ via minimizing

$$\min_s g^T s + \frac{1}{2} s^T H s$$

i.e. our subproblem without the constraint, and with H positive definite.

Idea: run CG without modification until:

- Solution $Hs = -g$ reached (interior global minimizer)
- $\|s\| > \Delta$ reached — shrink step to $\|s\| = \Delta$ and terminate
- Negative curvature detected — follow this direction until $\|s\| = \Delta$ and terminate

This is cheap to implement for large-scale problems (Hessian-vector products only), and provides a reasonable model decrease in practice. Always achieves Cauchy decrease.

1. Main MBDFO Algorithm
2. **Convergence Theory**
3. Interpolation Model Construction

Main Algorithm

Let's now prove **global convergence** (i.e. converges for any $\mathbf{x}_0$) of our algorithm

- Given $\mathbf{x}_k$ and Δ_k , build fully linear model m_k
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:
 - (Very successful) If $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = 2\Delta_k$
 - (Successful) If $0.1 \leq \rho_k < 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = \Delta_k$
 - (Unsuccessful) If $\rho_k < 0.1$ or $\|\mathbf{g}_k\| < \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k$ and $\Delta_{k+1} = \frac{1}{2}\Delta_k$

Main Algorithm

Let's now prove **global convergence** (i.e. converges for any $\mathbf{x}_0$) of our algorithm

- Given $\mathbf{x}_k$ and Δ_k , build fully linear model m_k
- Calculate a step, $\mathbf{s}_k \approx \arg \min_{\|\mathbf{s}\| \leq \Delta_k} m_k(\mathbf{x}_k + \mathbf{s})$
- Evaluate $f(\mathbf{x}_k + \mathbf{s}_k)$ and ratio $\rho_k = [f(\mathbf{x}_k) - f(\mathbf{x}_k + \mathbf{s}_k)]/[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]$
- Accept/reject the step and update trust-region radius:
 - (Very successful) If $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = 2\Delta_k$
 - (Successful) If $0.1 \leq \rho_k < 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{s}_k$ and $\Delta_{k+1} = \Delta_k$
 - (Unsuccessful) If $\rho_k < 0.1$ or $\|\mathbf{g}_k\| < \mu\Delta_k$, set $\mathbf{x}_{k+1} = \mathbf{x}_k$ and $\Delta_{k+1} = \frac{1}{2}\Delta_k$

Assumptions:

- ∇f L -Lipschitz continuous and f bounded below, $f(\mathbf{x}) \geq f_{\text{low}}$ for all $\mathbf{x} \in \mathbb{R}^n$
- Step $\mathbf{s}_k$ always achieves Cauchy decrease in the model
- Uniformly bounded model Hessians, $\|H_k\| + 1 \leq \kappa_H$ for all k

Convergence: Step 1

Lemma

If $\mathbf{g}_k \neq \mathbf{0}$ and $\Delta_k \leq C_0 \|\mathbf{g}_k\|$, where $C_0 := \min \left(\frac{0.3}{4\kappa_{\text{mf}}}, \frac{1}{\kappa_H}, \frac{1}{\mu} \right)$, then iteration k is very successful (i.e. $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu \Delta_k$).

Convergence: Step 1

Lemma

If $\mathbf{g}_k \neq \mathbf{0}$ and $\Delta_k \leq C_0 \|\mathbf{g}_k\|$, where $C_0 := \min\left(\frac{0.3}{4\kappa_{\text{mf}}}, \frac{1}{\kappa_H}, \frac{1}{\mu}\right)$, then iteration k is very successful (i.e. $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$).

Proof: By definition of C_0 , we have $\Delta_k \leq \frac{\|\mathbf{g}_k\|}{\kappa_H} \leq \frac{\|\mathbf{g}_k\|}{\|H_k\|+1}$, so Cauchy decrease gives

$$m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k) \geq \frac{1}{2} \|\mathbf{g}_k\| \min\left(\Delta_k, \frac{\|\mathbf{g}_k\|}{\|H_k\|+1}\right) = \frac{1}{2} \|\mathbf{g}_k\| \Delta_k.$$

Convergence: Step 1

Lemma

If $\mathbf{g}_k \neq \mathbf{0}$ and $\Delta_k \leq C_0 \|\mathbf{g}_k\|$, where $C_0 := \min\left(\frac{0.3}{4\kappa_{\text{mf}}}, \frac{1}{\kappa_H}, \frac{1}{\mu}\right)$, then iteration k is very successful (i.e. $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$).

Proof: By definition of C_0 , we have $\Delta_k \leq \frac{\|\mathbf{g}_k\|}{\kappa_H} \leq \frac{\|\mathbf{g}_k\|}{\|H_k\|+1}$, so Cauchy decrease gives

$$m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k) \geq \frac{1}{2} \|\mathbf{g}_k\| \min\left(\Delta_k, \frac{\|\mathbf{g}_k\|}{\|H_k\|+1}\right) = \frac{1}{2} \|\mathbf{g}_k\| \Delta_k.$$

Then, by definition of ρ_k and using fully linear models, we have

$$|\rho_k - 1| = \frac{|[f(\mathbf{x}_k) - m_k(\mathbf{x}_k)] - [f(\mathbf{x}_k + \mathbf{s}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]|}{|m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)|} \leq \frac{2\kappa_{\text{mf}}\Delta_k^2}{\frac{1}{2}\|\mathbf{g}_k\|\Delta_k} \leq 4\kappa_{\text{mf}}C_0$$

Convergence: Step 1

Lemma

If $\mathbf{g}_k \neq \mathbf{0}$ and $\Delta_k \leq C_0 \|\mathbf{g}_k\|$, where $C_0 := \min\left(\frac{0.3}{4\kappa_{\text{mf}}}, \frac{1}{\kappa_H}, \frac{1}{\mu}\right)$, then iteration k is very successful (i.e. $\rho_k \geq 0.7$ and $\|\mathbf{g}_k\| \geq \mu\Delta_k$).

Proof: By definition of C_0 , we have $\Delta_k \leq \frac{\|\mathbf{g}_k\|}{\kappa_H} \leq \frac{\|\mathbf{g}_k\|}{\|H_k\|+1}$, so Cauchy decrease gives

$$m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k) \geq \frac{1}{2} \|\mathbf{g}_k\| \min\left(\Delta_k, \frac{\|\mathbf{g}_k\|}{\|H_k\| + 1}\right) = \frac{1}{2} \|\mathbf{g}_k\| \Delta_k.$$

Then, by definition of ρ_k and using fully linear models, we have

$$|\rho_k - 1| = \frac{|[f(\mathbf{x}_k) - m_k(\mathbf{x}_k)] - [f(\mathbf{x}_k + \mathbf{s}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)]|}{|m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)|} \leq \frac{2\kappa_{\text{mf}}\Delta_k^2}{\frac{1}{2}\|\mathbf{g}_k\|\Delta_k} \leq 4\kappa_{\text{mf}}C_0$$

so $|\rho_k - 1| \leq 0.3$, which implies $\rho_k \geq 0.7$. $\|\mathbf{g}_k\| \geq \mu\Delta_k$ is trivial by definition of C_0 .

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K-1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Proof: Result is trivial for $k = 0$, suppose it holds for some $k = 0, \dots, K - 1$. For a contradiction, suppose $\Delta_{k+1} < \Delta_{\min}(\epsilon)$.

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Proof: Result is trivial for $k = 0$, suppose it holds for some $k = 0, \dots, K - 1$. For a contradiction, suppose $\Delta_{k+1} < \Delta_{\min}(\epsilon)$.

Then $\Delta_{k+1} < \Delta_{\min}(\epsilon) \leq \Delta_k$, i.e. iteration k must be unsuccessful.

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Proof: Result is trivial for $k = 0$, suppose it holds for some $k = 0, \dots, K - 1$. For a contradiction, suppose $\Delta_{k+1} < \Delta_{\min}(\epsilon)$.

Then $\Delta_{k+1} < \Delta_{\min}(\epsilon) \leq \Delta_k$, i.e. iteration k must be unsuccessful. By previous lemma, we must have $\Delta_k > C_0 \|\mathbf{g}_k\|$.

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Proof: Result is trivial for $k = 0$, suppose it holds for some $k = 0, \dots, K - 1$. For a contradiction, suppose $\Delta_{k+1} < \Delta_{\min}(\epsilon)$.

Then $\Delta_{k+1} < \Delta_{\min}(\epsilon) \leq \Delta_k$, i.e. iteration k must be unsuccessful. By previous lemma, we must have $\Delta_k > C_0 \|\mathbf{g}_k\|$.

Since model is fully linear,

$$\epsilon \leq \|\nabla f(\mathbf{x}_k)\| \leq \|\mathbf{g}_k\| + \|\nabla f(\mathbf{x}_k) - \mathbf{g}_k\| < C_0^{-1} \Delta_k + \kappa_{\text{mg}} \Delta_k$$

Convergence: Step 2

Lemma

If $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K-1$, then $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all $k = 0, \dots, K$, where $\Delta_{\min}(\epsilon) := \min(\Delta_0, \frac{\epsilon/2}{\kappa_{\text{mg}} + C_0^{-1}})$. Note $\Delta_{\min}(\epsilon) = \mathcal{O}(\epsilon)$.

Proof: Result is trivial for $k = 0$, suppose it holds for some $k = 0, \dots, K-1$. For a contradiction, suppose $\Delta_{k+1} < \Delta_{\min}(\epsilon)$.

Then $\Delta_{k+1} < \Delta_{\min}(\epsilon) \leq \Delta_k$, i.e. iteration k must be unsuccessful. By previous lemma, we must have $\Delta_k > C_0 \|\mathbf{g}_k\|$.

Since model is fully linear,

$$\epsilon \leq \|\nabla f(\mathbf{x}_k)\| \leq \|\mathbf{g}_k\| + \|\nabla f(\mathbf{x}_k) - \mathbf{g}_k\| < C_0^{-1} \Delta_k + \kappa_{\text{mg}} \Delta_k$$

Hence $\Delta_k > \frac{\epsilon}{\kappa_{\text{mg}} + C_0^{-1}}$, so $\Delta_{k+1} = \frac{1}{2} \Delta_k \geq \Delta_{\min}(\epsilon)$, a contradiction.

Convergence: Main Result

We can now prove a global convergence and worst-case complexity bound. Suppose $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$ — we seek an upper bound on K .

Convergence: Main Result

We can now prove a global convergence and worst-case complexity bound. Suppose $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$ — we seek an upper bound on K .

Partition iterations into (very) successful and unsuccessful, $\{0, \dots, K - 1\} = \mathcal{S} \cup \mathcal{U}$, and note previous lemma gives $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all such k .

Convergence: Main Result

We can now prove a global convergence and worst-case complexity bound. Suppose $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$ — we seek an upper bound on K .

Partition iterations into (very) successful and unsuccessful, $\{0, \dots, K - 1\} = \mathcal{S} \cup \mathcal{U}$, and note previous lemma gives $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all such k .

Successful iterations: Cauchy decrease and criticality ($\|\mathbf{g}_k\| \geq \mu\Delta_k$) give

$$\begin{aligned} m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k) &\geq \frac{1}{2} \|\mathbf{g}_k\| \min \left(\Delta_k, \frac{\|\mathbf{g}_k\|}{\|H_k\| + 1} \right) \\ &\geq \frac{1}{2} \mu \Delta_{\min}(\epsilon) \min \left(\Delta_{\min}(\epsilon), \frac{\mu \Delta_{\min}(\epsilon)}{\kappa_H} \right) = \mathcal{O}(\Delta_{\min}(\epsilon)^2) \end{aligned}$$

Convergence: Main Result

We can now prove a global convergence and worst-case complexity bound. Suppose $\|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ for all $k = 0, \dots, K - 1$ — we seek an upper bound on K .

Partition iterations into (very) successful and unsuccessful, $\{0, \dots, K - 1\} = \mathcal{S} \cup \mathcal{U}$, and note previous lemma gives $\Delta_k \geq \Delta_{\min}(\epsilon)$ for all such k .

Successful iterations: Cauchy decrease and criticality ($\|\mathbf{g}_k\| \geq \mu\Delta_k$) give

$$\begin{aligned} m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k) &\geq \frac{1}{2} \|\mathbf{g}_k\| \min \left(\Delta_k, \frac{\|\mathbf{g}_k\|}{\|H_k\| + 1} \right) \\ &\geq \frac{1}{2} \mu \Delta_{\min}(\epsilon) \min \left(\Delta_{\min}(\epsilon), \frac{\mu \Delta_{\min}(\epsilon)}{\kappa_H} \right) = \mathcal{O}(\Delta_{\min}(\epsilon)^2) \end{aligned}$$

So, for successful steps, $f(\mathbf{x}_k) - f(\mathbf{x}_{k+1}) \geq 0.1[m_k(\mathbf{x}_k) - m_k(\mathbf{x}_k + \mathbf{s}_k)] \geq C_1 \Delta_{\min}(\epsilon)^2$.

Convergence: Main Result

Successful iterations: since $\mathbf{x}_{k+1} = \mathbf{x}_k$ if unsuccessful, over K iterations we have

$$f(\mathbf{x}_0) - f(\mathbf{x}_K) = \sum_{k=0}^{K-1} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] = \sum_{k \in \mathcal{S}} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] \geq C_1 \Delta_{\min}(\epsilon)^2 \cdot |\mathcal{S}|$$

Convergence: Main Result

Successful iterations: since $\mathbf{x}_{k+1} = \mathbf{x}_k$ if unsuccessful, over K iterations we have

$$f(\mathbf{x}_0) - f(\mathbf{x}_K) = \sum_{k=0}^{K-1} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] = \sum_{k \in \mathcal{S}} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] \geq C_1 \Delta_{\min}(\epsilon)^2 \cdot |\mathcal{S}|$$

Bounding LHS by $f(\mathbf{x}_0) - f_{\text{low}}$ gives a bound on the number of successful iterations

$$|\mathcal{S}| \leq \frac{f(\mathbf{x}_0) - f_{\text{low}}}{C_1 \Delta_{\min}(\epsilon)^2} = \mathcal{O}(\epsilon^{-2})$$

Convergence: Main Result

Successful iterations: since $\mathbf{x}_{k+1} = \mathbf{x}_k$ if unsuccessful, over K iterations we have

$$f(\mathbf{x}_0) - f(\mathbf{x}_K) = \sum_{k=0}^{K-1} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] = \sum_{k \in \mathcal{S}} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] \geq C_1 \Delta_{\min}(\epsilon)^2 \cdot |\mathcal{S}|$$

Bounding LHS by $f(\mathbf{x}_0) - f_{\text{low}}$ gives a bound on the number of successful iterations

$$|\mathcal{S}| \leq \frac{f(\mathbf{x}_0) - f_{\text{low}}}{C_1 \Delta_{\min}(\epsilon)^2} = \mathcal{O}(\epsilon^{-2})$$

Unsuccessful iterations: we have $\Delta_k \geq \Delta_{\min}(\epsilon)$, and $\Delta_{k+1} \leq 2\Delta_k$ for $k \in \mathcal{S}$ and $\Delta_{k+1} = 0.5\Delta_k$ for $k \in \mathcal{U}$.

Convergence: Main Result

Successful iterations: since $\mathbf{x}_{k+1} = \mathbf{x}_k$ if unsuccessful, over K iterations we have

$$f(\mathbf{x}_0) - f(\mathbf{x}_K) = \sum_{k=0}^{K-1} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] = \sum_{k \in \mathcal{S}} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] \geq C_1 \Delta_{\min}(\epsilon)^2 \cdot |\mathcal{S}|$$

Bounding LHS by $f(\mathbf{x}_0) - f_{\text{low}}$ gives a bound on the number of successful iterations

$$|\mathcal{S}| \leq \frac{f(\mathbf{x}_0) - f_{\text{low}}}{C_1 \Delta_{\min}(\epsilon)^2} = \mathcal{O}(\epsilon^{-2})$$

Unsuccessful iterations: we have $\Delta_k \geq \Delta_{\min}(\epsilon)$, and $\Delta_{k+1} \leq 2\Delta_k$ for $k \in \mathcal{S}$ and $\Delta_{k+1} = 0.5\Delta_k$ for $k \in \mathcal{U}$. Over K iterations, we get

$$\Delta_{\min}(\epsilon) \leq \Delta_K \leq 2^{|\mathcal{S}| - |\mathcal{U}|} \Delta_0$$

Convergence: Main Result

Successful iterations: since $\mathbf{x}_{k+1} = \mathbf{x}_k$ if unsuccessful, over K iterations we have

$$f(\mathbf{x}_0) - f(\mathbf{x}_K) = \sum_{k=0}^{K-1} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] = \sum_{k \in \mathcal{S}} [f(\mathbf{x}_k) - f(\mathbf{x}_{k+1})] \geq C_1 \Delta_{\min}(\epsilon)^2 \cdot |\mathcal{S}|$$

Bounding LHS by $f(\mathbf{x}_0) - f_{\text{low}}$ gives a bound on the number of successful iterations

$$|\mathcal{S}| \leq \frac{f(\mathbf{x}_0) - f_{\text{low}}}{C_1 \Delta_{\min}(\epsilon)^2} = \mathcal{O}(\epsilon^{-2})$$

Unsuccessful iterations: we have $\Delta_k \geq \Delta_{\min}(\epsilon)$, and $\Delta_{k+1} \leq 2\Delta_k$ for $k \in \mathcal{S}$ and $\Delta_{k+1} = 0.5\Delta_k$ for $k \in \mathcal{U}$. Over K iterations, we get

$$\Delta_{\min}(\epsilon) \leq \Delta_K \leq 2^{|\mathcal{S}| - |\mathcal{U}|} \Delta_0$$

So, $|\mathcal{U}| \leq |\mathcal{S}| + \log_2(\Delta_0/\Delta_{\min}(\epsilon)) = \mathcal{O}(\epsilon^{-2}) + \mathcal{O}(\log(\epsilon^{-1})) = \mathcal{O}(\epsilon^{-2})$.

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

- $\mathcal{O}(\epsilon^{-2})$ iterations is standard if we have exact gradients e.g. [Cartis et al., 2022]

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

- $\mathcal{O}(\epsilon^{-2})$ iterations is standard if we have exact gradients e.g. [Cartis et al., 2022]
- The penalty is in the size of the constant!

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

- $\mathcal{O}(\epsilon^{-2})$ iterations is standard if we have exact gradients e.g. [Cartis et al., 2022]
- The penalty is in the size of the constant! If $\kappa_m := \max(\kappa_{mf}, \kappa_{mg})$, we have
 - $\mathcal{O}((\kappa_m + \kappa_H)\epsilon^{-2})$ iterations for gradient-based trust-region methods
 - $\mathcal{O}(\kappa_H(\kappa_m + \kappa_H)^2\epsilon^{-2})$ iterations for MBDFO

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

- $\mathcal{O}(\epsilon^{-2})$ iterations is standard if we have exact gradients e.g. [Cartis et al., 2022]
- The penalty is in the size of the constant! If $\kappa_m := \max(\kappa_{mf}, \kappa_{mg})$, we have
 - $\mathcal{O}((\kappa_m + \kappa_H)\epsilon^{-2})$ iterations for gradient-based trust-region methods
 - $\mathcal{O}(\kappa_H(\kappa_m + \kappa_H)^2\epsilon^{-2})$ iterations for MBDFO

The difference comes from $\|\mathbf{g}_k\| \geq \mu\Delta_{\min}(\epsilon)$ vs. $\|\mathbf{g}_k\| = \|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ in Cauchy decrease for successful steps.

Convergence: Main Result

All together, we have a **worst-case complexity** bound and **global convergence**:

Theorem

$\|\nabla f(\mathbf{x}_k)\| < \epsilon$ first occurs after $\mathcal{O}(\epsilon^{-2})$ iterations. Hence $\liminf \|\nabla f(\mathbf{x}_k)\| = 0$.

How does this compare to derivative-based methods?

- $\mathcal{O}(\epsilon^{-2})$ iterations is standard if we have exact gradients e.g. [Cartis et al., 2022]
- The penalty is in the size of the constant! If $\kappa_m := \max(\kappa_{mf}, \kappa_{mg})$, we have
 - $\mathcal{O}((\kappa_m + \kappa_H)\epsilon^{-2})$ iterations for gradient-based trust-region methods
 - $\mathcal{O}(\kappa_H(\kappa_m + \kappa_H)^2\epsilon^{-2})$ iterations for MBDFO

The difference comes from $\|\mathbf{g}_k\| \geq \mu\Delta_{\min}(\epsilon)$ vs. $\|\mathbf{g}_k\| = \|\nabla f(\mathbf{x}_k)\| \geq \epsilon$ in Cauchy decrease for successful steps. **Also, κ_m is usually larger for MBDFO.**

Termination

Since we don't know $\|\nabla f(\mathbf{x}_k)\|$, we can't terminate when this gets sufficiently small.

In practice, we terminate the algorithm:

Termination

Since we don't know $\|\nabla f(\mathbf{x}_k)\|$, we can't terminate when this gets sufficiently small.

In practice, we terminate the algorithm:

- After some maximum budget (usually # objective evaluations)
- When Δ_k is sufficiently small

Termination

Since we don't know $\|\nabla f(\mathbf{x}_k)\|$, we can't terminate when this gets sufficiently small.

In practice, we terminate the algorithm:

- After some maximum budget (usually # objective evaluations)
- When Δ_k is sufficiently small

The latter is justified by the following result:

Theorem

We have $\lim_{k \rightarrow \infty} \Delta_k = 0$, and when $\Delta_{k+1} < \Delta_{\min}$ for the first time we have $\|\nabla f(\mathbf{x}_k)\| \leq \mathcal{O}(\Delta_{\min})$.

1. Main MBDFO Algorithm
2. Convergence Theory
3. **Interpolation Model Construction**

Interpolation Models

Our algorithm now has convergence theory, provided we can construct fully linear models at every iteration.

Interpolation Models

Our algorithm now has convergence theory, provided we can construct fully linear models at every iteration.

In practice, we use linear or quadratic interpolation over points near $\mathbf{x}_k$ to define the model m_k .

Interpolation Models

Our algorithm now has convergence theory, provided we can construct fully linear models at every iteration.

In practice, we use linear or quadratic interpolation over points near $\mathbf{x}_k$ to define the model m_k . If the model is a good approximation to a Taylor series, then it is fully linear.

Interpolation Models

Our algorithm now has convergence theory, provided we can construct fully linear models at every iteration.

In practice, we use linear or quadratic interpolation over points near $\mathbf{x}_k$ to define the model m_k . If the model is a good approximation to a Taylor series, then it is fully linear.

Theorem [LR, 2025]

If m_k is a quadratic model such that

$$|m_k(\mathbf{y}) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y} - \mathbf{x}_k)| \leq \kappa \Delta_k^2$$

for all $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, and $\|H_k\| \leq \kappa_H$, then m_k is fully linear with $\kappa_{\text{mf}} = \kappa + \frac{L}{2}$ and $\kappa_{\text{mg}} = 2\kappa + L + 2\kappa_H$.

Linear Interpolation

First, let's build a **linear interpolation model**:

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k).$$

Linear Interpolation

First, let's build a **linear interpolation model**:

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k).$$

We have $p = n + 1$ degrees of freedom, so assume we have interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ where we know the value of f .

Linear Interpolation

First, let's build a **linear interpolation model**:

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k).$$

We have $p = n + 1$ degrees of freedom, so assume we have interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ where we know the value of f .

The unique model satisfying $m_k(\mathbf{y}_t) = f(\mathbf{y}_t)$ is given by the $p \times p$ linear system

$$\begin{bmatrix} 1 & (\mathbf{y}_1 - \mathbf{x}_k)^T \\ \vdots & \vdots \\ 1 & (\mathbf{y}_p - \mathbf{x}_k)^T \end{bmatrix} \begin{bmatrix} c_k \\ \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

Linear Interpolation

First, let's build a **linear interpolation model**:

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k).$$

We have $p = n + 1$ degrees of freedom, so assume we have interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ where we know the value of f .

The unique model satisfying $m_k(\mathbf{y}_t) = f(\mathbf{y}_t)$ is given by the $p \times p$ linear system

$$\begin{bmatrix} 1 & (\mathbf{y}_1 - \mathbf{x}_k)^T \\ \vdots & \vdots \\ 1 & (\mathbf{y}_p - \mathbf{x}_k)^T \end{bmatrix} \begin{bmatrix} c_k \\ \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

We shall see that we want $\|\mathbf{y}_t - \mathbf{x}_k\| = \mathcal{O}(\Delta_k)$.

Linear Interpolation

First, let's build a **linear interpolation model**:

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) = c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k).$$

We have $p = n + 1$ degrees of freedom, so assume we have interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ where we know the value of f .

The unique model satisfying $m_k(\mathbf{y}_t) = f(\mathbf{y}_t)$ is given by the $p \times p$ linear system

$$\begin{bmatrix} 1 & (\mathbf{y}_1 - \mathbf{x}_k)^T \\ \vdots & \vdots \\ 1 & (\mathbf{y}_p - \mathbf{x}_k)^T \end{bmatrix} \begin{bmatrix} c_k \\ \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

We shall see that we want $\|\mathbf{y}_t - \mathbf{x}_k\| = \mathcal{O}(\Delta_k)$. But $\Delta_k \rightarrow 0$, so the matrix becomes increasingly ill-conditioned!

Linear Interpolation

Instead, we rescale the linear system and solve

$$\underbrace{\begin{bmatrix} 1 & \hat{\mathbf{s}}_1^T \\ \vdots & \vdots \\ 1 & \hat{\mathbf{s}}_p^T \end{bmatrix}}_{=:\hat{\mathbf{M}}} \begin{bmatrix} c_k \\ \Delta_k \cdot \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

where $\hat{\mathbf{s}}_t := (\mathbf{y}_t - \mathbf{x}_k) / \Delta_k$.

Linear Interpolation

Instead, we rescale the linear system and solve

$$\underbrace{\begin{bmatrix} 1 & \hat{\mathbf{s}}_1^T \\ \vdots & \vdots \\ 1 & \hat{\mathbf{s}}_p^T \end{bmatrix}}_{=:\hat{M}} \begin{bmatrix} c_k \\ \Delta_k \cdot \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

where $\hat{\mathbf{s}}_t := (\mathbf{y}_t - \mathbf{x}_k) / \Delta_k$.

Is $\hat{M}$ ill-conditioned?

Linear Interpolation

Instead, we rescale the linear system and solve

$$\underbrace{\begin{bmatrix} 1 & \hat{\mathbf{s}}_1^T \\ \vdots & \vdots \\ 1 & \hat{\mathbf{s}}_p^T \end{bmatrix}}_{=:\hat{M}} \begin{bmatrix} c_k \\ \Delta_k \cdot \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

where $\hat{\mathbf{s}}_t := (\mathbf{y}_t - \mathbf{x}_k) / \Delta_k$.

Is $\hat{M}$ ill-conditioned? Choosing $\|\mathbf{y}_t - \mathbf{x}_k\| = \mathcal{O}(\Delta_k)$ ensures $\|\hat{M}\|$ is not too large, but $\|\hat{M}^{-1}\|$ may still be large!

Linear Interpolation

Instead, we rescale the linear system and solve

$$\underbrace{\begin{bmatrix} 1 & \hat{\mathbf{s}}_1^T \\ \vdots & \vdots \\ 1 & \hat{\mathbf{s}}_p^T \end{bmatrix}}_{=:\hat{M}} \begin{bmatrix} c_k \\ \Delta_k \cdot \mathbf{g}_k \end{bmatrix} = \begin{bmatrix} f(\mathbf{y}_1) \\ \vdots \\ f(\mathbf{y}_p) \end{bmatrix}$$

where $\hat{\mathbf{s}}_t := (\mathbf{y}_t - \mathbf{x}_k)/\Delta_k$.

Is $\hat{M}$ ill-conditioned? Choosing $\|\mathbf{y}_t - \mathbf{x}_k\| = \mathcal{O}(\Delta_k)$ ensures $\|\hat{M}\|$ is not too large, but $\|\hat{M}^{-1}\|$ may still be large!

Our fully linear error bounds reflect this...

Theorem

If $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ for all t and $\hat{M}$ is invertible, then the model is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(L\sqrt{n}\beta^2\|\hat{M}^{-1}\|_\infty)$.

Theorem

If $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ for all t and $\hat{M}$ is invertible, then the model is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(L\sqrt{n}\beta^2\|\hat{M}^{-1}\|_\infty)$.

Proof: Since $f(\mathbf{y}_t) = m_k(\mathbf{y}_t)$, Taylor error bounds give

$$|m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)| \leq \frac{L}{2}\|\mathbf{y}_t - \mathbf{x}_k\|^2 \leq \frac{L}{2}\beta^2\Delta_k^2.$$

Theorem

If $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ for all t and $\hat{M}$ is invertible, then the model is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(L\sqrt{n}\beta^2\|\hat{M}^{-1}\|_\infty)$.

Proof: Since $f(\mathbf{y}_t) = m_k(\mathbf{y}_t)$, Taylor error bounds give

$$|m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)| \leq \frac{L}{2}\|\mathbf{y}_t - \mathbf{x}_k\|^2 \leq \frac{L}{2}\beta^2\Delta_k^2.$$

Now pick any $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, with $\hat{\mathbf{s}} := (\mathbf{y} - \mathbf{x}_k)/\Delta_k \in B(\mathbf{0}, 1)$.

Theorem

If $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ for all t and $\hat{M}$ is invertible, then the model is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(L\sqrt{n}\beta^2\|\hat{M}^{-1}\|_\infty)$.

Proof: Since $f(\mathbf{y}_t) = m_k(\mathbf{y}_t)$, Taylor error bounds give

$$|m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)| \leq \frac{L}{2}\|\mathbf{y}_t - \mathbf{x}_k\|^2 \leq \frac{L}{2}\beta^2\Delta_k^2.$$

Now pick any $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, with $\hat{\mathbf{s}} := (\mathbf{y} - \mathbf{x}_k)/\Delta_k \in B(\mathbf{0}, 1)$. Since $\hat{M}$ is invertible, there exists a unique solution to

$$\hat{M}^T \mathbf{v} = \begin{bmatrix} 1 \\ \hat{\mathbf{s}} \end{bmatrix} \quad \text{or} \quad \sum_{t=1}^p v_t \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} = \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix}$$

Theorem

If $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ for all t and $\hat{M}$ is invertible, then the model is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(L\sqrt{n}\beta^2\|\hat{M}^{-1}\|_\infty)$.

Proof: Since $f(\mathbf{y}_t) = m_k(\mathbf{y}_t)$, Taylor error bounds give

$$|m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)| \leq \frac{L}{2}\|\mathbf{y}_t - \mathbf{x}_k\|^2 \leq \frac{L}{2}\beta^2\Delta_k^2.$$

Now pick any $\mathbf{y} \in B(\mathbf{x}_k, \Delta_k)$, with $\hat{\mathbf{s}} := (\mathbf{y} - \mathbf{x}_k)/\Delta_k \in B(\mathbf{0}, 1)$. Since $\hat{M}$ is invertible, there exists a unique solution to

$$\hat{M}^T \mathbf{v} = \begin{bmatrix} 1 \\ \hat{\mathbf{s}} \end{bmatrix} \quad \text{or} \quad \sum_{t=1}^p v_t \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} = \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix}$$

Note $\|\mathbf{v}\|_1 \leq \|\hat{M}^{-T}\|_1(1 + \|\hat{\mathbf{s}}\|_1) \leq \|\hat{M}^{-1}\|_\infty(1 + \sqrt{n})$, using $\|\hat{\mathbf{s}}\| \leq 1$.

Linear Interpolation

$$\begin{aligned} & |m_k(\mathbf{y}) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y} - \mathbf{x}_k)| \\ &= \left| \begin{bmatrix} c \\ \mathbf{g}_k \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix} - \begin{bmatrix} f(\mathbf{x}_k) \\ \nabla f(\mathbf{x}_k) \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix} \right| \\ &= \left| \sum_{t=1}^p v_t \left(\begin{bmatrix} c \\ \mathbf{g}_k \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} - \begin{bmatrix} f(\mathbf{x}_k) \\ \nabla f(\mathbf{x}_k) \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} \right) \right| \\ &= \left| \sum_{t=1}^p v_t [m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)] \right| \\ &\leq \frac{L}{2} \beta^2 \Delta_k^2 \cdot \|\mathbf{v}\|_1 \end{aligned}$$

Linear Interpolation

$$\begin{aligned} & |m_k(\mathbf{y}) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y} - \mathbf{x}_k)| \\ &= \left| \begin{bmatrix} c \\ \mathbf{g}_k \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix} - \begin{bmatrix} f(\mathbf{x}_k) \\ \nabla f(\mathbf{x}_k) \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y} - \mathbf{x}_k \end{bmatrix} \right| \\ &= \left| \sum_{t=1}^p v_t \left(\begin{bmatrix} c \\ \mathbf{g}_k \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} - \begin{bmatrix} f(\mathbf{x}_k) \\ \nabla f(\mathbf{x}_k) \end{bmatrix}^T \begin{bmatrix} 1 \\ \mathbf{y}_t - \mathbf{x}_k \end{bmatrix} \right) \right| \\ &= \left| \sum_{t=1}^p v_t [m_k(\mathbf{y}_t) - f(\mathbf{x}_k) - \nabla f(\mathbf{x}_k)^T(\mathbf{y}_t - \mathbf{x}_k)] \right| \\ &\leq \frac{L}{2} \beta^2 \Delta_k^2 \cdot \|\mathbf{v}\|_1 \end{aligned}$$

Then apply bound on $\|\mathbf{v}\|_1$ and “Taylor approximation implies fully linear” theorem.

Example Sets

For example, finite differencing with stepsize $h = \Delta_k$ is equivalent to linear interpolation with points $\{\mathbf{x}_k, \mathbf{x}_k + \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k + \Delta_k \mathbf{e}_n\}$.

Example Sets

For example, finite differencing with stepsize $h = \Delta_k$ is equivalent to linear interpolation with points $\{\mathbf{x}_k, \mathbf{x}_k + \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k + \Delta_k \mathbf{e}_n\}$. This gives $\beta = 1$ and $\|\hat{M}^{-1}\|_\infty = 2$, and so $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n})$.

Example Sets

For example, finite differencing with stepsize $h = \Delta_k$ is equivalent to linear interpolation with points $\{\mathbf{x}_k, \mathbf{x}_k + \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k + \Delta_k \mathbf{e}_n\}$. This gives $\beta = 1$ and $\|\hat{M}^{-1}\|_\infty = 2$, and so $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n})$.

Corollary

Using this interpolation set, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(n\epsilon^{-2})$ iterations or $\mathcal{O}(n^2\epsilon^{-2})$ evaluations.

Example Sets

For example, finite differencing with stepsize $h = \Delta_k$ is equivalent to linear interpolation with points $\{\mathbf{x}_k, \mathbf{x}_k + \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k + \Delta_k \mathbf{e}_n\}$. This gives $\beta = 1$ and $\|\hat{M}^{-1}\|_\infty = 2$, and so $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n})$.

Corollary

Using this interpolation set, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(n\epsilon^{-2})$ iterations or $\mathcal{O}(n^2\epsilon^{-2})$ evaluations.

By taking perturbations of size $\Delta_k/\sqrt{n}$, we can improve this (in terms of explicit n dependence) to a (likely) optimal result.

Example Sets

For example, finite differencing with stepsize $h = \Delta_k$ is equivalent to linear interpolation with points $\{\mathbf{x}_k, \mathbf{x}_k + \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k + \Delta_k \mathbf{e}_n\}$. This gives $\beta = 1$ and $\|\hat{M}^{-1}\|_\infty = 2$, and so $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n})$.

Corollary

Using this interpolation set, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(n\epsilon^{-2})$ iterations or $\mathcal{O}(n^2\epsilon^{-2})$ evaluations.

By taking perturbations of size $\Delta_k/\sqrt{n}$, we can improve this (in terms of explicit n dependence) to a (likely) optimal result.

Corollary

Using coordinate perturbations of size $\Delta_k/\sqrt{n}$, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(\epsilon^{-2})$ iterations or $\mathcal{O}(n\epsilon^{-2})$ evaluations.

Quadratic Interpolation

In practice, linear models are not particularly good*, but quadratic models are generally quite effective.

* exc., e.g. [Cartis & LR, 2019]

Quadratic Interpolation

In practice, linear models are not particularly good*, but quadratic models are generally quite effective.

* exc., e.g. [Cartis & LR, 2019]

To build a quadratic model

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k) + \frac{1}{2}(\mathbf{y} - \mathbf{x}_k)^T H_k(\mathbf{y} - \mathbf{x}_k)$$

requires $p = 1 + n + n(n+1)/2 = (n+1)(n+2)/2$ degrees of freedom/interpolation points.

Quadratic Interpolation

In practice, linear models are not particularly good*, but quadratic models are generally quite effective.

* exc., e.g. [Cartis & LR, 2019]

To build a quadratic model

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k) + \frac{1}{2}(\mathbf{y} - \mathbf{x}_k)^T H_k(\mathbf{y} - \mathbf{x}_k)$$

requires $p = 1 + n + n(n+1)/2 = (n+1)(n+2)/2$ degrees of freedom/interpolation points.

Even for moderate n , this can mean many interpolation points, and $\mathcal{O}(p^3) = \mathcal{O}(n^6)$ cost of solving linear systems.

Quadratic Interpolation

In practice, linear models are not particularly good*, but quadratic models are generally quite effective.

* exc., e.g. [Cartis & LR, 2019]

To build a quadratic model

$$f(\mathbf{y}) \approx m_k(\mathbf{y}) := c_k + \mathbf{g}_k^T(\mathbf{y} - \mathbf{x}_k) + \frac{1}{2}(\mathbf{y} - \mathbf{x}_k)^T H_k(\mathbf{y} - \mathbf{x}_k)$$

requires $p = 1 + n + n(n+1)/2 = (n+1)(n+2)/2$ degrees of freedom/interpolation points.

Even for moderate n , this can mean many interpolation points, and $\mathcal{O}(p^3) = \mathcal{O}(n^6)$ cost of solving linear systems.

However, we do get higher accuracy (**fully quadratic** $\approx$ 2nd order Taylor accurate) which allows convergence to second-order optimal points ($\|\nabla f(\mathbf{x})\| = 0$ and $\nabla^2 f(\mathbf{x}) \succeq 0$).

Quadratic Interpolation

A practically successful middle ground is **underdetermined quadratic interpolation** (i.e. $n + 1 < p < (n + 1)(n + 2)/2$ interpolation points).

Quadratic Interpolation

A practically successful middle ground is **underdetermined quadratic interpolation** (i.e. $n + 1 < p < (n + 1)(n + 2)/2$ interpolation points).

Interpolation system is underdetermined — infinitely many possible models satisfy the interpolation conditions. Two common choices, both with theory:

Quadratic Interpolation

A practically successful middle ground is **underdetermined quadratic interpolation** (i.e. $n + 1 < p < (n + 1)(n + 2)/2$ interpolation points).

Interpolation system is underdetermined — infinitely many possible models satisfy the interpolation conditions. Two common choices, both with theory:

- Minimum norm solution of linear system [Conn et al., 2008]
- Minimum Frobenius norm (MFN) interpolation [Powell, 2004]

Quadratic Interpolation

A practically successful middle ground is **underdetermined quadratic interpolation** (i.e. $n + 1 < p < (n + 1)(n + 2)/2$ interpolation points).

Interpolation system is underdetermined — infinitely many possible models satisfy the interpolation conditions. Two common choices, both with theory:

- Minimum norm solution of linear system [Conn et al., 2008]
- Minimum Frobenius norm (MFN) interpolation [Powell, 2004]

The most widely used MBDFO software use MFN interpolation.

Minimum Frobenius Norm Interpolation

For **MFN**, given interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ (with $n + 1 \leq p \leq (n + 1)(n + 2)/2$), derive a quadratic model via

$$\min_{\mathbf{c}_k, \mathbf{g}_k, H_k} \|H_k\|_F^2 \quad \text{s.t.} \quad H_k = H_k^T \text{ and } m_k(\mathbf{y}_t) = f(\mathbf{y}_t) \forall t = 1, \dots, p$$

Minimum Frobenius Norm Interpolation

For **MFN**, given interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ (with $n + 1 \leq p \leq (n + 1)(n + 2)/2$), derive a quadratic model via

$$\min_{c_k, \mathbf{g}_k, H_k} \|H_k\|_F^2 \quad \text{s.t.} \quad H_k = H_k^T \text{ and } m_k(\mathbf{y}_t) = f(\mathbf{y}_t) \forall t = 1, \dots, p$$

In practice, minimum change Hessian, $\min \|H_k - H_{k-1}\|_F^2$ also works well, but no theory.

Minimum Frobenius Norm Interpolation

For **MFN**, given interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ (with $n + 1 \leq p \leq (n + 1)(n + 2)/2$), derive a quadratic model via

$$\min_{c_k, \mathbf{g}_k, H_k} \|H_k\|_F^2 \quad \text{s.t.} \quad H_k = H_k^T \text{ and } m_k(\mathbf{y}_t) = f(\mathbf{y}_t) \forall t = 1, \dots, p$$

In practice, minimum change Hessian, $\min \|H_k - H_{k-1}\|_F^2$ also works well, but no theory. Idea motivated by quasi-Newton updating (min change in Hessian s.t. secant equation).

Minimum Frobenius Norm Interpolation

For **MFN**, given interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ (with $n + 1 \leq p \leq (n + 1)(n + 2)/2$), derive a quadratic model via

$$\min_{c_k, \mathbf{g}_k, H_k} \|H_k\|_F^2 \quad \text{s.t.} \quad H_k = H_k^T \text{ and } m_k(\mathbf{y}_t) = f(\mathbf{y}_t) \forall t = 1, \dots, p$$

In practice, minimum change Hessian, $\min \|H_k - H_{k-1}\|_F^2$ also works well, but no theory. Idea motivated by quasi-Newton updating (min change in Hessian s.t. secant equation).

This problem is a convex, equality-constrained QP, so can be solved via a linear system (size $p + n + 1$, symmetric saddle point system).

Minimum Frobenius Norm Interpolation

For **MFN**, given interpolation points $\mathbf{y}_1, \dots, \mathbf{y}_p$ (with $n + 1 \leq p \leq (n + 1)(n + 2)/2$), derive a quadratic model via

$$\min_{c_k, \mathbf{g}_k, H_k} \|H_k\|_F^2 \quad \text{s.t.} \quad H_k = H_k^T \text{ and } m_k(\mathbf{y}_t) = f(\mathbf{y}_t) \forall t = 1, \dots, p$$

In practice, minimum change Hessian, $\min \|H_k - H_{k-1}\|_F^2$ also works well, but no theory. Idea motivated by quasi-Newton updating (min change in Hessian s.t. secant equation).

This problem is a convex, equality-constrained QP, so can be solved via a linear system (size $p + n + 1$, symmetric saddle point system).

If $p = n + 1$ get linear models, if $p = (n + 1)(n + 2)/2$ get (unique) fully quadratic interpolant.

Minimum Frobenius Norm Interpolation

From saddle point system theory, we get

Lemma [Benzi et al., 2005]

If the MFN interpolation matrix $\hat{F}$ (scaled via $(\mathbf{y}_t - \mathbf{x}_k)/\Delta_k$) is invertible, then the equivalent linear regression system $\hat{M} \in \mathbb{R}^{p \times (n+1)}$ is full column rank.

Minimum Frobenius Norm Interpolation

From saddle point system theory, we get

Lemma [Benzi et al., 2005]

If the MFN interpolation matrix $\hat{F}$ (scaled via $(\mathbf{y}_t - \mathbf{x}_k)/\Delta_k$) is invertible, then the equivalent linear regression system $\hat{M} \in \mathbb{R}^{p \times (n+1)}$ is full column rank.

The fully linear bound is:

Theorem [LR, 2025], originally [Conn et al., 2008]

If all $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ and $\hat{F}$ is invertible, then the MFN model:

Minimum Frobenius Norm Interpolation

From saddle point system theory, we get

Lemma [Benzi et al., 2005]

If the MFN interpolation matrix $\hat{F}$ (scaled via $(\mathbf{y}_t - \mathbf{x}_k)/\Delta_k$) is invertible, then the equivalent linear regression system $\hat{M} \in \mathbb{R}^{p \times (n+1)}$ is full column rank.

The fully linear bound is:

Theorem [LR, 2025], originally [Conn et al., 2008]

If all $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ and $\hat{F}$ is invertible, then the MFN model:

- Has bounded Hessian, $\|H_k\| \leq \kappa_H = \mathcal{O}(Lp\beta^4 \|\hat{F}^{-1}\|_\infty)$.

Minimum Frobenius Norm Interpolation

From saddle point system theory, we get

Lemma [Benzi et al., 2005]

If the MFN interpolation matrix $\hat{F}$ (scaled via $(\mathbf{y}_t - \mathbf{x}_k)/\Delta_k$) is invertible, then the equivalent linear regression system $\hat{M} \in \mathbb{R}^{p \times (n+1)}$ is full column rank.

The fully linear bound is:

Theorem [LR, 2025], originally [Conn et al., 2008]

If all $\|\mathbf{y}_t - \mathbf{x}_k\| \leq \beta \Delta_k$ and $\hat{F}$ is invertible, then the MFN model:

- Has bounded Hessian, $\|H_k\| \leq \kappa_H = \mathcal{O}(Lp\beta^4 \|\hat{F}^{-1}\|_\infty)$.
- Is fully linear with $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}((L + \kappa_H)\sqrt{n}\beta^2 \|\hat{M}^\dagger\|_\infty)$.

Minimum Frobenius Norm Interpolation

As a simple example, if we do MFN with $p = 2n + 1$ (a popular choice) and coordinate perturbations, $\{\mathbf{x}_k, \mathbf{x}_k \pm \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k \pm \Delta_k \mathbf{e}_n\}$, then we get central difference gradients and diagonal H_k (3-point stencil for second derivatives).

Minimum Frobenius Norm Interpolation

As a simple example, if we do MFN with $p = 2n + 1$ (a popular choice) and coordinate perturbations, $\{\mathbf{x}_k, \mathbf{x}_k \pm \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k \pm \Delta_k \mathbf{e}_n\}$, then we get central difference gradients and diagonal H_k (3-point stencil for second derivatives).

The above result gives $\kappa_H = \mathcal{O}(n^2)$ and $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n} \kappa_H)$, but a more refined analysis gives $\kappa_H = \mathcal{O}(1)$.

Minimum Frobenius Norm Interpolation

As a simple example, if we do MFN with $p = 2n + 1$ (a popular choice) and coordinate perturbations, $\{\mathbf{x}_k, \mathbf{x}_k \pm \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k \pm \Delta_k \mathbf{e}_n\}$, then we get central difference gradients and diagonal H_k (3-point stencil for second derivatives).

The above result gives $\kappa_H = \mathcal{O}(n^2)$ and $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n} \kappa_H)$, but a more refined analysis gives $\kappa_H = \mathcal{O}(1)$.

Corollary

Using this interpolation set, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(n\epsilon^{-2})$ iterations or $\mathcal{O}(n^2\epsilon^{-2})$ evaluations.

Minimum Frobenius Norm Interpolation

As a simple example, if we do MFN with $p = 2n + 1$ (a popular choice) and coordinate perturbations, $\{\mathbf{x}_k, \mathbf{x}_k \pm \Delta_k \mathbf{e}_1, \dots, \mathbf{x}_k \pm \Delta_k \mathbf{e}_n\}$, then we get central difference gradients and diagonal H_k (3-point stencil for second derivatives).

The above result gives $\kappa_H = \mathcal{O}(n^2)$ and $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(\sqrt{n} \kappa_H)$, but a more refined analysis gives $\kappa_H = \mathcal{O}(1)$.

Corollary

Using this interpolation set, we get $\|\nabla f(\mathbf{x}_k)\| < \epsilon$ after at most $\mathcal{O}(n\epsilon^{-2})$ iterations or $\mathcal{O}(n^2\epsilon^{-2})$ evaluations.

Similar finite differencing-type choices with perturbations scaled by n can again recover $\kappa_{\text{mf}}, \kappa_{\text{mg}} = \mathcal{O}(1)$, i.e. $\mathcal{O}(\epsilon^{-2})$ iterations and $\mathcal{O}(n\epsilon^{-2})$ evaluations.

Next steps

This approach to model construction is simple and gives nice theory, but requires resampling all interpolation points at every iteration (wasteful when evaluations are expensive).

Next steps

This approach to model construction is simple and gives nice theory, but requires resampling all interpolation points at every iteration (wasteful when evaluations are expensive).

Tomorrow, we look at more efficient interpolation set management — reusing many/most interpolation points from the previous iteration.

Next steps

This approach to model construction is simple and gives nice theory, but requires resampling all interpolation points at every iteration (wasteful when evaluations are expensive).

Tomorrow, we look at more efficient interpolation set management — reusing many/most interpolation points from the previous iteration.

We will also look at extensions of this algorithmic framework (e.g. constrained optimization).

M. BENZI, G. H. GOLUB, AND J. LIESEN, *Numerical solution of saddle point problems*, Acta Numerica, 14 (2005), pp. 1–137.

C. CARTIS, N. I. M. GOULD, AND P. L. TOINT, *Evaluation Complexity of Algorithms for Nonconvex Optimization: Theory, Computation and Perspectives*, no. 30 in MOS-SIAM Series on Optimization, MOS/SIAM, Philadelphia, 2022.

C. CARTIS AND L. ROBERTS, *A derivative-free Gauss–Newton method*, Mathematical Programming Computation, 11 (2019), pp. 631–674.

A. R. CONN, N. I. M. GOULD, AND P. L. TOINT, *Trust-Region Methods*, vol. 1 of MPS-SIAM Series on Optimization, MPS/SIAM, Philadelphia, 2000.

A. R. CONN, K. SCHEINBERG, AND L. N. VICENTE, *Geometry of sample sets in derivative-free optimization: Polynomial regression and underdetermined interpolation*, IMA Journal of Numerical Analysis, 28 (2008), pp. 721–748.

A. R. CONN, K. SCHEINBERG, AND L. N. VICENTE, *Global convergence of general derivative-free trust-region algorithms to first- and second-order critical points*, SIAM Journal on Optimization, 20 (2009), pp. 387–415.

- A. R. CONN, K. SCHEINBERG, AND L. N. VICENTE, *Introduction to Derivative-Free Optimization*, vol. 8 of MPS-SIAM Series on Optimization, MPS/SIAM, Philadelphia, 2009.
- M. J. D. POWELL, *Least Frobenius norm updating of quadratic models that satisfy interpolation conditions*, Mathematical Programming, 100 (2004), pp. 183–215.
- L. ROBERTS, *Introduction to model-based derivative-free optimization*.
arXiv preprint [arXiv:2510.04473](https://arxiv.org/abs/2510.04473), 2025.