

Nonconvex optimization algorithms with complexity guarantees

Clément W. Royer

Soutenance d'Habilitation à Diriger des Recherches

September 22, 2025

Dauphine | PSL  **LAMSADE**
UNIVERSITÉ PARIS UMR 7243



- **PhD** (2016)
Université de Toulouse.
- **Postdoc** (2016 - 2019)
University of Wisconsin-Madison.
- **Faculty** (2019 -)
Université Paris Dauphine-PSL.

Collaborators since the PhD

Co-authors (including **junior researchers** and **group members**)

- **North America** Albert S. Berahas, Kwassi Joseph Dzahini, Michael J. O'Neill, Akwum Onwunta, Daniel P. Robinson, Oumaima Sohab, Stephen J. Wright (USA); Youssef Diouane, Warren Hare, Gabriel Jarry-Bolduc (Canada).
- **Dauphine** Rémi Chan-Renous-Legoubin, Yann Chevaleyre, Denis Cornaz, Florentin Goyens, Sébastien Kerleau, Laurent Meunier.
- **France** Jérémy Rapin, Olivier Teytaud.
- **Europe** Vladimir Kunc, Vyacheslav Kungurtsev (Czech Republic); Francesco Rinaldi, Damiano Zeffiro (Italy).
- **Africa** El Houcine Bergou (Morocco).
- **Oceania** Lindon Roberts (Australia).

Collaborators since the PhD

Co-authors (including **junior researchers** and **group members**)

- **North America** Albert S. Berahas, Kwassi Joseph Dzahini, Michael J. O'Neill, Akwum Onwunta, Daniel P. Robinson, Oumaima Sohab, Stephen J. Wright (USA); Youssef Diouane, Warren Hare, Gabriel Jarry-Bolduc (Canada).
- **Dauphine** Rémi Chan-Renous-Legoubin, Yann Chevaleyre, Denis Cornaz, Florentin Goyens, Sébastien Kerleau, Laurent Meunier.
- **France** Jérémy Rapin, Olivier Teytaud.
- **Europe** Vladimir Kunc, Vyacheslav Kungurtsev (Czech Republic); Francesco Rinaldi, Damiano Zeffiro (Italy).
- **Africa** El Houcine Bergou (Morocco).
- **Oceania** Lindon Roberts (Australia).

Other collaborations

Alexandre Allauzen, Antonin Chambolle, Irène Waldspurger, Florian Yger.



The group @ EUROPT 2025

- Bastien Cavarretta (PhD 2024 -)
- Sébastien Kerleau
(PhD → November 2025)
- Iskander Legheraba
(PhD → September 2025)
- Annette Dumas
(Postdoc → September 2025).



The group @ EUROPT 2025

- Bastien Cavarretta (PhD 2024 -)
- Sébastien Kerleau (PhD → November 2025)
- Iskander Legheraba (PhD → September 2025)
- Annette Dumas (Postdoc → September 2025).

Other group members

- Gaetano Agazzotti (Master, 2025 -)
- Florentin Goyens (Postdoc, 2022 - 2024)
- Rémi Chan-Renous Legoubin (Master, 2021).
- Thomas Georges, Marc Kaspar, Christian Kayo, Eloi Martin, Luca Solbiati (Master, [2021,2024]).

Funding sources



PROGRAMME
DE RECHERCHE
INTELLIGENCE
ARTIFICIELLE

PR[AI]RIE

PaRis Artificial Intelligence Research Institute

Numerical optimization

- Design algorithms to minimize a (continuous) function.

$$\underset{\boldsymbol{x} \in \mathbb{R}^n}{\text{minimize}} \ f(\boldsymbol{x}).$$

- Prove theoretical guarantees.
- Validate practical performance.

Numerical optimization

- Design algorithms to minimize a (continuous) function.

$$\underset{\boldsymbol{x} \in \mathbb{R}^n}{\text{minimize}} \ f(\boldsymbol{x}).$$

- Prove theoretical guarantees.
- Validate practical performance.

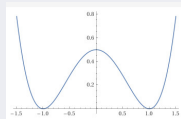
Recurring theme: Complexity analysis in nonconvex settings.

Nonconvex setting

Main problem

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \ f(\mathbf{x})$$

f smooth ($\mathcal{C}^1$ or $\mathcal{C}^2$) and **nonconvex**.

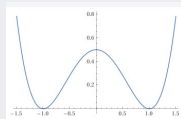


Main problem

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \ f(\mathbf{x})$$

f smooth ($\mathcal{C}^1$ or $\mathcal{C}^2$) and **nonconvex**.

- NP-hard to solve, even locally.
- Settle to find approximate stationary points.

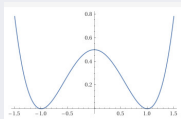


Main problem

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \ f(\mathbf{x})$$

f smooth ($\mathcal{C}^1$ or $\mathcal{C}^2$) and **nonconvex**.

- NP-hard to solve, even locally.
- Settle to find approximate stationary points.



Approximate stationary points

First-order For $\epsilon \in (0, 1)$,

$$\|\nabla f(\mathbf{x})\| \leq \epsilon.$$

Second-order For $\epsilon, \epsilon_H \in (0, 1)^2$,

$$\|\nabla f(\mathbf{x})\| \leq \epsilon, \quad \nabla^2 f(\mathbf{x}) \succeq -\epsilon_H \mathbf{I}.$$

Problem minimize $x \in \mathbb{R}^n$ $f(x)$.

Solving the problem

- **Algorithm** Iterative process $\{x_0, x_1, \dots\}$.
- **Cost** Iterations, evaluations of f and derivatives, etc.

Problem minimize $\mathbf{x} \in \mathbb{R}^n$ $f(\mathbf{x})$.

Solving the problem

- **Algorithm** Iterative process $\{\mathbf{x}_0, \mathbf{x}_1, \dots\}$.
- **Cost** Iterations, evaluations of f and derivatives, etc.

Complexity

Given ϵ, ϵ_H and an algorithm $\{\mathbf{x}_j\}_j$, find worst-case cost of the algorithm to reach $\mathbf{x}_J$ such that

- 1 $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$, or
- 2 $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon, \quad \nabla^2 f(\mathbf{x}_J) \succeq -\epsilon_H \mathbf{I}.$

Canonical example: Gradient descent

Input $x_0 \in \mathbb{R}^n$.

For $j = 0, 1, 2, \dots$

- Compute $\alpha_j > 0$.
- Set $\mathbf{x}_{j+1} = \mathbf{x}_j - \alpha_j \nabla f(\mathbf{x}_j)$.

Canonical example: Gradient descent

Input $x_0 \in \mathbb{R}^n$.

For $j = 0, 1, 2, \dots$

- Compute $\alpha_j > 0$.
 - Set $\mathbf{x}_{j+1} = \mathbf{x}_j - \alpha_j \nabla f(\mathbf{x}_j)$.
-
- Stepsize α_j : Chosen constant or with line search.
 - **Cost** Iterations/Calls to ∇f + Calls to f (Line search).

Complexity of gradient descent

Gradient descent $\mathbf{x}_{j+1} = \mathbf{x}_j - \alpha_j \nabla f(\mathbf{x}_j)$.

Goal $\|\nabla f(\mathbf{x})\| \leq \epsilon$.

Complexity theorem

Assume ∇f Lipschitz continuous, f bounded below. Then the method runs in at most

$$\mathcal{O}(\epsilon^{-2})$$

iterations/calls to ∇f .

Complexity of gradient descent

Gradient descent $\mathbf{x}_{j+1} = \mathbf{x}_j - \alpha_j \nabla f(\mathbf{x}_j)$.

Goal $\|\nabla f(\mathbf{x})\| \leq \epsilon$.

Complexity theorem

Assume ∇f Lipschitz continuous, f bounded below. Then the method runs in at most

$$\mathcal{O}(\epsilon^{-2})$$

iterations/calls to ∇f .

- Holds for fixed/adaptive α_j choices.
- With line search, cost in calls to f is also $\mathcal{O}(\epsilon^{-2})$.
- Sharp bounds in terms of ϵ .

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
- Study special classes of nonconvex problems.

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
- Study special classes of nonconvex problems.

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
- Checks to detect nonconvex behavior.

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
- Study special classes of nonconvex problems.

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
- Checks to detect nonconvex behavior.

- **Derivative-free methods**

- Improve complexity with randomness.
- Use geometry to handle difficult settings.

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
- Study special classes of nonconvex problems.

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
- Checks to detect nonconvex behavior.

- **Derivative-free methods**

- Improve complexity with randomness.
- Use geometry to handle difficult settings.

This presentation: Focus on contributions with students/postdocs.

- 1 Newton-type methods
- 2 Conjugate gradient methods
- 3 Direct-search methods
- 4 Perspectives

- 1 Newton-type methods
- 2 Conjugate gradient methods
- 3 Direct-search methods
- 4 Perspectives

$$\underset{x \in \mathbb{R}^n}{\text{minimize}} \ f(x)$$

$f \in \mathcal{C}^2$ bounded below and nonconvex.

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} f(\mathbf{x})$$

$f \in \mathcal{C}^2$ bounded below and nonconvex.

Goal: Reach (ϵ, ϵ_H) -point

$$\|\nabla f(\mathbf{x})\| \leq \epsilon \quad \text{and} \quad \nabla^2 f(\mathbf{x}) \succeq -\epsilon_H \mathbf{I}.$$

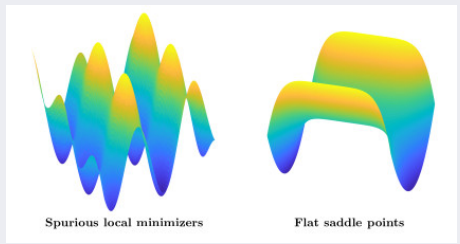
- For convex functions: Second condition always true $\Rightarrow$ Close to a global minimum!
- For nonconvex functions: Depends on the function **landscape**.

Goal Reach (ϵ, ϵ_H) -points

Bad instances

(ϵ, ϵ_H) -points can be close to

- Local, non-global minima.
- High-order saddle points.



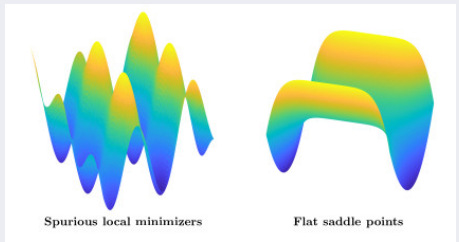
Landscape in nonconvex optimization

Goal Reach (ϵ, ϵ_H) -points

Bad instances

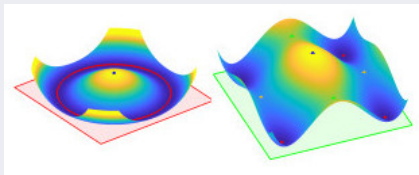
(ϵ, ϵ_H) -points can be close to

- Local, non-global minima.
- High-order saddle points.



Nice instances

- (ϵ, ϵ_H) -points close to local minima.
- All local minima are global.



Figures: J. Wright and Y. Ma, *High-Dimensional Data Analysis with Low-Dimensional Models*, 2022.

Nicest case in optimization: Strongly convex functions

- $f : \mathbb{R}^n \rightarrow \mathbb{R}$ $\mathcal{C}^2$ is μ -strongly convex if

$$\nabla^2 f(\mathbf{x}) \succeq \mu \mathbf{I}.$$

- Unique minimum $\mathbf{x}^* \in \mathbb{R}^n$.

One consequence of μ -strong convexity

For any $\gamma > 0$,

$$\|\nabla f(\mathbf{x})\| \leq \gamma \quad \Rightarrow \quad \|\mathbf{x} - \mathbf{x}^*\| \leq \frac{2\gamma}{\mu}.$$

Definition (Adapted from Ge et al '17)

f is $(\gamma, \lambda, \mu, \delta)$ -strict saddle with $\gamma, \lambda, \mu, \delta > 0$ if at any $\mathbf{x} \in \mathbb{R}^n$, one of the above holds:

- ❶ $\|\nabla f(\mathbf{x})\| \geq \gamma$;
- ❷ $\lambda_{\min}(\nabla^2 f(\mathbf{x})) \leq -\beta$;
- ❸ There exists $\mathbf{x}^*$ local minimum of f such that

$$\|\mathbf{x} - \mathbf{x}^*\| \leq \delta \quad \text{and} \quad \nabla^2 f(\mathbf{y}) \succeq \gamma \mathbf{I} \succ 0 \quad \forall \mathbf{y}, \|\mathbf{y} - \mathbf{x}^*\| \leq 2\delta.$$

Definition (Adapted from Ge et al '17)

f is $(\gamma, \lambda, \mu, \delta)$ -strict saddle with $\gamma, \lambda, \mu, \delta > 0$ if at any $\mathbf{x} \in \mathbb{R}^n$, one of the above holds:

- 1 $\|\nabla f(\mathbf{x})\| \geq \gamma$;
- 2 $\lambda_{\min}(\nabla^2 f(\mathbf{x})) \leq -\beta$;
- 3 There exists $\mathbf{x}^*$ local minimum of f such that

$$\|\mathbf{x} - \mathbf{x}^*\| \leq \delta \quad \text{and} \quad \nabla^2 f(\mathbf{y}) \succeq \gamma \mathbf{I} \succ 0 \quad \forall \mathbf{y}, \|\mathbf{y} - \mathbf{x}^*\| \leq 2\delta.$$

- Similar in spirit to phase/regional complexity frameworks.
- Includes strongly convex as a special case.

Trust region for minimize $_{x \in \mathbb{R}^n} f(\boldsymbol{x})$

Inputs $\boldsymbol{x}_0 \in \mathbb{R}^n$, $\Delta_0 > 0$, $\eta > 0$.

For $j = 0, 1, 2, \dots$

Trust region for minimize $x \in \mathbb{R}^n \ f(x)$

Inputs $x_0 \in \mathbb{R}^n$, $\Delta_0 > 0$, $\eta > 0$.

For $j = 0, 1, 2, \dots$

- 1 Define $m_j(x_j + s) := \nabla f(x_j)^T s + \frac{1}{2} s^T \nabla^2 f(x_j) s$ and compute

$$s_j \in \underset{\|s\| \leq \Delta_j}{\operatorname{argmin}} m_j(x_j + s).$$

Trust region for minimize $x \in \mathbb{R}^n \ f(x)$

Inputs $x_0 \in \mathbb{R}^n$, $\Delta_0 > 0$, $\eta > 0$.

For $j = 0, 1, 2, \dots$

- 1 Define $m_j(x_j + s) := \nabla f(x_j)^T s + \frac{1}{2} s^T \nabla^2 f(x_j) s$ and compute

$$s_j \in \underset{\|s\| \leq \Delta_j}{\operatorname{argmin}} m_j(x_j + s).$$

- 2 Compute $\rho_j = \frac{f(x_j) - f(x_j + s_j)}{m_j(x_j) - m_j(x_j + s_j)}$.

Trust region for minimize $x \in \mathbb{R}^n \ f(x)$

Inputs $x_0 \in \mathbb{R}^n$, $\Delta_0 > 0$, $\eta > 0$.

For $j = 0, 1, 2, \dots$

- 1 Define $m_j(x_j + s) := \nabla f(x_j)^T s + \frac{1}{2} s^T \nabla^2 f(x_j) s$ and compute

$$s_j \in \underset{\|s\| \leq \Delta_j}{\operatorname{argmin}} m_j(x_j + s).$$

- 2 Compute $\rho_j = \frac{f(x_j) - f(x_j + s_j)}{m_j(x_j) - m_j(x_j + s_j)}$.
- 3 If $\rho_j \geq \eta$, set $x_{j+1} = x_j + s_j$ and $\Delta_{j+1} = 2\Delta_j$.
- 4 Otherwise, set $x_{j+1} = x_j$ and $\Delta_{j+1} = 0.5\Delta_j$.

Trust region for minimize $x \in \mathbb{R}^n \ f(x)$

Inputs $x_0 \in \mathbb{R}^n$, $\Delta_0 > 0$, $\eta > 0$.

For $j = 0, 1, 2, \dots$

- 1 Define $m_j(x_j + s) := \nabla f(x_j)^T s + \frac{1}{2} s^T \nabla^2 f(x_j) s$ and compute

$$s_j \in \underset{\|s\| \leq \Delta_j}{\operatorname{argmin}} m_j(x_j + s).$$

- 2 Compute $\rho_j = \frac{f(x_j) - f(x_j + s_j)}{m_j(x_j) - m_j(x_j + s_j)}$.
- 3 If $\rho_j \geq \eta$, set $x_{j+1} = x_j + s_j$ and $\Delta_{j+1} = 2\Delta_j$.
- 4 Otherwise, set $x_{j+1} = x_j$ and $\Delta_{j+1} = 0.5\Delta_j$.

- No Hessian term in m_j : Gradient descent.
- **Cost** Iterations.

Complexity results for trust region

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$ and $\nabla^2 f(\mathbf{x}_J) \succeq -\epsilon_H \mathbf{I}$.

General nonconvex f

$$J = \mathcal{O} \left(\max \{ \epsilon^{-2}, \epsilon_H^{-3} \} \right).$$

Strict saddle f (w/ Florentin Goyens)

If f is $(\gamma, \lambda, \mu, \delta)$ -strict saddle,

$$J = J_f + J_\epsilon, \quad \begin{cases} J_f &= \mathcal{O} \left(\max \{ \gamma^{-2} \lambda^{-1}, \gamma^{-2} \mu^{-1}, \lambda^{-3}, \mu^{-3}, \mu^{-2} \delta^{-1} \} \right) \\ J_\epsilon &= \log \log \left[\mathcal{O} \left(\mu \epsilon^{-1} \right) \right]. \end{cases}$$

Complexity results for trust region

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$ and $\nabla^2 f(\mathbf{x}_J) \succeq -\epsilon_H \mathbf{I}$.

General nonconvex f

$$J = \mathcal{O} \left(\max \{ \epsilon^{-2}, \epsilon_H^{-3} \} \right).$$

Strict saddle f (w/ Florentin Goyens)

If f is $(\gamma, \lambda, \mu, \delta)$ -strict saddle,

$$J = J_f + J_\epsilon, \quad \begin{cases} J_f &= \mathcal{O} \left(\max \{ \gamma^{-2} \lambda^{-1}, \gamma^{-2} \mu^{-1}, \lambda^{-3}, \mu^{-3}, \mu^{-2} \delta^{-1} \} \right) \\ J_\epsilon &= \log \log \left[\mathcal{O} \left(\mu \epsilon^{-1} \right) \right]. \end{cases}$$

Complexity results for trust region

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$ and $\nabla^2 f(\mathbf{x}_J) \succeq -\epsilon_H \mathbf{I}$.

General nonconvex f

$$J = \mathcal{O} \left(\max \{ \epsilon^{-2}, \epsilon_H^{-3} \} \right).$$

Strict saddle f (w/ Florentin Goyens)

If f is $(\gamma, \lambda, \mu, \delta)$ -strict saddle,

$$J = J_f + J_\epsilon, \quad \begin{cases} J_f &= \mathcal{O} \left(\max \{ \gamma^{-2} \lambda^{-1}, \gamma^{-2} \mu^{-1}, \lambda^{-3}, \mu^{-3}, \mu^{-2} \delta^{-1} \} \right) \\ J_\epsilon &= \log \log \left[\mathcal{O} \left(\mu \epsilon^{-1} \right) \right]. \end{cases}$$

- Second term vanishes for large ϵ .
- $\log \log$ dependency in ϵ not on ϵ_H ($\sim$ strongly convex case)!

Improved complexity using Newton steps

Improved complexity using Newton steps

- For generic nonconvex f , regularize trust-region models!
(w/ F. E. Curtis, D. P. Robinson, S. J. Wright).

$$\mathcal{O}(\max\{\epsilon^{-2}, \epsilon_H^{-3}\}) \longrightarrow \mathcal{O}(\max\{\epsilon^{-2}\epsilon_H, \epsilon_H^{-3}\}).$$

Improved complexity using Newton steps

- For generic nonconvex f , regularize trust-region models!
(w/ *F. E. Curtis, D. P. Robinson, S. J. Wright*).

$$\mathcal{O}(\max\{\epsilon^{-2}, \epsilon_H^{-3}\}) \longrightarrow \mathcal{O}(\max\{\epsilon^{-2}\epsilon_H, \epsilon_H^{-3}\}).$$

- Line-search variants (w/ *M. O'Neill, S. J. Wright*).

Improved complexity using Newton steps

- For generic nonconvex f , regularize trust-region models!
(w/ *F. E. Curtis, D. P. Robinson, S. J. Wright*).

$$\mathcal{O}(\max\{\epsilon^{-2}, \epsilon_H^{-3}\}) \longrightarrow \mathcal{O}(\max\{\epsilon^{-2}\epsilon_H, \epsilon_H^{-3}\}).$$

- Line-search variants (w/ *M. O'Neill, S. J. Wright*).

Structured nonconvex problems

Improved complexity using Newton steps

- For generic nonconvex f , regularize trust-region models!
(w/ *F. E. Curtis, D. P. Robinson, S. J. Wright*).

$$\mathcal{O}(\max\{\epsilon^{-2}, \epsilon_H^{-3}\}) \longrightarrow \mathcal{O}(\max\{\epsilon^{-2}\epsilon_H, \epsilon_H^{-3}\}).$$

- Line-search variants (w/ *M. O'Neill, S. J. Wright*).

Structured nonconvex problems

- Manifold constraints: minimize $_{x \in \mathcal{M}}$ $f(x)$, $\mathcal{M} = \mathbb{S}^{n-1}, \mathbb{C}^n, \dots$
(w/ *Florentin Goyens*)

Improved complexity using Newton steps

- For generic nonconvex f , regularize trust-region models!
(w/ F. E. Curtis, D. P. Robinson, S. J. Wright).

$$\mathcal{O}(\max\{\epsilon^{-2}, \epsilon_H^{-3}\}) \longrightarrow \mathcal{O}(\max\{\epsilon^{-2}\epsilon_H, \epsilon_H^{-3}\}).$$

- Line-search variants (w/ M. O'Neill, S. J. Wright).

Structured nonconvex problems

- Manifold constraints: minimize $_{\mathbf{x} \in \mathcal{M}}$ $f(\mathbf{x})$, $\mathcal{M} = \mathbb{S}^{n-1}, \mathbb{C}^n, \dots$
(w/ Florentin Goyens)
- Least-squares problems: minimize $_{\mathbf{x} \in \mathbb{R}^n}$ $f(\mathbf{x}) = \frac{1}{2} \|\mathbf{r}(\mathbf{x})\|^2$.
(PhD Iskander Legheraba).

- 1 Newton-type methods
- 2 Conjugate gradient methods
- 3 Direct-search methods
- 4 Perspectives

Setup

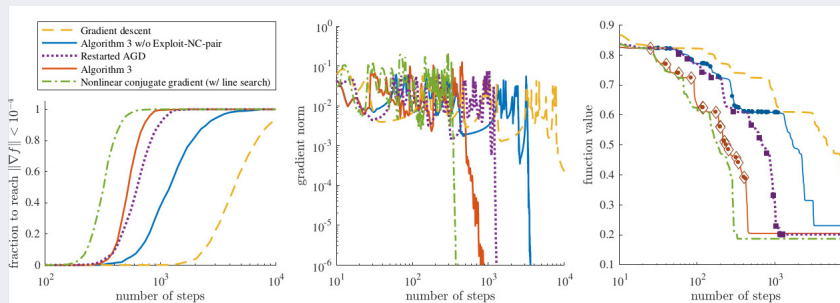
Problem minimize $_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x})$, $f \mathcal{C}^1$ nonconvex.

Goal $\|\nabla f(\mathbf{x})\| \leq \epsilon$.

Problem minimize $\mathbf{x} \in \mathbb{R}^n$ $f(\mathbf{x})$, $f \in \mathcal{C}^1$ nonconvex.

Goal $\|\nabla f(\mathbf{x})\| \leq \epsilon$.

An observation (from Carmon et al '17)



- Gradient descent (for $f \in \mathcal{C}^1$): $\mathcal{O}(\epsilon^{-2})$ calls to ∇f .
- Algorithm 3/AGD (for $f \in \mathcal{C}^2$): $\tilde{\mathcal{O}}(\epsilon^{-7/4})$ calls to ∇f .
- **All outperformed by standard nonlinear CG on a nonconvex regression task!**

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

① Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^T \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.
- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.
- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.
- 4 If $\nabla f(\mathbf{x}_{j+1})^\top \mathbf{d}_{j+1} \geq 0$, set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1})$.

Basic nonlinear CG framework

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.
- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.
- 4 If $\nabla f(\mathbf{x}_{j+1})^\top \mathbf{d}_{j+1} \geq 0$, set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1})$.

- β_{j+1} chosen using standard formulas (PRP+), restart condition practical.

No complexity guarantees!

Basic nonlinear CG framework

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.
- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.
- 4 If $\nabla f(\mathbf{x}_{j+1})^\top \mathbf{d}_{j+1} \geq 0$, set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1})$.

- β_{j+1} chosen using standard formulas (PRP+), restart condition practical.

No complexity guarantees!

- **Our approach** Change restart condition!

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.
- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$, $p \geq 0$, $\kappa \in (0, 1)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.

- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.

- 4 If $\nabla f(\mathbf{x}_{j+1})^\top \mathbf{d}_{j+1} \geq -\kappa \|\nabla f(\mathbf{x}_{j+1})\|^{1+p}$ or $\|\mathbf{d}_{j+1}\| \geq \kappa \|\nabla f(\mathbf{x}_{j+1})\|^{\frac{1+p}{2}}$,
set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1})$.

Nonlinear CG with restart conditions

Init $\mathbf{x}_0 \in \mathbb{R}^n$, $(\eta, \theta) \in (0, 1)^2$, $\mathbf{d}_0 = -\nabla f(\mathbf{x}_0)$, $p \geq 0$, $\kappa \in (0, 1)$.

For $j = 0, 1, 2, \dots$

- 1 Compute $\alpha_j \in \{1, \theta, \theta^2, \dots\}$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) + \eta \alpha_j \nabla f(\mathbf{x}_j)^\top \mathbf{d}_j.$$

- 2 Set $\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{d}_j$.

- 3 Set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1}) + \beta_{j+1} \mathbf{d}_j$.

- 4 If $\nabla f(\mathbf{x}_{j+1})^\top \mathbf{d}_{j+1} \geq -\kappa \|\nabla f(\mathbf{x}_{j+1})\|^{1+p}$ or $\|\mathbf{d}_{j+1}\| \geq \kappa \|\nabla f(\mathbf{x}_{j+1})\|^{\frac{1+p}{2}}$,
set $\mathbf{d}_{j+1} = -\nabla f(\mathbf{x}_{j+1})$.

- Restarted iterations: Take gradient steps.
- Non-restarted: Guarantees on step quality.

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem (w/ Rémi Chan–Renous–Legoubin)

For restarted NCG, J is at most

$$\underbrace{\mathcal{O}(\epsilon^{-2})}_{\text{restarted}} + \underbrace{\mathcal{O}(\epsilon^{-(1+p)})}_{\text{non restarted}}$$

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem (w/ Rémi Chan–Renous–Legoubin)

For restarted NCG, J is at most

$$\underbrace{\mathcal{O}(\epsilon^{-2})}_{\text{restarted}} + \underbrace{\mathcal{O}(\epsilon^{-(1+p)})}_{\text{non restarted}}$$

- The bound is still $\mathcal{O}(\epsilon^{-2})$ (like gradient descent)
- But better results for **non-restarted iterations** for $p \in [0, 1)$!

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem (w/ Rémi Chan–Renous–Legoubin)

For restarted NCG, J is at most

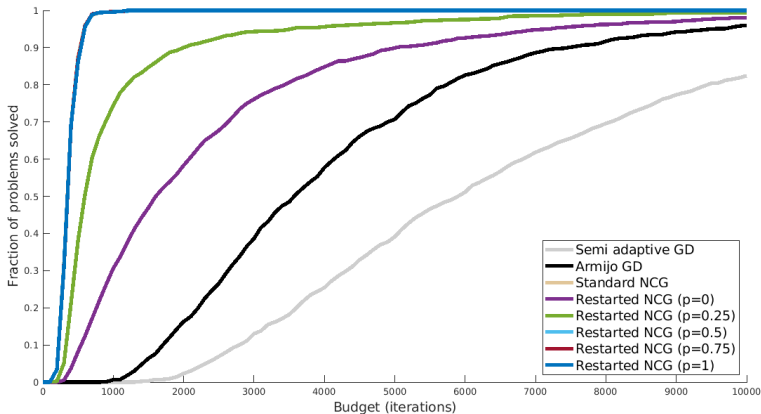
$$\underbrace{\mathcal{O}(\epsilon^{-2})}_{\text{restarted}} + \underbrace{\mathcal{O}(\epsilon^{-(1+p)})}_{\text{non restarted}}$$

- The bound is still $\mathcal{O}(\epsilon^{-2})$ (like gradient descent)
- But better results for **non-restarted iterations** for $p \in [0, 1)$!

How do restarted methods behave in practice?

Comparison

- Two gradient-based methods, Standard linear CG, restarted variants.
- Nonconvex regression instances from (Carmon et al '17).



Restarted methodology (w/ A. S. Berahas, M. O'Neill)

- Applied to quasi-Newton schemes.
- Adapted to noisy optimization.

→ *Restarting can be applied to yield complexity guarantees without harming the practical performance!*

Restarted methodology (w/ A. S. Berahas, M. O'Neill)

- Applied to quasi-Newton schemes.
- Adapted to noisy optimization.

→ *Restarting can be applied to yield complexity guarantees without harming the practical performance!*

Special case: quadratic functions (w/ M. O'Neill, S. J. Wright)

- Nonconvex quadratics from Newton-type methods.
- Can get complexity for **linear** conjugate gradient applied to those problems.

→ *Extra checks in algorithms allow to detect and exploit nonconvexity!*

- 1 Newton-type methods
- 2 Conjugate gradient methods
- 3 Direct-search methods**
- 4 Perspectives

Derivative-free paradigm

$$\underset{\boldsymbol{x} \in \mathbb{R}^n}{\text{minimize}} \quad f(\boldsymbol{x}), \quad f \in \mathcal{C}^1 \quad \text{nonconvex.}$$

Blackbox/Derivative-free optimization

- **Derivatives unavailable for algorithmic use.**
- Only access to values of f .
- Those can take a long time to compute!

Derivative-free paradigm

$$\underset{\mathbf{x} \in \mathbb{R}^n}{\text{minimize}} \quad f(\mathbf{x}), \quad f \in \mathcal{C}^1 \quad \text{nonconvex.}$$

Blackbox/Derivative-free optimization

- **Derivatives unavailable for algorithmic use.**
- Only access to values of f .
- Those can take a long time to compute!

Complexity aspects

- **Goal** Find $\mathbf{x}_k$ such that $\|\nabla f(\mathbf{x}_k)\| \leq \epsilon$.
- **Cost** Number of **function evaluations**.

Basic direct-search framework

Inputs $x_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- Choose $\mathcal{D}_j \subset \mathbb{R}^n$.

Basic direct-search framework

Inputs $\mathbf{x}_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- Choose $\mathcal{D}_j \subset \mathbb{R}^n$.
- If $\exists \mathbf{d}_j \in \mathcal{D}_j$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) - \alpha_j^2 \|\mathbf{d}_j\|^2$$

set $\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{d}_j$, $\alpha_{j+1} := 2\alpha_j$.

Inputs $\mathbf{x}_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- Choose $\mathcal{D}_j \subset \mathbb{R}^n$.
- If $\exists \mathbf{d}_j \in \mathcal{D}_j$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) - \alpha_j^2 \|\mathbf{d}_j\|^2$$

set $\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{d}_j$, $\alpha_{j+1} := 2\alpha_j$.

- Otherwise, set $\mathbf{x}_{j+1} := \mathbf{x}_j$, $\alpha_{j+1} := 0.5\alpha_j$.

Basic direct-search framework

Inputs $\mathbf{x}_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- **Choose a PSS** $\mathcal{D}_j \subset \mathbb{R}^n$.
- If $\exists \mathbf{d}_j \in \mathcal{D}_j$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) - \alpha_j^2 \|\mathbf{d}_j\|^2$$

set $\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{d}_j$, $\alpha_{j+1} := 2\alpha_j$.

- Otherwise, set $\mathbf{x}_{j+1} := \mathbf{x}_j$, $\alpha_{j+1} := 0.5\alpha_j$.

Positive Spanning Set (PSS)

$$\mathcal{D} \subset \mathbb{R}^n \text{ PSS} \quad \Longleftrightarrow \quad \text{cm}(\mathcal{D}) = \max_{\mathbf{d} \in \mathcal{D}} \min_{\|\mathbf{v}\|=1} \frac{\mathbf{d}^T \mathbf{v}}{\|\mathbf{d}\|} > 0.$$

$\text{cm}(\mathcal{D})$: Cosine measure.

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem

If $\mathcal{D}_j = \mathcal{D} \ \forall j$, the method takes at most

$$\mathcal{O} \left(|\mathcal{D}| \operatorname{cm}(\mathcal{D})^{-2} \epsilon^{-2} \right)$$

calls to f .

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem

If $\mathcal{D}_j = \mathcal{D} \ \forall j$, the method takes at most

$$\mathcal{O} \left(|\mathcal{D}| \operatorname{cm}(\mathcal{D})^{-2} \epsilon^{-2} \right)$$

calls to f .

- $|\mathcal{D}|$ and $\operatorname{cm}(\mathcal{D})$ depend on dimension n !
- Best known bound: $\mathcal{O}(n^2 \epsilon^{-2})$.

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem

If $\mathcal{D}_j = \mathcal{D} \ \forall j$, the method takes at most

$$\mathcal{O} \left(|\mathcal{D}| \operatorname{cm}(\mathcal{D})^{-2} \epsilon^{-2} \right)$$

calls to f .

- $|\mathcal{D}|$ and $\operatorname{cm}(\mathcal{D})$ depend on dimension n !
- Best known bound: $\mathcal{O}(n^2 \epsilon^{-2})$.

What if certain calls to f take too long to compute?

Direct search with stragglers

Inputs $\mathbf{x}_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- Choose $\mathcal{D}_j \subset \mathbb{R}^n$.
- If $\exists \mathbf{d}_j \in \mathcal{D}_j \setminus \mathcal{S}_j$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) - \alpha_j^2 \|\mathbf{d}_j\|^2$$

set $\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{d}_j$, $\alpha_{j+1} := 2\alpha_j$.

- Otherwise, set $\mathbf{x}_{j+1} := \mathbf{x}_j$, $\alpha_{j+1} := 0.5\alpha_j$.

- $\mathcal{S}_j$ Unknown straggler evaluations (at most k).

Direct search with stragglers

Inputs $\mathbf{x}_0 \in \mathbb{R}^n$, $\alpha_0 > 0$.

For $j = 0, 1, \dots$

- Choose a PkSS $\mathcal{D}_j \subset \mathbb{R}^n$.
- If $\exists \mathbf{d}_j \in \mathcal{D}_j \setminus \mathcal{S}_j$ such that

$$f(\mathbf{x}_j + \alpha_j \mathbf{d}_j) < f(\mathbf{x}_j) - \alpha_j^2 \|\mathbf{d}_j\|^2$$

set $\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{d}_j$, $\alpha_{j+1} := 2\alpha_j$.

- Otherwise, set $\mathbf{x}_{j+1} := \mathbf{x}_j$, $\alpha_{j+1} := 0.5\alpha_j$.

- $\mathcal{S}_j$ Unknown straggler evaluations (at most k).
- Fix: Use richer $\mathcal{D}_j$ s!

Definition

$\mathcal{D} \subset \mathbb{R}^n$ is a $PkSS$ with $k \geq 1$ if

- Any $\mathcal{N} \subset \mathcal{D}$ with $|\mathcal{N}| = |\mathcal{D}| - k + 1$ is a PSS.
- Removing $k - 1$ vectors from $\mathcal{D}$ does not change its PSS nature.

Positive k -spanning sets ($PkSS$)

Definition

$\mathcal{D} \subset \mathbb{R}^n$ is a $PkSS$ with $k \geq 1$ if

- Any $\mathcal{N} \subset \mathcal{D}$ with $|\mathcal{N}| = |\mathcal{D}| - k + 1$ is a PSS.
- Removing $k - 1$ vectors from $\mathcal{D}$ does not change its PSS nature.

The k -cosine measure (w/ W. Hare, G. Jarry-Bolduc, Sébastien Kerleau)

- For any $\mathcal{D} \subset \mathbb{R}^n$, the k -cosine measure of $\mathcal{D}$ is

$$\text{cm}_k(\mathcal{D}) = \min_{\substack{\mathcal{N} \subset \mathcal{D} \\ |\mathcal{N}| = |\mathcal{D}| - k + 1}} \text{cm}(\mathcal{N}).$$

- $\mathcal{D} \text{ } PkSS \iff \text{cm}_k(\mathcal{D}) > 0.$

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem (*PhD Sébastien Kerleau*)

Suppose $\mathcal{D}_j = \mathcal{D}$ PkSS for all j . Then, the method takes at most

$$\mathcal{O}(|\mathcal{D}| \text{cm}_k(\mathcal{D})^{-2} \epsilon^{-2})$$

calls to f .

Goal $\|\nabla f(\mathbf{x}_J)\| \leq \epsilon$.

Theorem (*PhD Sébastien Kerleau*)

Suppose $\mathcal{D}_j = \mathcal{D}$ PkSS for all j . Then, the method takes at most

$$\mathcal{O}(|\mathcal{D}| \text{cm}_k(\mathcal{D})^{-2} \epsilon^{-2})$$

calls to f .

- $|\mathcal{D}|$ and $\text{cm}_k(\mathcal{D})$ depend on k and n .
- Best found bound: $\mathcal{O}(kn^2\epsilon^{-2})$.

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin

· 23 janvier · 🌐

...

In progress: adding <https://github.com/lindonroberts/directsearch> inside Nevergrad.

In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin

· 23 janvier · 🌐

...

In progress: adding <https://github.com/lindonroberts/directsearch> inside Nevergrad.

In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

- Surveyed the area (Dzahini et al '25)
- Dimension dependencies can be improved through randomization!

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin

· 23 janvier · 🌐

...

In progress: adding <https://github.com/lindonroberts/directsearch> inside Nevergrad.

In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

- Surveyed the area (Dzahini et al '25)

→ Dimension dependencies can be improved through randomization!

Geometrical structures in direct search

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin · 23 janvier · 🌐

In progress: adding <https://github.com/lindonroberts/directsearch> inside Nevergrad.
In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

- Surveyed the area (Dzahini et al '25)
- Dimension dependencies can be improved through randomization!

Geometrical structures in direct search

- Decomposition of PSSs along subspaces through graph theory arguments (Cornaz et al '25).

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin

· 23 janvier · 🌐

...

In progress: adding <https://github.com/londonroberts/directsearch> inside Nevergrad.

In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

- Surveyed the area (Dzahini et al '25)
- Dimension dependencies can be improved through randomization!

Geometrical structures in direct search

- Decomposition of PSSs along subspaces through graph theory arguments (Cornaz et al '25).
- Generation of Pk SSs through polytopes (Kerleau '25).

Complexity and direct search

- Randomized subspace variant (Roberts & Royer '23)



Olivier Teytaud

Admin · 23 janvier · 🌐

In progress: adding <https://github.com/londonroberts/directsearch> inside Nevergrad.
In particular there is an excellent stochastic direct search method. I don't know exactly the algorithm (yet).
Thanks guys for this excellent code!

- Surveyed the area (Dzahini et al '25)

→ Dimension dependencies can be improved through randomization!

Geometrical structures in direct search

- Decomposition of PSSs along subspaces through graph theory arguments (Cornaz et al '25).
- Generation of Pk SSs through polytopes (Kerleau '25).

→ Discrete structures matter for direct search!

- 1 Newton-type methods
- 2 Conjugate gradient methods
- 3 Direct-search methods
- 4 Perspectives

Complexity $\longrightarrow$ **Structures**

Complexity $\longrightarrow$ Structures

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
- Study special classes of nonconvex problems.

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
- Checks to detect nonconvex behavior.

- **Derivative-free methods**

- Improve complexity with randomness.
- Use geometry to handle difficult settings.

Complexity $\longrightarrow$ Structures

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.

- Optimal bounds within the class of second-order methods**

- Study special classes of nonconvex problems.

- Structured nonconvexity**

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
 - Checks to detect nonconvex behavior.

- **Derivative-free methods**

- Improve complexity with randomness.
 - Use geometry to handle difficult settings.

Complexity $\longrightarrow$ Structures

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
Optimal bounds within the class of second-order methods
- Study special classes of nonconvex problems.
Structured nonconvexity

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
Complexity obtained without hurting performance
- Checks to detect nonconvex behavior.
Leverage algorithmic structure

- **Derivative-free methods**

- Improve complexity with randomness.
- Use geometry to handle difficult settings.

Complexity $\longrightarrow$ Structures

- **Newton-type methods**

- Obtain sharp complexity with Newton steps.
Optimal bounds within the class of second-order methods
- Study special classes of nonconvex problems.
Structured nonconvexity

- **Conjugate gradient methods**

- Restart conditions in nonlinear case.
Complexity obtained without hurting performance
- Checks to detect nonconvex behavior.
Leverage algorithmic structure

- **Derivative-free methods**

- Improve complexity with randomness.
Subspace variants with complexity
- Use geometry to handle difficult settings.
Graph/Polytope structures

Research structures





Goal: Optimize moves during renovation
(with S. Airiau, L. Galand, J. Lang, S. Toubaline)



Goal: Optimize moves during renovation
(with S. Airiau, L. Galand, J. Lang, S. Toubaline)

- Linear Program (simplicity) → Other models are nonconvex!



Goal: Optimize moves during renovation
(with S. Airiau, L. Galand, J. Lang, S. Toubaline)

- Linear Program (simplicity) → Other models are nonconvex!
- Continuous LP solvers struggle → Revisit algorithms!



Goal: Optimize moves during renovation
(with S. Airiau, L. Galand, J. Lang, S. Toubaline)

- Linear Program (simplicity)→ Other models are nonconvex!
- Continuous LP solvers struggle→ Revisit algorithms!
- Model parameters to tune→ Derivative-free methods!



Goal: Optimize moves during renovation
(with S. Airiau, L. Galand, J. Lang, S. Toubaline)

- Linear Program (simplicity)→ Other models are nonconvex!
- Continuous LP solvers struggle→ Revisit algorithms!
- Model parameters to tune→ Derivative-free methods!

Complexity and **structure** key to practical efficiency

Thank you for your attention!