

Random Gradient-Free Optimization in Infinite Dimensional Spaces

Daniel Csillag
FGV EMAp

June 17, 2026

Joint work with Caio Peixoto, Bernardo Costa and Yuri Saporito

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Loss & gradient operator:

$$F : \mathbb{R}^d \rightarrow \mathbb{R}, \quad \nabla F : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Loss & gradient operator:

$$F : \mathbb{R}^d \rightarrow \mathbb{R}, \quad \nabla F : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Optimize with gradient descent:

$$\theta_{t+1} = \theta_t - \eta \nabla F(\theta_t)$$

(or some variant thereof, e.g. Adam)

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Loss & gradient operator:

$$F : \mathbb{R}^d \rightarrow \mathbb{R}, \quad \nabla F : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Optimize with gradient descent:

$$\theta_{t+1} = \theta_t - \eta \nabla F(\theta_t)$$

(or some variant thereof, e.g. Adam)

$$\operatorname{argmin}_{h \in \mathcal{H}} \mathcal{R}(h), \quad \mathcal{H} \text{ Hilbert}$$

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Loss & gradient operator:

$$F : \mathbb{R}^d \rightarrow \mathbb{R}, \quad \nabla F : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Optimize with gradient descent:

$$\theta_{t+1} = \theta_t - \eta \nabla F(\theta_t)$$

(or some variant thereof, e.g. Adam)

$$\operatorname{argmin}_{h \in \mathcal{H}} \mathcal{R}(h), \quad \mathcal{H} \text{ Hilbert}$$

Loss & gradient operator:

$$\mathcal{R} : \mathcal{H} \rightarrow \mathbb{R}, \quad \nabla \mathcal{R} : \mathcal{H} \rightarrow \mathcal{H}$$

Functional Optimization

$$\operatorname{argmin}_{f : \mathcal{X} \rightarrow \mathcal{Y}} \mathcal{R}(f)$$

$$\operatorname{argmin}_{\theta \in \mathbb{R}^d} F(\theta), \quad F(\theta) = \mathcal{R}(f_\theta)$$

Loss & gradient operator:

$$F : \mathbb{R}^d \rightarrow \mathbb{R}, \quad \nabla F : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Optimize with gradient descent:

$$\theta_{t+1} = \theta_t - \eta \nabla F(\theta_t)$$

(or some variant thereof, e.g. Adam)

$$\operatorname{argmin}_{h \in \mathcal{H}} \mathcal{R}(h), \quad \mathcal{H} \text{ Hilbert}$$

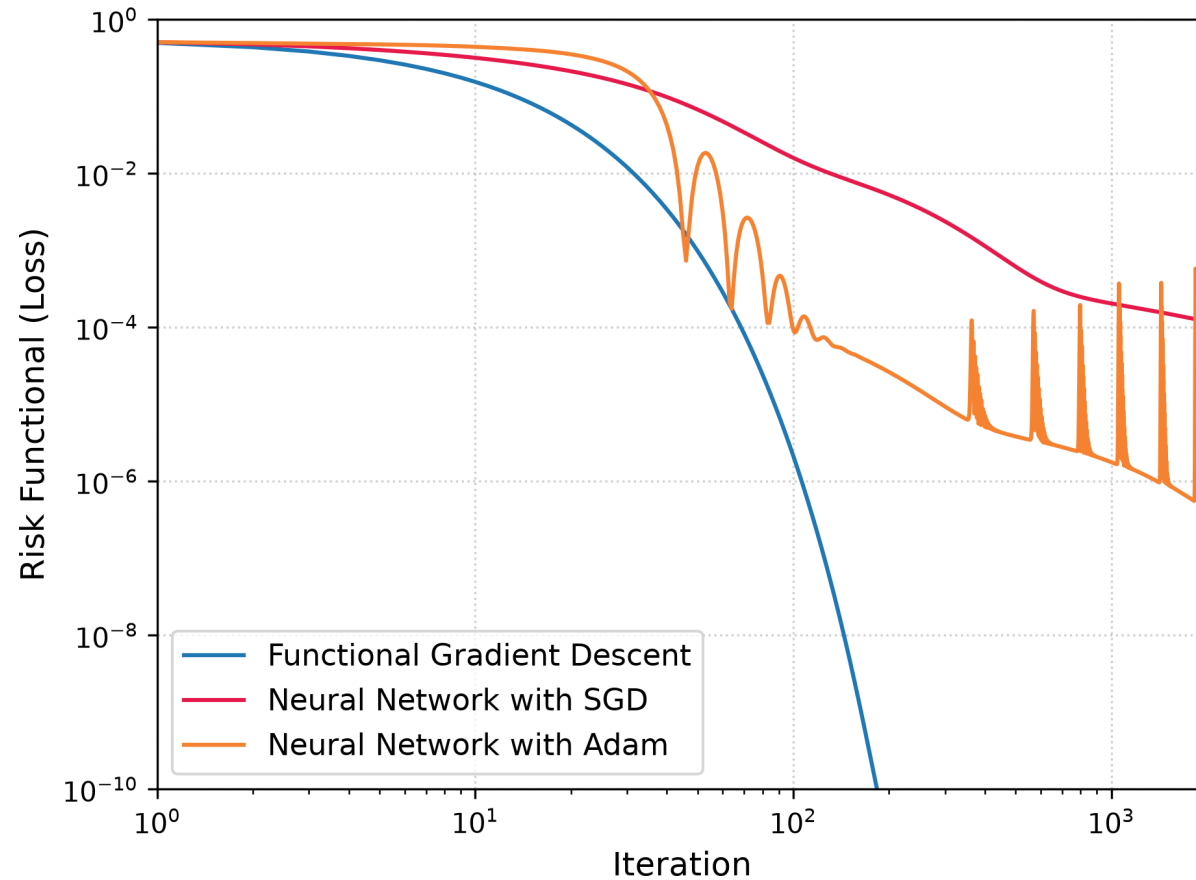
Loss & gradient operator:

$$\mathcal{R} : \mathcal{H} \rightarrow \mathbb{R}, \quad \nabla \mathcal{R} : \mathcal{H} \rightarrow \mathcal{H}$$

Optimize with *functional* gradient descent:

$$h_{t+1} = h_t - \eta \nabla \mathcal{R}(h_t)$$

Functional Optimization



Functional Gradients are Hard

Problem: it's really hard to compute functional gradients.

Functional Gradients are Hard

Problem: it's really hard to compute functional gradients.

Running example: Solving a PDE

$$\begin{cases} L[u](x) = 0 & x \in \Omega \\ u(x) = f(x) & x \in \partial\Omega \end{cases}$$

where L is a differential operator; consider the risk functional

$$\mathcal{R}(h) : \mathbb{H}^s(\Omega) \rightarrow \mathbb{R}$$

$$h \mapsto \|L[h]\|_{L^2(\Omega)}^2 + \|h - f\|_{L^2(\partial\Omega)}^2$$

Functional Gradients are Hard

Full gradients are ***hard***: $\nabla \mathcal{R}(h)$ is the unique $g \in \mathbb{H}^s(\Omega)$ such that

$$\sum_{j=0}^s \int_{\Omega} \mathbf{D}^{(j)} g(x) \cdot \mathbf{D}^{(j)} v(x) \, dx = \int_{\Omega} L[h](x) L[v](x) \, dx + \int_{\partial\Omega} (h(x) - f(x)) v(x) \, \sigma(dx)$$

for all $v \in \mathbb{H}^s(\Omega)$.

Functional Gradients are Hard

Full gradients are **hard**: $\nabla \mathcal{R}(h)$ is the unique $g \in \mathbb{H}^s(\Omega)$ such that

$$\sum_{j=0}^s \int_{\Omega} D^{(j)} g(x) \cdot D^{(j)} v(x) \, dx = \int_{\Omega} L[h](x) L[v](x) \, dx + \int_{\partial\Omega} (h(x) - f(x)) v(x) \, \sigma(dx)$$

for all $v \in \mathbb{H}^s(\Omega)$.

But directional derivatives are easy:

$$\begin{aligned} D\mathcal{R}(h; v) &:= \lim_{\delta \rightarrow 0} \frac{\mathcal{R}(h + \delta v) - \mathcal{R}(h)}{\delta} \\ &= \left[\frac{d}{d\delta} \mathcal{R}(h + \delta v) \right]_{\delta=0} \end{aligned}$$

Grad. Descent with Directional Derivatives

Q: Can we use just directional derivatives for optimization?

Grad. Descent with Directional Derivatives

Q: Can we use just directional derivatives for optimization?

A: In finite dimensions, yes: Random 'Gradient-Free' Optimization

Grad. Descent with Directional Derivatives

Q: Can we use just directional derivatives for optimization?

A: In finite dimensions, yes: Random ‘Gradient-Free’ Optimization

To estimate $\nabla \mathcal{R}(h)$, sample $v_1, \dots, v_M \sim \mathbb{P}_v$, and compute:

$$h, v_1, \dots, v_M \in \mathbb{R}^d$$

$$\hat{g} = \frac{1}{M} \sum_{m=1}^M \mathrm{D}\mathcal{R}(h; v_m) v_m$$

Grad. Descent with Directional Derivatives

Q: Can we use just directional derivatives for optimization?

A: In finite dimensions, yes: Random ‘Gradient-Free’ Optimization

To estimate $\nabla \mathcal{R}(h)$, sample $v_1, \dots, v_M \sim \mathbb{P}_v$, and compute:

$$h, v_1, \dots, v_M \in \mathbb{R}^d$$

$$\hat{g} = \frac{1}{M} \sum_{m=1}^M \mathrm{D}\mathcal{R}(h; v_m) v_m$$

$$\bullet \quad \mathbb{E}[v] = 0 \text{ and } \mathrm{Cov}[v] = I \quad \implies \quad \mathbb{E}[\hat{g}] = \nabla L(h).$$

$$(\text{e.g. } \mathbb{P}_v = \mathcal{N}(0, I))$$

Grad. Descent with Directional Derivatives

Q: Can we use just directional derivatives for optimization?

A: In finite dimensions, yes: Random ‘Gradient-Free’ Optimization

To estimate $\nabla \mathcal{R}(h)$, sample $v_1, \dots, v_M \sim \mathbb{P}_v$, and compute:

$$h, v_1, \dots, v_M \in \mathbb{R}^d$$

$$\hat{g} = \frac{1}{M} \sum_{m=1}^M \mathrm{D}\mathcal{R}(h; v_m) v_m$$

$$\bullet \mathbb{E}[v] = 0 \text{ and } \mathrm{Cov}[v] = I \implies \mathbb{E}[\hat{g}] = \nabla L(h). \quad (\text{e.g. } \mathbb{P}_v = \mathcal{N}(0, I))$$

However, this breaks in infinite dimensions:

- Variance blows up as $d \rightarrow +\infty$
- There is no isotropic Gaussian $\mathcal{N}(0, I)$, nor analogous distributions
- Even if we change the distribution, $\mathbb{E}[\|v\|^2] < +\infty \implies \mathrm{Cov}[v] \neq I$.

Infinite Dim. Grad.-Free Descent

Ingredients:

- A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
$$\mathcal{B} \text{ is linearly independent} \quad \text{and} \quad \text{span}(\mathcal{B}) \text{ is dense in } \mathcal{H}$$
- A distribution $\mathbb{P}_K$ over the positive naturals;

Infinite Dim. Grad.-Free Descent

Ingredients:

- A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
 $\mathcal{B}$ is linearly independent and $\text{span}(\mathcal{B})$ is dense in $\mathcal{H}$
- A distribution $\mathbb{P}_K$ over the positive naturals;

The procedure: To estimate $\nabla \mathcal{R}(h)$:

- Sample a dimension $K \sim \mathbb{P}_K$, and choose a batch size $M_K \in \mathbb{N}$ (e.g. $\lceil K/2 \rceil$)

Infinite Dim. Grad.-Free Descent

Ingredients:

- A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
 $\mathcal{B}$ is linearly independent and $\text{span}(\mathcal{B})$ is dense in $\mathcal{H}$
- A distribution $\mathbb{P}_K$ over the positive naturals;

The procedure: To estimate $\nabla \mathcal{R}(h)$:

- Sample a dimension $K \sim \mathbb{P}_K$, and choose a batch size $M_K \in \mathbb{N}$ (e.g. $\lceil K/2 \rceil$)
- Sample weights $\mathbf{w}^{(1)}, \dots, \mathbf{w}^{(M_K)} \sim \mathcal{N}(0, \Sigma_K)$ in $\mathbb{R}^K$, with a tailored Σ_K

Infinite Dim. Grad.-Free Descent

Ingredients:

- A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
$$\mathcal{B} \text{ is linearly independent} \quad \text{and} \quad \text{span}(\mathcal{B}) \text{ is dense in } \mathcal{H}$$
- A distribution $\mathbb{P}_K$ over the positive naturals;

The procedure: To estimate $\nabla \mathcal{R}(h)$:

- Sample a dimension $K \sim \mathbb{P}_K$, and choose a batch size $M_K \in \mathbb{N}$ (e.g. $\lceil K/2 \rceil$)
- Sample weights $\mathbf{w}^{(1)}, \dots, \mathbf{w}^{(M_K)} \sim \mathcal{N}(0, \Sigma_K)$ in $\mathbb{R}^K$, with a tailored Σ_K
- Compute

$$\hat{g} = \frac{\lambda_K}{M_K} \sum_{m=1}^{M_K} \text{D}\mathcal{R}(h; v_m) v_m, \quad \text{where} \quad v_i := \sum_{j=1}^K \mathbf{w}_j^{(i)} b_j$$

Infinite Dim. Grad.-Free Descent

Ingredients:

- A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
$$\mathcal{B} \text{ is linearly independent} \quad \text{and} \quad \text{span}(\mathcal{B}) \text{ is dense in } \mathcal{H}$$
- A distribution $\mathbb{P}_K$ over the positive naturals;

The procedure: To estimate $\nabla \mathcal{R}(h)$:

- Sample a dimension $K \sim \mathbb{P}_K$, and choose a batch size $M_K \in \mathbb{N}$ (e.g. $\lceil K/2 \rceil$)
- Sample weights $\mathbf{w}^{(1)}, \dots, \mathbf{w}^{(M_K)} \sim \mathcal{N}(0, \Sigma_K)$ in $\mathbb{R}^K$, with a tailored Σ_K
- Compute

$$\hat{g} = \frac{\lambda_K}{M_K} \sum_{m=1}^{M_K} \text{D}\mathcal{R}(h; v_m) v_m, \quad \text{where} \quad v_i := \sum_{j=1}^K \mathbf{w}_j^{(i)} b_j$$

Infinite Dim. Grad.-Free Descent

Proposition. $\hat{g}$ is an unbiased estimator of a preconditioned gradient:

$$\mathbb{E}[\hat{g}] = C \nabla \mathcal{R}(h),$$

where $C : \mathcal{H} \rightarrow \mathcal{H}$ is a positive-definite diagonal operator determined by $(\lambda_1, \lambda_2, \dots)$.

Infinite Dim. Grad.-Free Descent

Proposition. $\hat{g}$ is an unbiased estimator of a preconditioned gradient:

$$\mathbb{E}[\hat{g}] = C \nabla \mathcal{R}(h),$$

where $C : \mathcal{H} \rightarrow \mathcal{H}$ is a positive-definite diagonal operator determined by $(\lambda_1, \lambda_2, \dots)$.

If we take $\lambda_k \equiv 1$, then we get $C = I$. But the **variance explodes**:

Proposition. If $\lambda_k \equiv 1$, then

$$\mathbb{E}[\|\hat{g}\|^2] \geq \sum_{i=1}^{\infty} \frac{\langle \nabla \mathcal{R}(h), e_i \rangle^2}{t_i}.$$

Infinite Dim. Grad.-Free Descent

Proposition. $\hat{g}$ is an unbiased estimator of a preconditioned gradient:

$$\mathbb{E}[\hat{g}] = C \nabla \mathcal{R}(h),$$

where $C : \mathcal{H} \rightarrow \mathcal{H}$ is a positive-definite diagonal operator determined by $(\lambda_1, \lambda_2, \dots)$.

So we take $\lambda_k = \mathbb{P}[K \geq k]$:

Proposition. Take $\lambda_k = \mathbb{P}[K \geq k]$ and $M_k = \lceil k/c \rceil$, for some $c > 0$. We then have

$$\mathbb{E}[\|\hat{g}\|_C^2] \leq 2(1 + 2c) \|\nabla \mathcal{R}(h)\|^2,$$

where $\|\hat{g}\|_C := \|C^{-\frac{1}{2}} \hat{g}\|$.

Convergence

Theorem. If the risk is convex, has L -Lipschitz gradients and the minimizer h^* is reasonably regular, then for a sufficiently small learning rate α ,

$$\mathbb{E}[\mathcal{R}(\hat{h})] - \mathcal{R}(h^*) \leq \|h^* - h_1\|_C^2 \cdot \left(\frac{1}{2N\alpha} + \alpha(1 + 2c)L^2 \right),$$

where $\hat{h}$ is the averaged iterate $\frac{1}{N}(h_1 + \dots + h_N)$.

Convergence

Theorem. If the risk is convex, has L -Lipschitz gradients and the minimizer h^* is reasonably regular, then for a sufficiently small learning rate α ,

$$\mathbb{E}[\mathcal{R}(\hat{h})] - \mathcal{R}(h^*) \leq \|h^* - h_1\|_C^2 \cdot \left(\frac{1}{2N\alpha} + \alpha(1 + 2c)L^2 \right),$$

where $\hat{h}$ is the averaged iterate $\frac{1}{N}(h_1 + \dots + h_N)$.

Corollary. Choosing $\alpha = O(1/\sqrt{N})$, we get

$$\mathbb{E}[\mathcal{R}(\hat{h})] - \mathcal{R}(h^*) \leq O\left(\frac{1}{\sqrt{N}}\right).$$

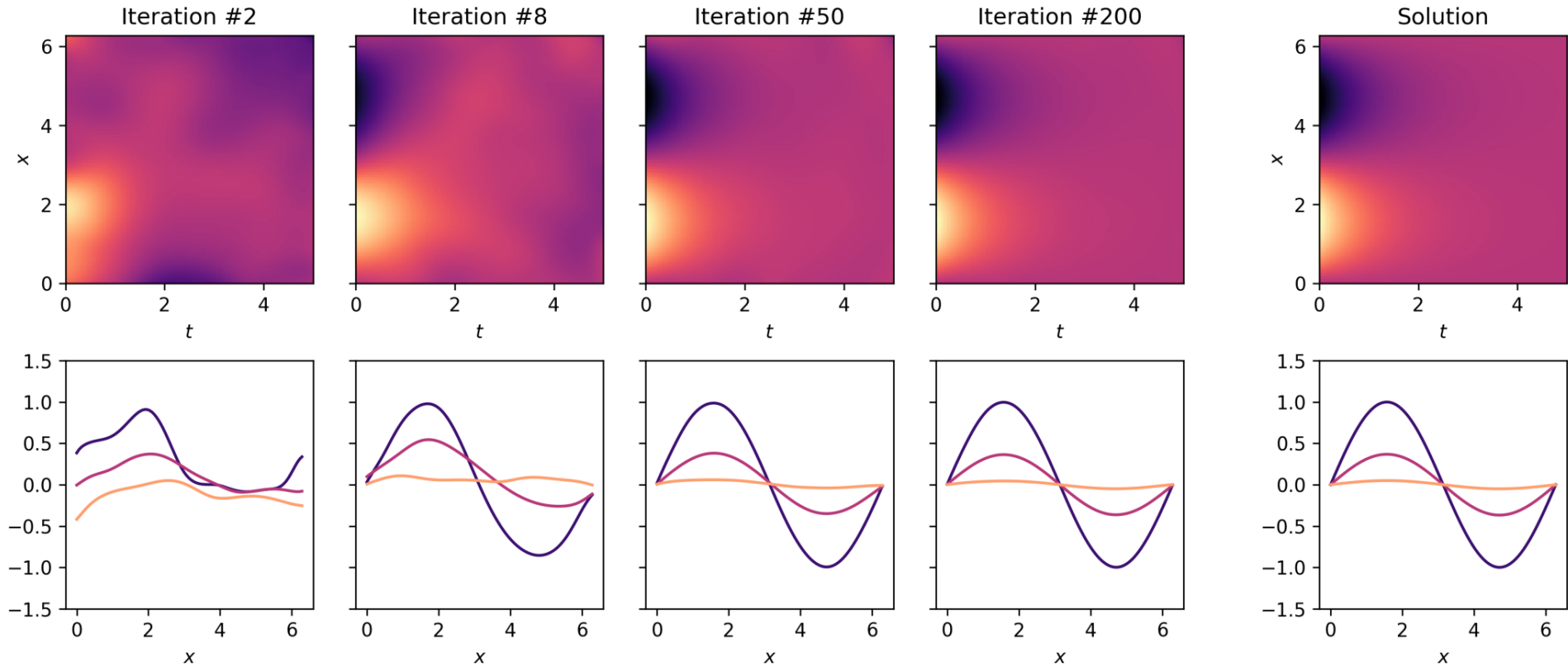
Example 1: Heat Equation

A simple heat equation:

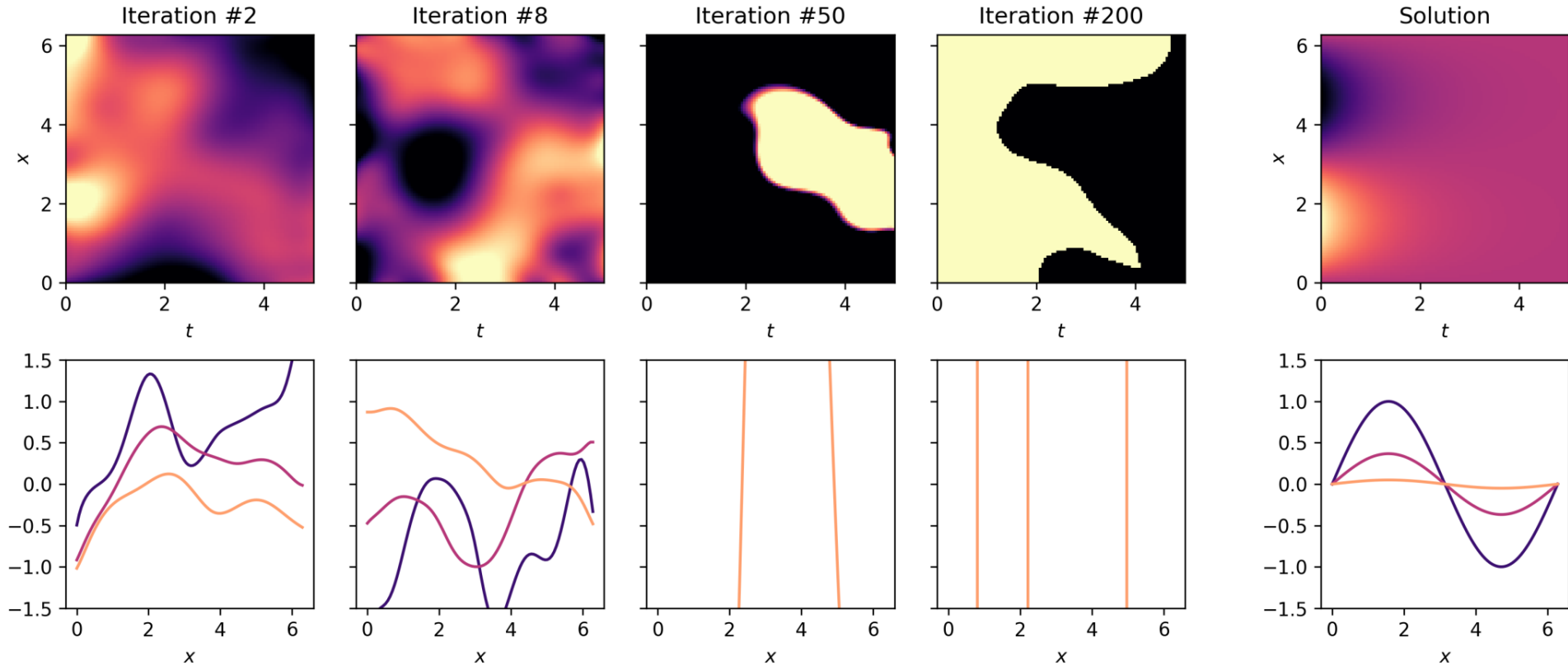
$$\begin{cases} \partial_t u(t, x) = \Delta u(t, x) & t \in (0, T), x \in (0, 2\pi) \\ u(0, x) = \sin x & x \in (0, 2\pi) \\ u(t, 0) = u(t, 2\pi) = 0 & t \in (0, T) \end{cases}$$

Unique solution given by $u(t, x) = e^{-t} \sin x$

Example 1: Heat Equation

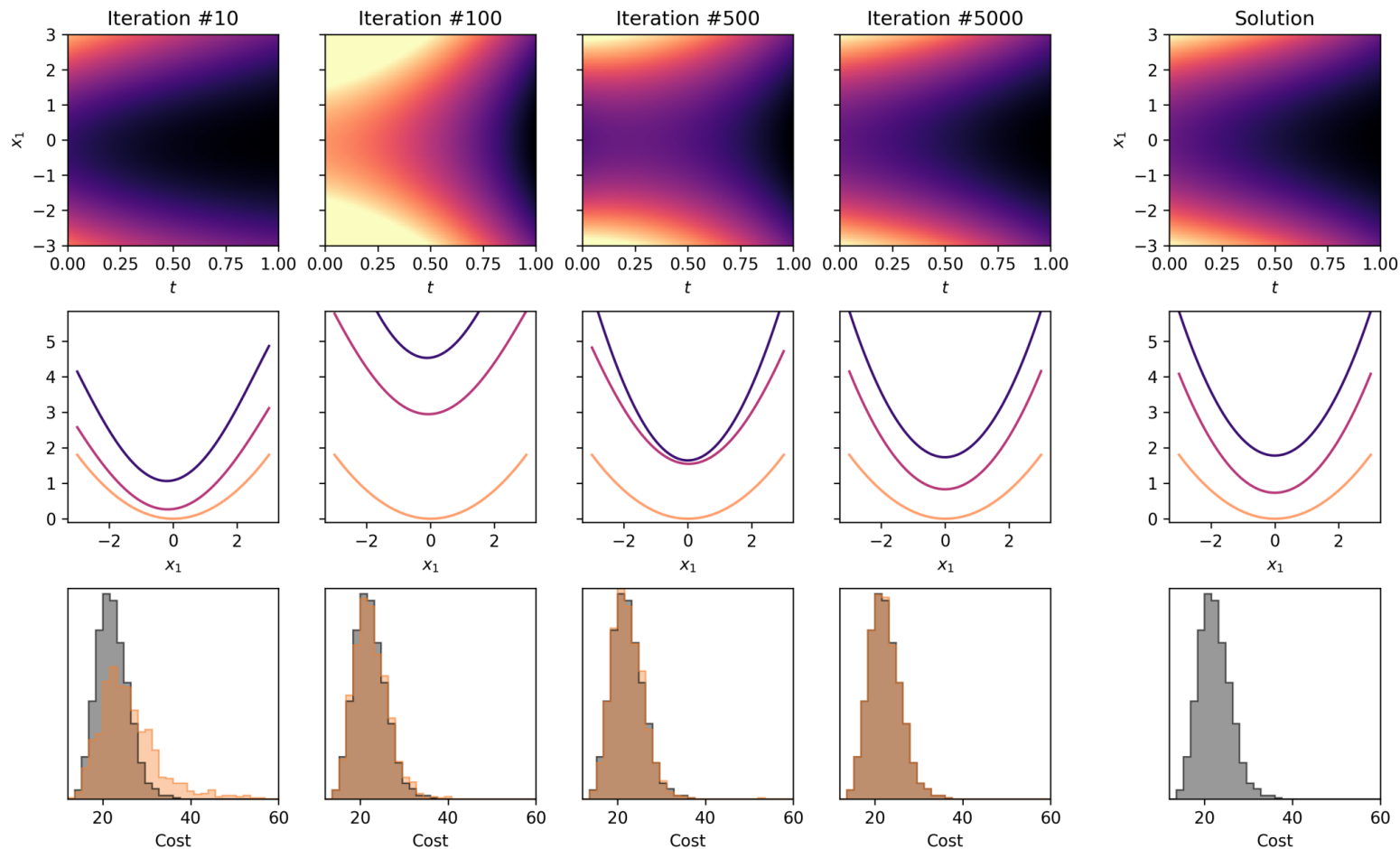


Example 1: Heat Equation



without preconditioning ($\lambda_k \equiv 1$)

Example 2: HJB Equation in 5+1 dims.



Conclusion

- Functional gradient descent is effective, but computing functional gradients is hard
 - Our method allows for estimation of functional gradients with autodiff
- Also in the paper: more general statements, some impossibility results, approximation error of $\mathbb{P}_K$, prebasis for Sobolev spaces, more examples, *etc.*
- **Extensions:** accelerated variants? proximal variants? variance reduction?

Caio Peixoto, Daniel Csillag, Bernardo Costa and Yuri Saporito. *Random Gradient-Free Optimization in Infinite Dimensional Spaces*

See also: Daniel Csillag, Rodrigo Schuller, Pedro Dall'Antonia, Leonidas Guibas, Luiz Velho, Tiago Novello. *Functional Gradient Descent with Adaptive Representations*

Conclusion

- Functional gradient descent is effective, but computing functional gradients is hard
 - Our method allows for estimation of functional gradients with autodiff
- Also in the paper: more general statements, some impossibility results, approximation error of $\mathbb{P}_K$, prebasis for Sobolev spaces, more examples, *etc.*
- **Extensions:** accelerated variants? proximal variants? variance reduction?

Caio Peixoto, Daniel Csillag, Bernardo Costa and Yuri Saporito. *Random Gradient-Free Optimization in Infinite Dimensional Spaces*

See also: Daniel Csillag, Rodrigo Schuller, Pedro Dall'Antonia, Leonidas Guibas, Luiz Velho, Tiago Novello. *Functional Gradient Descent with Adaptive Representations*

supplementary

Infinite Dim. Grad.-Free Descent

Key ingredients:

1. A prebasis for the Hilbert space: $\mathcal{B} = \{b_1, b_2, \dots\}$ such that
$$\mathcal{B} \text{ is linearly independent} \quad \text{and} \quad \text{span}(\mathcal{B}) \text{ is dense in } \mathcal{H}.$$
2. Tail probabilities $t_i := \mathbb{P}[K \geq i]$.

To sample $v|K$:

1. Compute the Gram matrix $[A_K]_{i,j} := \langle b_i, b_j \rangle$
2. Find its Cholesky decomposition $\mathbf{R}_K^\top \mathbf{R}_K = A_K$.
(Equivalent to finding the generalized QR decomposition $Q_K \mathbf{R}_K$ of $[b_1 \ \dots \ b_K]$.)
3. Sample weights $\mathbf{w} \sim \mathcal{N}(0, \mathbf{R}_K^{-1} \mathbf{T}_K^{-1} \mathbf{R}_K^{-\top})$, where $\mathbf{T}_K := \text{diag}(t_1, \dots, t_K)$, and take

$$v = \sum_{i=1}^K w_i b_i.$$