

Towards Efficient Machine Learning: Algorithms, Theoretical Foundations, and Applications

Shaocong Ma

April, 2026

Computer Science, University of Maryland, College Park

AI is revolutionizing our world!



Autonomous Driving



Recommendation System



Risk Management



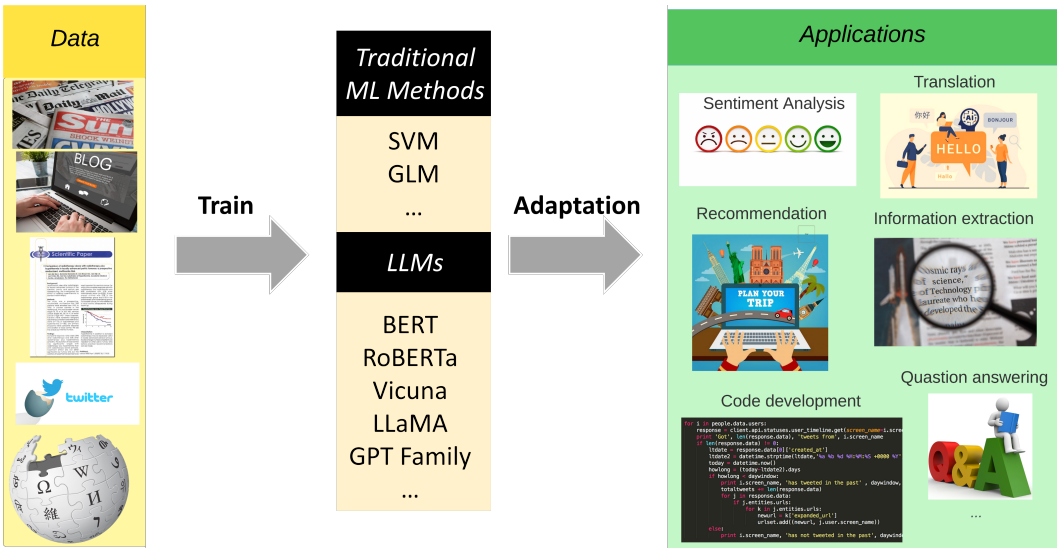
Personalized Education



Disease Detection



Content Creation



Data



Date collected	Pist	Species	Sex	Weight
1/9/79	1	DM	M	40
1/9/79	1	DM	F	36
1/9/79	1	DS	F	133
1/23/79	1	DM	F	39
1/23/79	2	DM	M	43
1/23/79	2	DS	F	144
3/13/79	2	DM	F	51
3/13/79	2	DM	F	44
3/13/79	2	DS	F	146

Train



VLMs

CLIP
BLIP
FLAVA
Sora
DALL-E
Gemini
...

Adaptation



Applications

Image/Vedio generation



Visual search



Data generation

Page	Target	Embedding	Pairwise
1	1.00000000	0.0000	0.0000
2	1.00000000	0.0000	0.0000
3	1.00000000	0.0000	0.0000
4	1.00000000	0.0000	0.0000
5	1.00000000	0.0000	0.0000
6	1.00000000	0.0000	0.0000
7	1.00000000	0.0000	0.0000
8	1.00000000	0.0000	0.0000
9	1.00000000	0.0000	0.0000
10	1.00000000	0.0000	0.0000
11	1.00000000	0.0000	0.0000
12	1.00000000	0.0000	0.0000
13	1.00000000	0.0000	0.0000
14	1.00000000	0.0000	0.0000
15	1.00000000	0.0000	0.0000
16	1.00000000	0.0000	0.0000
17	1.00000000	0.0000	0.0000
18	1.00000000	0.0000	0.0000
19	1.00000000	0.0000	0.0000
20	1.00000000	0.0000	0.0000

Image analysis



Photo editing



VLA Models



RT-2

OpenVLA



Helix

...

Noisy Data



Model Actions

Applications

Embodied AI



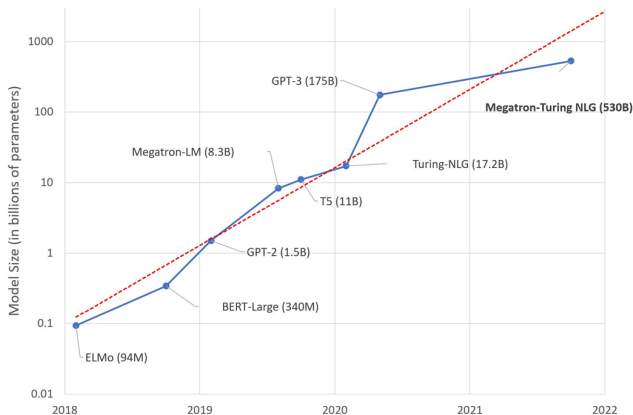
Robots



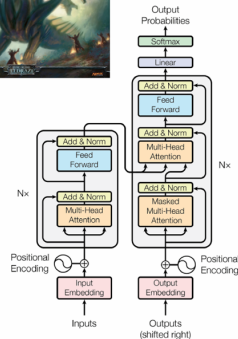
Autonomous Driving



Challenge 1: Increasing Model Sizes



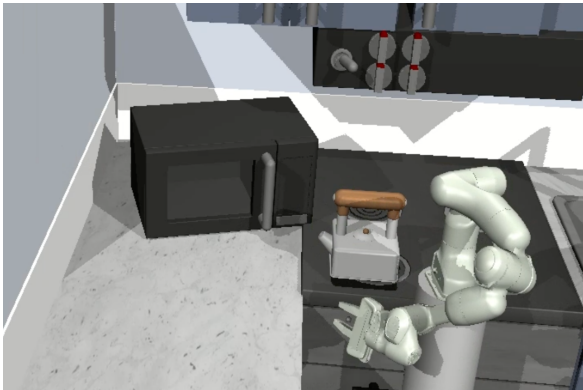
Source: <https://huggingface.co/blog/large-language-models>



Question:

How to save the computational resources?

Challenge 2: Noisy or Adversarial Data

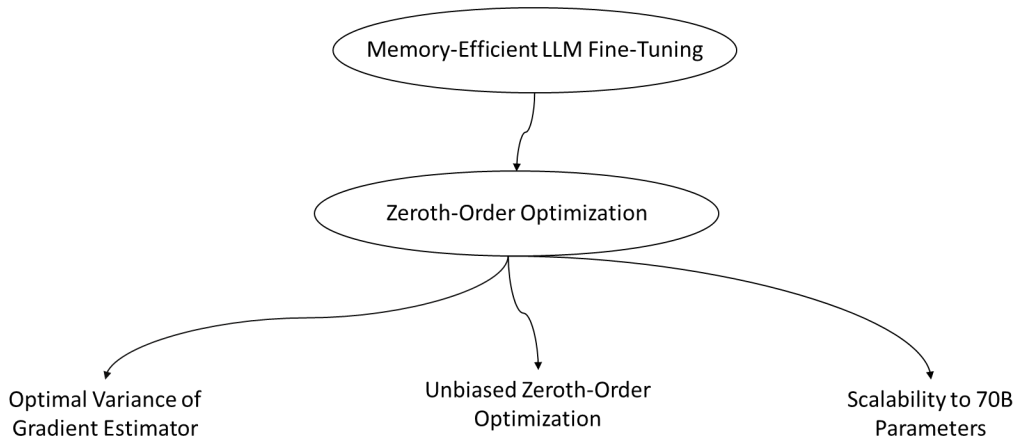


Source: Gu, Shangding, et al. "Robust Gymnasium: A Unified Modular Benchmark for Robust Reinforcement Learning." ICLR 2025.

Question:

How to efficiently train a robust ML model?

Memory-Efficient LLM Fine-Tuning: Outline



Motivation: Why Fine-Tuning A Local LLM?

■ Data privacy:

- Healthcare data
- Trading decision
- Computer control history
- ...



Motivation: Why Fine-Tuning A Local LLM?

■ Data privacy:

- Healthcare data
- Trading decision
- Computer control history
- ...



■ Domain Adaptation:

- Medical Diagnosis
- Chip design
- Specific programming language
- ...



Challenge: The “Memory Wall” Challenge in LLM Fine-tuning

- A single GPU cannot handle backpropagation for entire large models.

Table 1: VRAM Requirements and GPU Configuration

Model Size	First-Order (Full FT)	Est. GPU Setup
OPT-1.3B	≈ 27 GB	$1 \times \text{A100}$
OPT-6.7B	≈ 156 GB	$2 \times \text{A100}$
OPT-13B	≈ 356 GB	$4 \times \text{A100}$
OPT-30B	≈ 633 GB	$8 \times \text{A100}$

Source: Malladi, Sadhika, et al. “Fine-tuning language models with just forward passes.” NeurIPS 2023.

Question:

How to save the computational resources?

Zeroth-Order Optimization (ZOO)

Optimization problem:

$$\min_{x \in \mathbb{R}^d} f(x).$$

Gradient Descent:

$$x' \leftarrow x - \eta \nabla f(x)$$

Notation:

- $f(x)$: The loss function.
- η : The learning rate.

Zeroth-Order Optimization (ZOO)

Optimization problem:

$$\min_{x \in \mathbb{R}^d} f(x).$$

Gradient Descent:

$$x' \leftarrow x - \eta \nabla f(x)$$

Notation:

- $f(x)$: The loss function.
- η : The learning rate.

Memory-Consuming: Deep Neural Network $\frac{\partial f}{\partial x} = \frac{\partial f}{\partial L_N} \frac{\partial L_N}{\partial L_{N-1}} \cdots \frac{\partial L_1}{\partial x}$.

Maintain all intermediate states for backpropagation.

Zeroth-Order Optimization (ZOO)

Core Formula (Two-Point Estimator):

$$\nabla f(x) \approx \hat{\nabla} f(x) = \frac{f(x + \mu v) - f(x)}{\mu} v$$
$$x' \leftarrow x - \eta \hat{\nabla} f(x)$$

Notation:

- $f(x)$: The loss function.
- v : A random perturbation vector (e.g., drawn from a Gaussian distribution $\mathcal{N}(0, I_d)$).
- μ : The perturbation stepsize.

Zeroth-Order Optimization (ZOO)

Core Formula (Two-Point Estimator):

$$\nabla f(x) \approx \hat{\nabla} f(x) = \frac{f(x + \mu v) - f(x)}{\mu} v$$
$$x' \leftarrow x - \eta \hat{\nabla} f(x)$$

Notation:

- $f(x)$: The loss function.
- v : A random perturbation vector (e.g., drawn from a Gaussian distribution $\mathcal{N}(0, I_d)$).
- μ : The perturbation stepsize.

A high-level explanation:

- $\frac{f(x + \mu v) - f(x)}{\mu} > 0 \implies$ Loss increases $\implies$ Move to the opposite direction of v ;
- $\frac{f(x + \mu v) - f(x)}{\mu} < 0 \implies$ Loss decreases $\implies$ Move to the direction of v .

Advantages and Challenges of ZOO

Core Advantage:

- Requires only the **Forward Pass**.
- Memory footprint is comparable to **Inference** only.

Advantages and Challenges of ZOO

Core Advantage:

- Requires only the **Forward Pass**.
- Memory footprint is comparable to **Inference** only.

Key Challenges (Focus of this Talk):

1. **High Variance:** Gradient estimates are volatile/noisy.
2. **Biased:** Finite difference methods rely on approximations.

Advantages and Challenges of ZOO

Core Advantage:

- Requires only the **Forward Pass**.
- Memory footprint is comparable to **Inference** only.

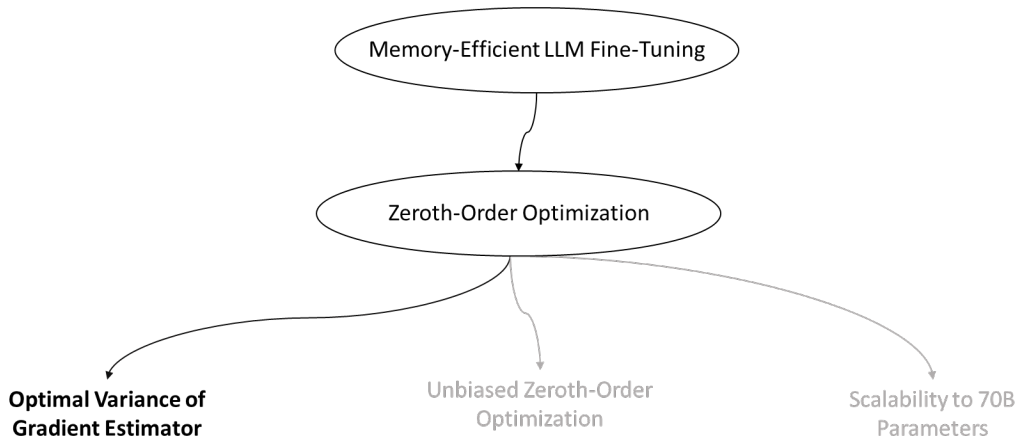
Key Challenges (Focus of this Talk):

1. **High Variance:** Gradient estimates are volatile/noisy.
2. **Biased:** Finite difference methods rely on approximations.

Roadmap: Two theoretical works improving ZOO

1. Derive the condition for achieving the optimal variance.
2. Propose a unbiased gradient estimator family.

Memory-Efficient LLM Fine-Tuning: Outline



Minimum Variance: Directionally Aligned Perturbation (DAP)

Recap: Zero-Order Gradient Estimator

$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

where v is typically sampled from a Gaussian distribution.

Q: Is the standard choice of distribution actually optimal?

Minimum Variance: Directionally Aligned Perturbation (DAP)

Recap: Zero-Order Gradient Estimator

$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

where v is typically sampled from a Gaussian distribution.

Goal: Minimize Estimation Error

Find the optimal distribution V for the perturbation vector v :

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \right\|^2, \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [v v^\top] = I_d. \end{aligned}$$

Minimum Variance: Directionally Aligned Perturbation (DAP)

Recap: Zero-Order Gradient Estimator

$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

where v is typically sampled from a Gaussian distribution.

Goal: Minimize Estimation Error

Find the optimal distribution V for the perturbation vector v :

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \right\|^2, \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [vv^\top] = I_d. \end{aligned}$$

Optimization Challenges:

- **Functional space:** Optimization is taken over all probability distributions.
- **Constraints with an empty interior:** The empty interior precludes the use of Interior Point Methods.

Minimize Estimation Error (Part I)

Find the optimal distribution V for the perturbation vector v :

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \right\|^2, \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [vv^\top] = I_d. \end{aligned}$$

Taylor Expansion:

$$\begin{aligned} f(x + \mu v) - f(x) &\approx \mu \nabla f(x)^\top v + \mu^2 \cdot \text{Bias} \\ \frac{f(x + \mu v) - f(x)}{\mu} v &\approx vv^\top \nabla f(x) + \mu \cdot \text{Bias} \\ \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) &\approx (vv^\top - I_d) \nabla f(x) + \underbrace{\mu \cdot \text{Bias}}_{\text{Ignored}} \end{aligned}$$

Definition of Bias: $\mu \cdot \text{Bias} := \mathbb{E}[\hat{\nabla} f(x)] - \nabla f(x)$

Minimize Estimation Error (Part II)

Find the optimal distribution V for the perturbation vector v :

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} - \nabla f(x) \right\|^2, \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [vv^\top] = I_d. \end{aligned}$$

Plug in the objective function:

$$\mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} - \nabla f(x) \right\|^2 = \mathbb{E}_{v \sim V} \nabla f(x)^\top (vv^\top) \nabla f(x) - \|\nabla f(x)\|^2$$

Minimize Estimation Error (Part II)

Find the optimal distribution V for the perturbation vector v :

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \right\|^2, \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [v v^\top] = I_d. \end{aligned}$$

Plug in the objective function:

$$\mathbb{E}_{v \sim V} \left\| \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \right\|^2 = \mathbb{E}_{v \sim V} \nabla f(x)^\top (v v^\top) \nabla f(x) - \|\nabla f(x)\|^2$$

Simplified Objective:

$$\begin{aligned} \min_V \quad & \mathbb{E}_{v \sim V} a^\top (v v^\top) a \\ \text{s.t.} \quad & \mathbb{E}_{v \sim V} [v v^\top] = I_d. \end{aligned}$$

Minimize Estimation Error (Part II)

We analytically solve this functional optimization problem.

Derive the lower bound + Tell when the lower bound is achieved

Theorem

Let v be a random vector following the distribution V with $\mathbb{E}_{v \sim V} vv^\top = I_d$ and $a \in \mathbb{R}^d$ be a fixed vector. Then

$$d\|a\|^2 \leq \mathbb{E}_{v \sim V} a^\top (vv^\top)^2 a \leq d\|a\|^2 + \frac{\|a\|^2}{2} \left(\rho_V + \sqrt{\rho_V^2 + 4(d-1)\rho_V} \right)$$

where $\rho_V := \mathbb{E}\|v\|^4 - d^2$.

Ma, S. & Huang, H. Revisiting Zeroth-Order Optimization: Minimum-Variance Two-Point Estimators and Directionally Aligned Perturbations. ICLR, 2025. **Spotlight**

Minimize Estimation Error (Part II)

We analytically solve this functional optimization problem.

Derive the lower bound + Tell when the lower bound is achieved

Theorem

Let v be a random vector following the distribution V with $\mathbb{E}_{v \sim V} vv^\top = I_d$ and $a \in \mathbb{R}$ be a fixed vector. Then

$$d\|a\|^2 \leq \mathbb{E}_{v \sim V} a^\top (vv^\top)^2 a \leq d\|a\|^2 + \frac{\|a\|^2}{2} \left(\rho_V + \sqrt{\rho_V^2 + 4(d-1)\rho_V} \right)$$

where $\rho_V := \mathbb{E}\|v\|^4 - d^2$.

Ma, S. & Huang, H. Revisiting Zeroth-Order Optimization: Minimum-Variance Two-Point Estimators and Directionally Aligned Perturbations. ICLR, 2025. Spotlight

What kind of distribution actually achieves this lower bound?

DAPs: Directionally Aligned Perturbation

We analytically solve this functional optimization problem.

(Equality Condition)

■ Constant Magnitude Perturbations:

- $\mathbb{E}_{v \sim V}[vv^\top] = I_d$.
- $\|v\|$ is fixed.

■ Directionally Aligned Perturbations (DAPs):

- $\mathbb{E}_{v \sim V}[vv^\top] = I_d$.
- $\nabla f(x)^\top v$ is fixed.

$\implies$ Both estimators achieve the **minimum variance**.

$\implies$ DAPs have some nice properties.

Traditional Methods Cannot Identify the Important Directions

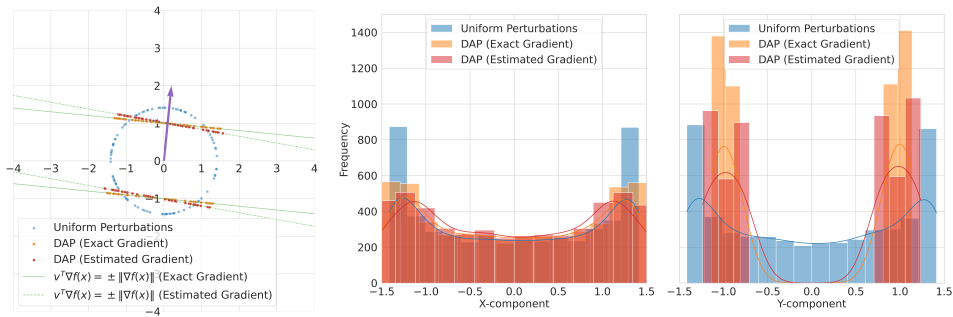


Figure 1: Illustration of the *directional alignment property* of DAP in $d = 2$ with estimating the gradient of $f(x) = x_1^2 + x_2^2$ at $x = \begin{bmatrix} 0.1 & 1 \end{bmatrix}^T$. Traditional estimator is **symmetric**, but we need a **non-symmetric** estimator.

Traditional Methods Cannot Identify the Important Directions

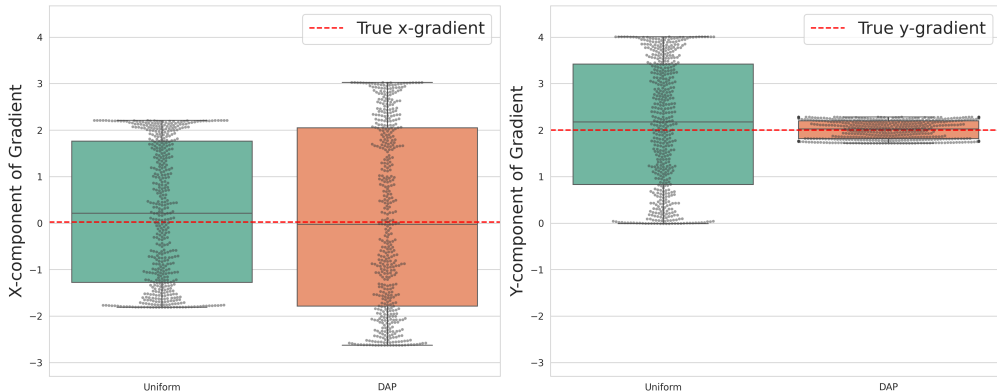


Figure 2: Comparison of gradient estimation performance with estimating the gradient of $f(x) = x_1^2 + x_2^2$ at $x = \begin{bmatrix} 0.1 & 1 \end{bmatrix}^T$ between uniform random perturbations and DAPs. The **non-symmetric** estimator is more accurate in the direction with larger gradient.

Applications in LLM Fine-Tuning

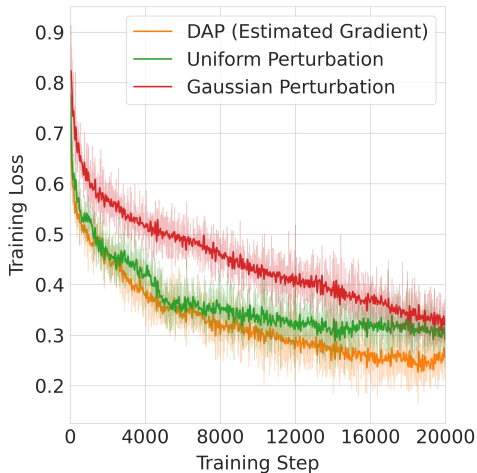


Figure 3: Comparison of training loss curves among different random perturbations on **Large Language Model Fine Tuning**.

Applications in Scientific Optimization

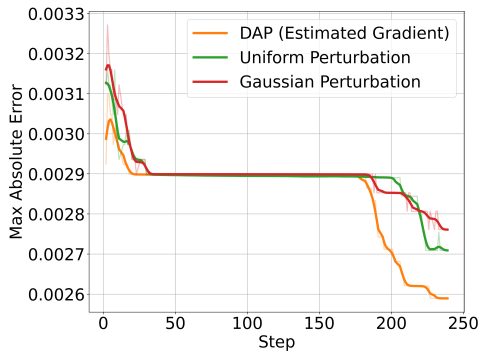


Figure 4: Comparison of training loss curves among different random perturbations on **Mesh Optimization for the Physical Numerical Solver**.

Summary

Derived the optimal distribution of v to achieve the minimum variance.

$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

Summary

Derived the optimal distribution of v to achieve the minimum variance.

$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

(Taylor approximation)

$$\frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \approx (vv^\top - I_d) \nabla f(x) + \underbrace{\mu \cdot \text{Bias}}_{\text{Ignored}}$$

Summary

Derived the optimal distribution of v to achieve the minimum variance.

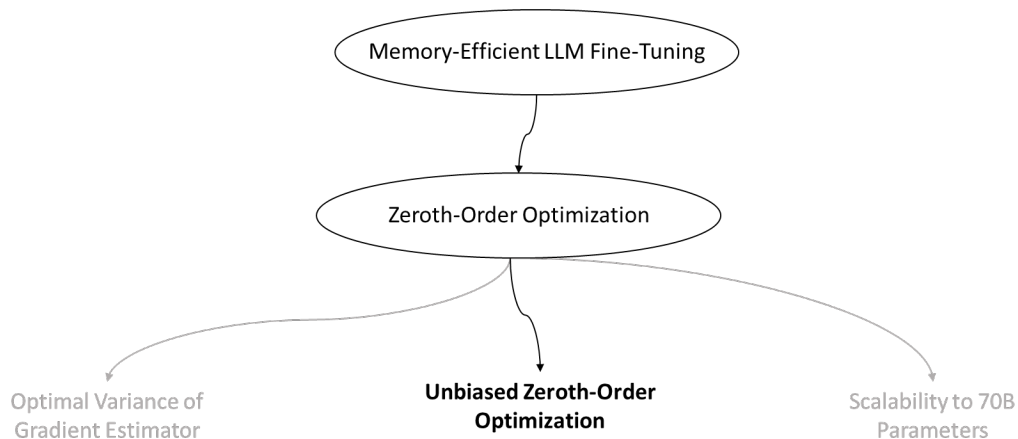
$$\nabla f(x) \approx \hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v$$

(Taylor approximation)

$$\frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) \approx (vv^\top - I_d) \nabla f(x) + \underbrace{\mu \cdot \text{Bias}}_{\text{Ignored}}$$

Is it possible to eliminate the bias completely?

Memory-Efficient LLM Fine-Tuning: Outline



Inherent Bias of Two-Point Estimator

Recap: Zero-Order Gradient Estimator

$$\hat{\nabla} f(x) := \frac{f(x + \mu v) - f(x)}{\mu} v.$$

Taylor Expansion:

$$\begin{aligned} f(x + \mu v) - f(x) &\approx \mu \nabla f(x)^\top v + \mu^2 \cdot \text{Bias} \\ \frac{f(x + \mu v) - f(x)}{\mu} v &\approx v v^\top \nabla f(x) + \mu \cdot \text{Bias} \\ \frac{f(x + \mu v) - f(x)}{\mu} v - \nabla f(x) &\approx (v v^\top - I_d) \nabla f(x) + \underbrace{\mu \cdot \text{Bias}}_{\text{Not Ignored?}} \end{aligned}$$

▷ When μ is large, the two-point estimator exhibits significant bias.

Unbiased Estimator based on Multi-Level Monte Carlo

Unbiased zeroth-order gradient estimator using only function evaluations.

- *Step 1.* Directional derivative along the direction v .

$$\nabla_v f(x) = \lim_{\mu \rightarrow 0} \frac{f(x + \mu v) - f(x)}{\mu}.$$

- *Step 2.* Telescoping series. Let $\mu_n \rightarrow 0$.

$$\begin{aligned} \nabla_v f(x) &= \frac{f(x + \mu_1 v) - f(x)}{\mu_1} \\ &\quad + \sum_{n=1}^{\infty} \left[\frac{f(x + \mu_{n+1} v) - f(x)}{\mu_{n+1}} - \frac{f(x + \mu_n v) - f(x)}{\mu_n} \right]. \end{aligned}$$

Unbiased Estimator based on Multi-Level Monte Carlo

- Step 3. Expectation representation.

Let $\sum_n p_n = 1$ and $0 < p_n < 1$.

$$\begin{aligned}\nabla_v f(x) = & \sum_{n=1}^{\infty} p_n \left[\frac{f(x + \mu_1 v) - f(x)}{\mu_1} \right. \\ & \left. + \frac{1}{p_n} \left(\frac{f(x + \mu_{n+1} v) - f(x)}{\mu_{n+1}} - \frac{f(x + \mu_n v) - f(x)}{\mu_n} \right) \right].\end{aligned}$$

Then $\nabla_v f(x)$ can be represented as

$$\mathbb{E}_{n \sim \{p_n\}_{n=1}^{\infty}} \left[\frac{f(x + \mu_1 v) - f(x)}{\mu_1} + \frac{1}{p_n} \left(\frac{f(x + \mu_{n+1} v) - f(x)}{\mu_{n+1}} - \frac{f(x + \mu_n v) - f(x)}{\mu_n} \right) \right].$$

$\Rightarrow$ Unbiased Estimator Family

Unbiased Estimator Family

- P_4 -Estimator:

$$P_4(n, v) := \frac{f(x + \mu_1 v) - f(x)}{\mu_1} + \frac{1}{p_n} \left(\frac{f(x + \mu_{n+1} v) - f(x)}{\mu_{n+1}} - \frac{f(x + \mu_n v) - f(x)}{\mu_n} \right)$$

- P_3 -Estimator:

$$P_3(n, v) := \frac{f(x + \mu_1 v) - f(x)}{\mu_1} U_2 + \frac{1}{p_n} \left(\frac{f(x + \mu_{n+1} v) - f(x)}{\mu_{n+1}} - \frac{f(x + \mu_n v) - f(x)}{\mu_n} \right) (1 - U_2)$$

where $U_2 \sim \text{Uniform}(\{0, 1\})$.

- We can also define P_2 -Estimator and P_1 -Estimator.

Ma, S. & Huang, H. On the Optimal Construction of Unbiased Gradient Estimators for Zeroth-Order Optimization. NeurIPS, 2025.

Spotlight

Unbiased Estimator Family: Variance

Theorem

Suppose that $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is second-order continuously differentiable and has L -Lipschitz continuous gradient. $\sum_{n=1}^{\infty} \mu_n < \infty$ and $V \sim \sqrt{d} \text{Uniform}(\mathbb{S}^{d-1})$. Define

$$\mu := \mu_1, \quad \varrho := \sum_{n=1}^{\infty} \frac{(\mu_{n+1} - \mu_n)^2}{p_n}, \quad \text{and} \quad \varphi := \sum_{n=1}^{\infty} \frac{\mu_n^2}{p_n}.$$

Then

$$\text{Var}[P_4(n, v)v] \leq (d-1)\|\nabla f(x)\|^2 + \frac{3L^2}{4}d^3\mu^2 + \frac{L^2d^3}{2}\varrho,$$

$$\text{Var}[P_3(n, v)v] \leq \text{Var}[P_4(n, v)v] + \frac{L^2}{8}d^3\mu^2 + \frac{L^2d^3}{8}\varrho.$$

▷ This variance results in the optimal oracle complexity.

Unbiased Estimator Family: Variance

Theorem

Let $\{\mu_n\}_{n=1}^{\infty}$ be a positive, decreasing sequence with $\sum_{n=1}^{\infty} \mu_n < \infty$, and let $\{p_n\}_{n=1}^{\infty}$ be a Probability Mass Function. Then

$$\varrho \geq \mu^2.$$

The equality holds if and only if

$$p_n = \frac{\mu_n - \mu_{n-1}}{\mu}.$$

▷ We obtain a simple and elegant relation to derive the optimal sequence $\{p_n\}$ and $\{\mu_n\}$.

- Geometric P_k -Estimator: $n \sim \text{Geom}(c)$ and $\mu_n = \mu_1 c^{n-1}$.
- Zipf's P_k -Estimator: $n \sim \text{Zipf}(s)$ ($s > 1$) and $\mu_n = \mu_1 \left[1 - \left(\sum_{j=1}^{n-1} \frac{1}{j^s} \right) / \zeta(s) \right]$.

Unbiased Estimator Leads to Better Accuracy

The quadratic loss $f_{\text{reg}} : \mathbb{R}^d \rightarrow \mathbb{R}$ and the logistic loss $f_{\text{cls}} : \mathbb{R}^d \rightarrow \mathbb{R}$:

$$f_{\text{reg}}(x) = x^\top A^\top A x, \quad f_{\text{cls}}(x) = \frac{1}{n} \sum_{i=1}^n \log(1 + \exp(-b_i \cdot (a_i^\top \cdot x))).$$

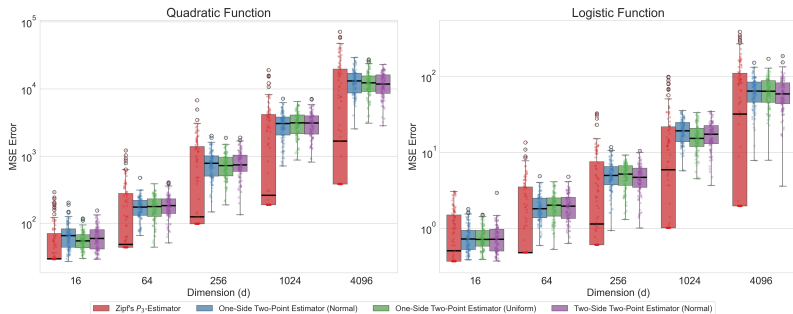


Figure 5: The MSE error (Left: f_{reg} , Right: f_{cls}) of different estimators.

Applications in LLM Fine-Tuning

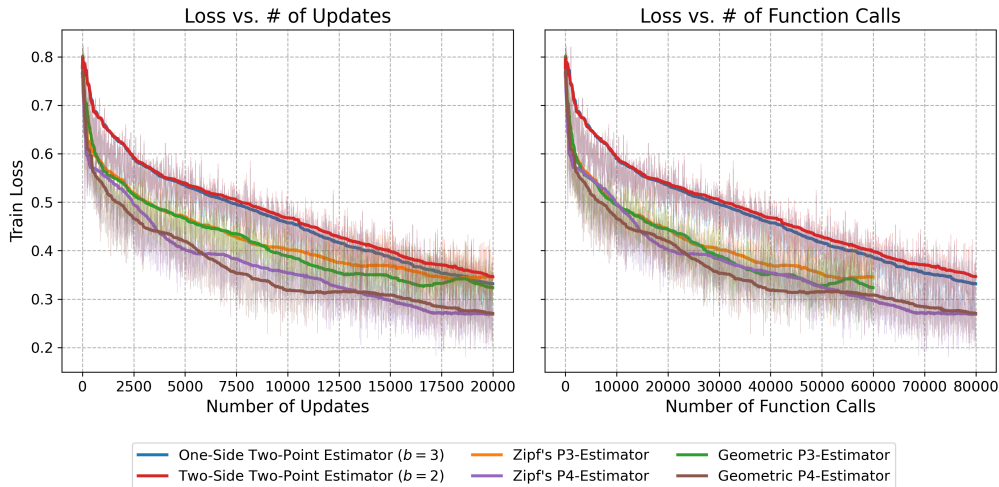


Figure 6: Fine-tuning the OPT-1.3B model on SST-2 using different gradient estimators.

Applications in LLM Fine-Tuning

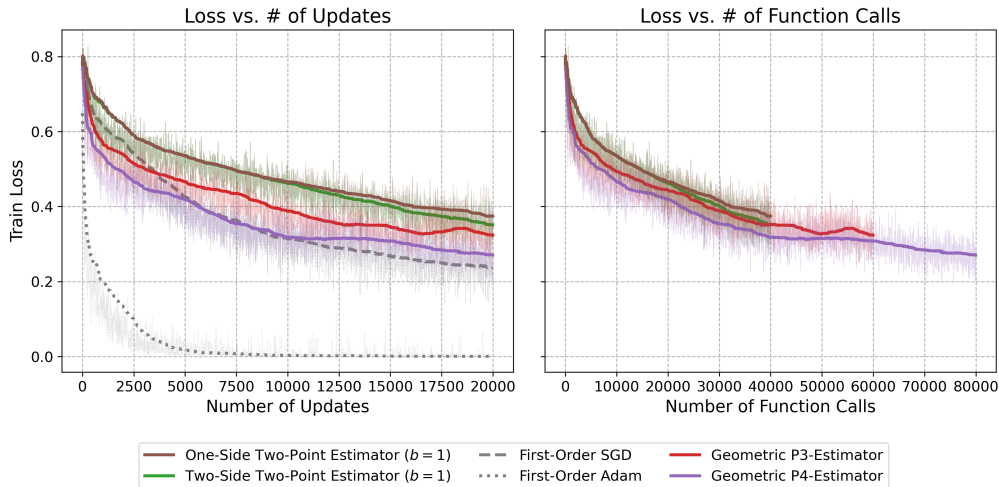


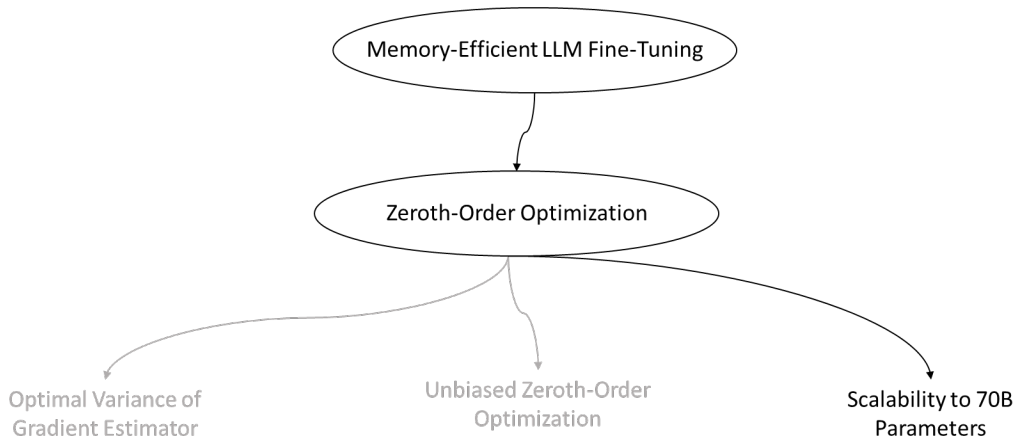
Figure 7: Fine-tuning the OPT-1.3B model on SST-2 using different gradient estimators.

- Constructed the family of unbiased zeroth-order gradient estimators.
- Provided the theoretical framework to minimize its variance.
- Validated its performance in synthetic and LLM experiments.

- Constructed the family of unbiased zeroth-order gradient estimators.
- Provided the theoretical framework to minimize its variance.
- Validated its performance in synthetic and LLM experiments.

Can we further scale up the Zeroth-Order Optimization method?

Memory-Efficient LLM Fine-Tuning: Outline



Geometric Constraints in Quantized LLMs

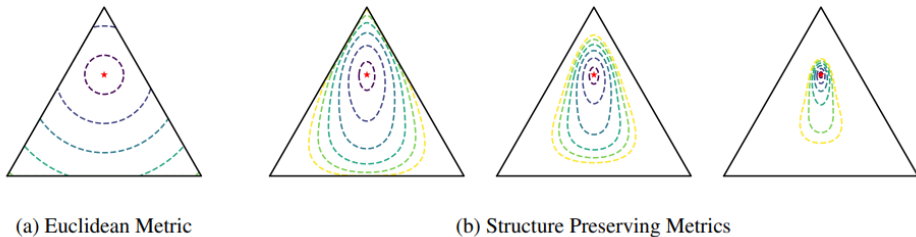


Figure 8: Visualization of different metrics on the probability simplex.

▷ Scale parameters in quantized LLMs form a geodesically incomplete Riemannian manifold. We propose structure preserving metrics to handle this issue.

Ma, S. & Huang, H. Riemannian Zeroth-Order Gradient Estimation with Structure-Preserving Metrics for Geodesically Incomplete Manifolds. ICLR, 2026.

Memory-Efficient Fine-Tuning of Quantized LLMs

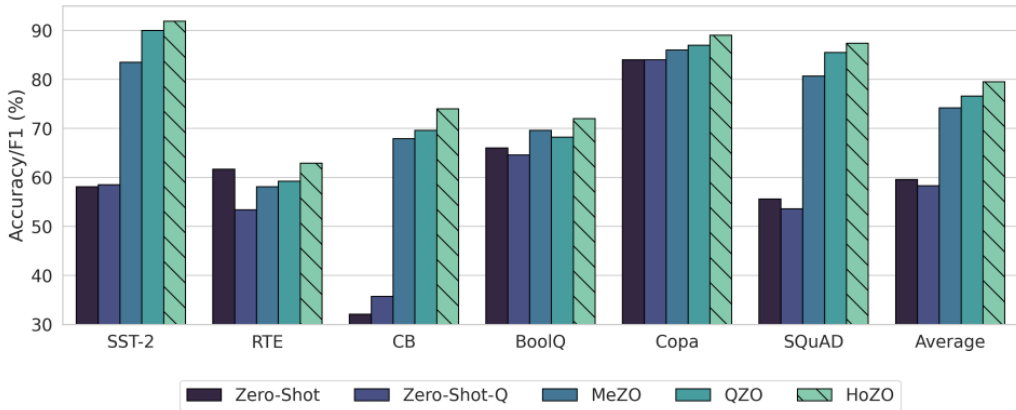


Figure 9: On the INT4 Llama-2-7B model across 6 downstream tasks, HoZO achieves consistently better performance than all baselines.

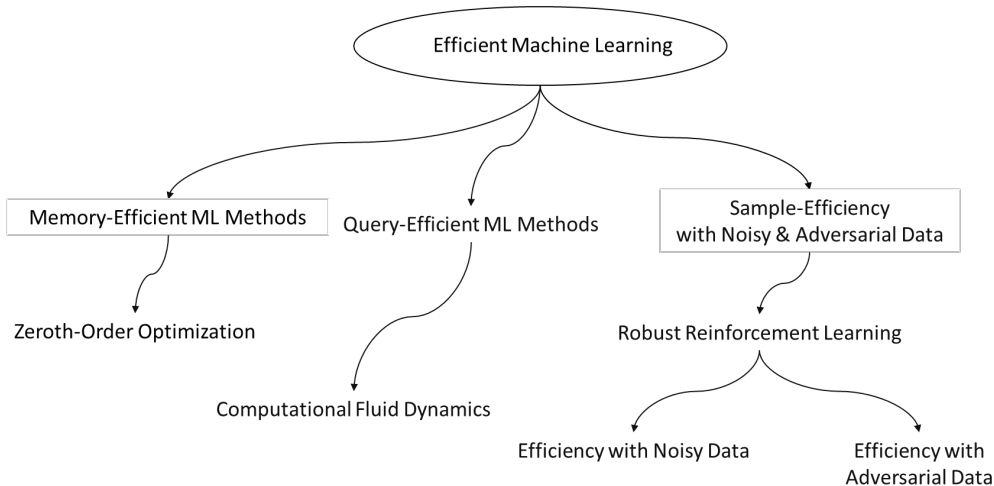
Memory-Efficient Fine-Tuning of Quantized LLMs

Model	Method	Model Precision	SST-2	RTE	CB	BoolQ	Copa	SQuAD	Average	
Llama-2-7b	Zero-Shot	16 bits	58.1*	61.7*	32.1*	66.0*	84.0	55.6*	59.6	-
	Zero-Shot-Q	4 bits	58.5*	53.4*	35.7*	64.6*	84.0	53.6*	58.3	↓1.3
	MeZO	16 bits	83.5*	58.1*	67.9*	69.6*	86.0	80.7*	74.2	↑14.6
	QZO	4 bits	90.0*	59.2*	69.6*	68.2*	87.0	85.5*	76.6	↑17.0
	HoZO	4 bits	<u>91.9</u>	<u>62.9</u>	<u>74.0</u>	<u>72.0</u>	<u>89.0</u>	<u>87.4</u>	<u>79.5</u>	↑19.9
Llama-2-13b	Zero-Shot	16 bits	61.1	50.9	44.0	74.1	89.0	63.7	63.8	-
	Zero-Shot-Q	4 bits	60.0	47.3	47.0	74.7	87.0	63.1	63.2	↓0.6
	MeZO	16 bits	90.7	58.5	<u>77.0</u>	81.6	<u>92.0</u>	87.5	81.2	↑17.4
	QZO	4 bits	91.9	62.8	<u>77.0</u>	<u>82.4</u>	<u>92.0</u>	89.2	82.5	↑18.7
	HoZO	4 bits	<u>92.4</u>	<u>68.6</u>	75.0	81.6	<u>92.0</u>	<u>89.3</u>	<u>83.2</u>	↑19.4
Llama-2-70b	Zero-Shot-Q	4 bits	56.4	60.6	47.0	74.7	92.0	71.4	67.0	-
	QZO	4 bits	90.6	<u>80.8</u>	82.0	<u>83.8</u>	93.0	90.4	86.8	↑19.8
	HoZO	4 bits	<u>91.5</u>	79.1	<u>83.0</u>	<u>83.8</u>	<u>95.0</u>	<u>91.2</u>	<u>87.3</u>	↑20.3

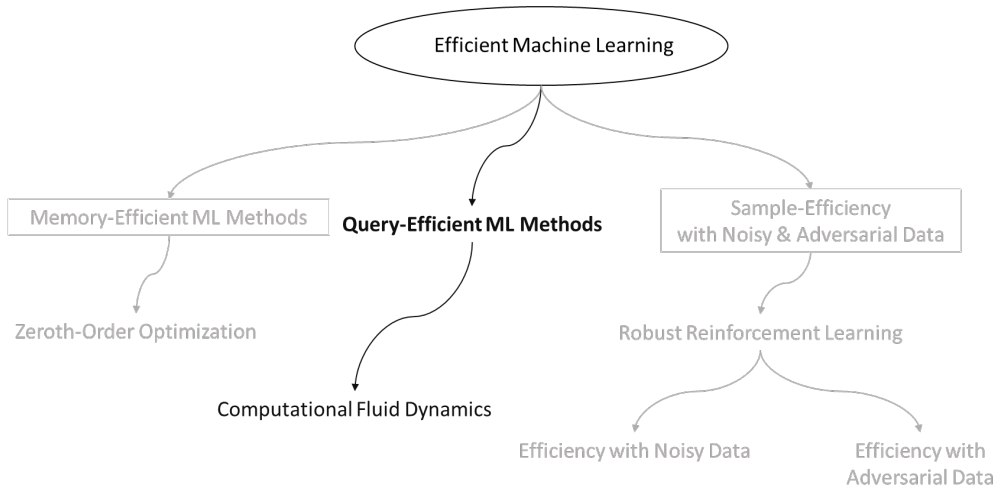
Ma, S. & Huang, H. Hyper-Octant Zeroth-Order Optimization: Fine-Tuning Quantized LLMs on the Positive Orthant Manifold.

Submitted to ICML, 2026.

My Other Research Summary



Query-Efficient ML Methods



Query-Efficient ML Methods

Motivation: A single query may take minutes or hours long.

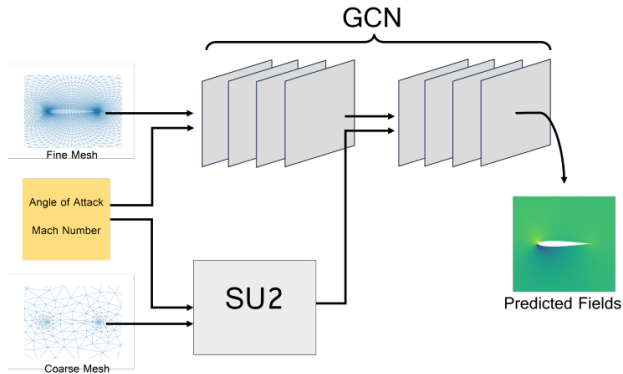


Figure 10: The CFD-GCN model. The PDE solver can be slow. How to improve the query efficiency?

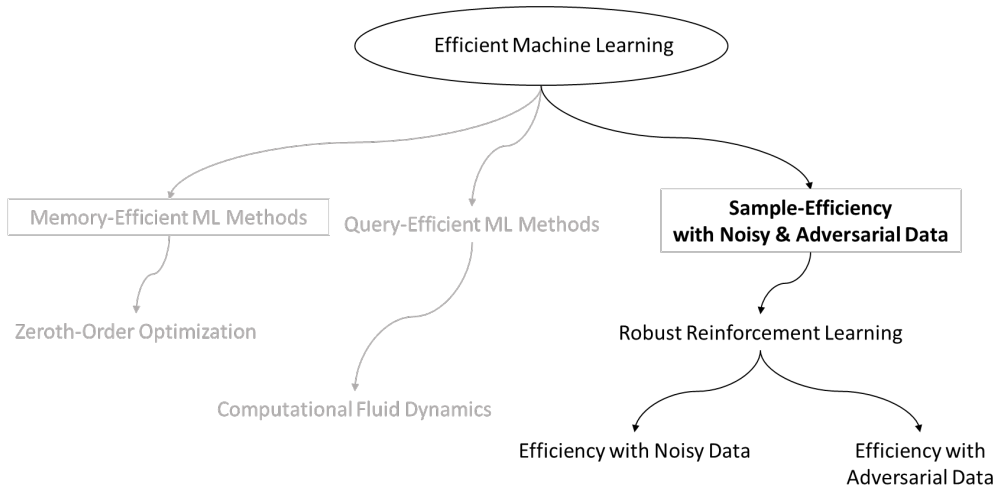
Our Solution: Estimate the non-differentiable part only.

$$\nabla_x f(g(x)) = \frac{\partial f}{\partial g} \times \underbrace{\frac{\widehat{\partial g}}{\partial x}}_{\text{Est. this}}$$

Ma, S., Diffenderfer, J., Kailkhura, B., & Zhou, Y. Deep learning of PDE correction and mesh adaption without automatic differentiation. Machine Learning, 2025.

Ma, S., Diffenderfer, J., Kailkhura, B., & Zhou, Y. End-to-End Mesh Optimization of a Hybrid Deep Learning Black-Box PDE Solver. NeurIPS, 2023 (ML4PS Workshop).

Sample-Efficiency with Noisy & Adversarial Data



Sample-Efficiency with Noisy Data

Motivation: Overcome sample-inefficiency caused by non-IID Markovian data.

Examples:

- Agent-environment interactions in RL.
- Continuous sensor streams in autonomous driving or IoT.
- Network traffic or queueing systems.

Sample-Efficiency with Noisy Data

Our Solution:

- Data Reshuffling.

Ma, S., & Zhou, Y. Understanding the impact of model incoherence on convergence of incremental sgd with random reshuffle. ICML, 2020.

- Larger Batch Size.

Ma, S., Chen, Z., Zhou, Y., Ji, K., & Liang, Y. Data sampling affects the complexity of online SGD over dependent data. UAI, 2022.

- Variance Reduction.

Ma, S., & Zhou, Y. Variance-Reduced Off-Policy TDC Learning: Non-Asymptotic Convergence Analysis. NeurIPS, 2020.

Ma, S., Chen, Z., & Zhou, Y. Greedy-GQ with Variance Reduction: Finite-time Analysis and Improved Complexity. ICLR, 2021.

Sample-Efficiency with Adversarial Data

Motivation:

- Environment mismatch between the simulation and training environments.
- Multi-agent systems.
- Adversarial attacking.

Question:

How to efficiently train a robust ML model?

Our Solution:

- Environment mismatch between the simulation and training environments.

Ma, S., & Huang, H. Robust Reinforcement Learning in Finance: Modeling Market Impact with Elliptic Uncertainty Sets. NeurIPS, 2025.

- Multi-agent systems.

Ma, S., Chen, Z., Zou, S. & Zhou, Y. Decentralized Robust V-Learning for Solving Markov Games with Model Uncertainty. JMLR, 2023.

- Adversarial attacking.

Chen, Z., Ma, S., & Zhou, Y. Accelerated Proximal Alternating Gradient-Descent-Ascent for Nonconvex Minimax Machine Learning. IEEE ISIT, 2022.

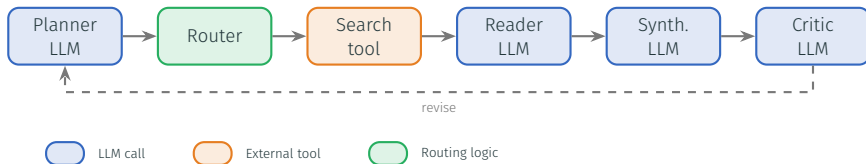
Future Work

Efficient Fine-Tuning of Agentic Systems: Motivation

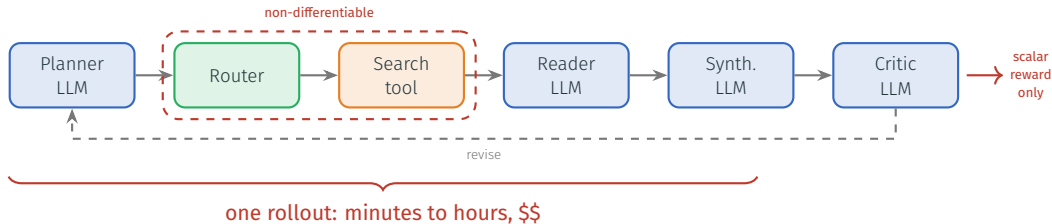


From research demos to deployed systems at scale.

Anatomy of an agent (e.g., a research agent)



Efficient Fine-Tuning of Agentic Systems: Challenges



Non-differentiable

Black-box tools;
discrete routing;
no gradients to BP.

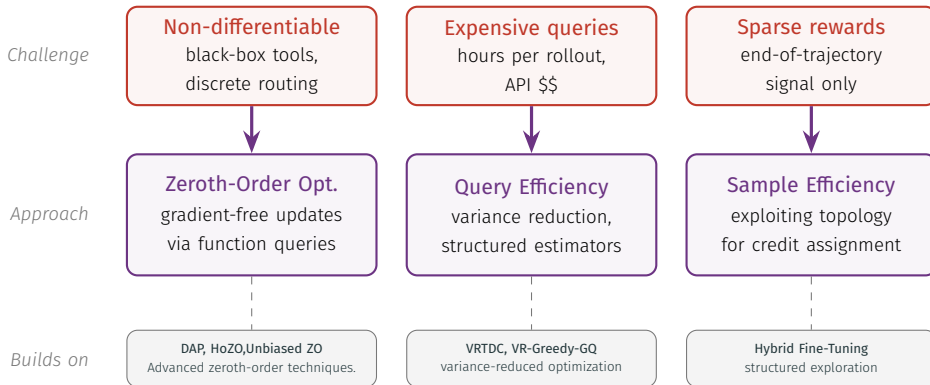
Expensive queries

Hours per rollout;
API \$\$ per call;
eval is the bottleneck.

Sparse rewards

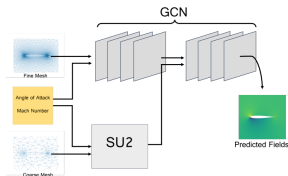
End-of-trajectory signal;
no per-step credit.

Efficient Fine-Tuning of Agentic Systems: From Challenges to Approaches



A unified zeroth-order perspective on fine-tuning agentic systems.

Efficient and Robust Machine Learning for Science/Healthcare



Motivation:

- The scientific software can be slow and non-differentiable $\Rightarrow$ Query Efficiency.
- Scientific/clinical measurements are noisy and heterogeneous $\Rightarrow$ Robustness.

Future directions:

- **Turbulence:** More challenging CFD problems.
- **Extension to protein prediction:** Molecular Dynamics + DNN, Mathematical Biology.
- **Catastrophic consequences in healthcare task:** Robustness under uncertainty.

Future Research Directions and Strategies

- Continue collaborations with
 - UMD, Univ. of Utah, USC, TAMU, UTA, OSU, ASU, Buffalo, *etc.*
 - Lawrence Livermore National Lab, MD Anderson Cancer Center, *etc.*
- Seek more collaborations with
 - Utah State University Math & Stat (Numerical PDE, Mathematical Biology, High-Dimensional Statistics), Physics, Chemistry & Biochemistry (CFD, AI4Science)
 - USU Space Dynamics Laboratory (SDL)
 - College of Veterinary Medicine, Ecology Center, University of Utah Health, Intermountain Health, Huntsman Cancer Institute (AI4Healthcare)
 - Local Industrials (Campbell Scientific, iFIT, Autoliv, Schreiber Foods, *etc.*)
 - General Industrials (Google, Amazon, UnitedHealthcare, Lightning AI, *etc.*)

Successful Funding Experiences in Proposal Writing

AI for Healthcare

- A Real-World Test Bed for Post-Market Surveillance and Stress Testing of AI-Enabled Imaging Tools
2025–2027, FDA, **\$1.2M**.
- Ultrascale Machine Learning to Empower Discovery in Alzheimer's Disease Biobanks
2026–2031 (Recommended), NIH center grant, **\$15M**.

Robust Machine Learning

- Advanced AI Framework to Improve Understanding and Prediction of Wildland Fire
2026–2028, NSF-RISE, **\$1,856,577**.

Other Writing Experiences

- NSF MFAI, NSF GCR, NSF PCL, NSF SLES, NSF/NIH SCH, NIH R01s.

My Future Research Funding Plan

- First several proposals: NSF MFAI, NSF Early Career Development, NSF-CISE, *etc.*
- Collaborate with colleagues at USU to seek: NSF Future CoRe (larger budget), NSF ACED, NSF AIMing, NIH-NIA R01, NIH-NIGMS R01, NIH-NIBIB R01, *etc.*

Teaching Experiences

- Teaching Assistant at UCSB
 - PSTAT 5A: Statistics
 - PSTAT 5LS: Statistics for Life Science
 - PSTAT 109: Statistics for Economics
 - PSTAT 172A: Actuarial Statistics
 - PSTAT 175: Survival Analysis
- Teaching Assistant at Univ. of Utah
 - ECE 3500: Fundamentals of Signals and Systems
- Co-teach at UMD
 - CMSC422: Introduction to Machine Learning

I can teach various courses:

- Lower-Level Undergraduate Courses:
 - Calculus & Multi-variable Calculus
 - Linear Algebra
 - Numerical Algorithms
 - Discrete Mathematics
 - Machine Learning & AI
- Upper-Level Undergraduate or Graduate Courses:
 - Advanced Machine Learning & Statistical Learning
 - Advanced Stochastic Algorithms
 - Reinforcement Learning
 - Actuarial Mathematics

Thank You!