

# Full-Batch Gradient Descent Outperforms One-Pass SGD: Sample Complexity Separation in Single-Index Learning

Filip Kovačević<sup>1</sup>, Hong Chang Ji<sup>2,\*</sup>, Denny Wu<sup>3,\*</sup>, Mahdi Soltanolkotabi<sup>4,\*</sup>, Marco Mondelli<sup>1,\*</sup>

*<sup>1</sup>Institute of Science and technology Austria*

*<sup>2</sup>Sung Kyun Kwan University*

*<sup>3</sup>New York University and Flatiron Institute*

*<sup>4</sup>University of Southern California*

\*Equal contribution

# Problem setup

| <u>Dataset</u>                       | <u>True model</u>     | <u>Hypothesis class</u>                       | <u>Empirical loss</u>                                                               |
|--------------------------------------|-----------------------|-----------------------------------------------|-------------------------------------------------------------------------------------|
| $\mathcal{D} = \{x_i, y_i\}_{i=1}^n$ | $y = f_{\theta^*}(x)$ | $\{f_{\theta} \mid \theta \in \mathbb{R}^d\}$ | $\mathcal{L}(\theta, \mathcal{D}) = \frac{1}{n} \sum_{i=1}^n l(\theta, (x_i, y_i))$ |

Main question

When does reusing data help **optimization**?

Optimization

Gradient descent

Full batch GD

VS

One pass SGD

$$\theta_t = \theta_{t-1} - \eta \nabla_{\theta} \mathcal{L}(\theta_{t-1}, \mathcal{D})$$

$$\theta_t = \theta_{t-1} - \eta \nabla_{\theta} l(\theta_{t-1}, (x_{t-1}, y_{t-1}))$$

# Data model

| <u>Dataset</u>                       | <u>True model</u>                                                            | <u>Hypothesis class</u>                                                  | <u>Empirical loss</u>                                                                                |
|--------------------------------------|------------------------------------------------------------------------------|--------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|
| $\mathcal{D} = \{x_i, y_i\}_{i=1}^n$ | $y = \sigma(\langle x, \theta^* \rangle)$<br>$\theta^* \in \mathbb{S}^{d-1}$ | $\{\sigma(\langle x, \theta \rangle) \mid \theta \in \mathbb{S}^{d-1}\}$ | $\mathcal{L}(\theta, \mathcal{D}) = -\frac{1}{n} \sum_{i=1}^n y_i \sigma(\langle x, \theta \rangle)$ |

## Goal

weak/strong recovery

$$|\langle \theta^*, \theta(\mathcal{D}) \rangle| \rightarrow c > 0$$

$$\|\theta(\mathcal{D}) - \theta^*\| \rightarrow 0$$



# Previous results

Theorem (informal)[Ben Arous et. al., 2021]

For  $\sigma(z) = z^2$  SGD on  $\mathbb{S}^{d-1}$  with  $n \ll d \log d$ ,  $|\langle \theta^*, \theta_n \rangle| \rightarrow 0$ .  
SGD on  $\mathbb{S}^{d-1}$  with  $n \gg d \log^2 d$ ,  $|\langle \theta^*, \theta_n \rangle| \rightarrow 1$ .

“GAP”



Theorem (informal)[Mondelli & Montanari, 2018, Barbier et. al., 2019]

For  $\sigma(z) = z^2$   $\forall$  algorithm with  $n \ll d$ ,  $|\langle \theta^*, \theta(\mathcal{D}) \rangle| \rightarrow 0$ .  
 $\exists$  algorithm with  $n \sim d$ ,  $|\langle \theta^*, \theta(\mathcal{D}) \rangle| \rightarrow c > 0$ .

# Full batch GF - phase retrieval

Full batch Gradient flow

$$\frac{d\theta_t}{dt} = -\nabla_{\mathbb{S}^{d-1}} \mathcal{L}(\theta_{t-1}, \mathcal{D})$$

Theorem (informal)[ours]

For  $\sigma(z) = z^2$  full-batch GF with  $n \ll d \log d$ ,  $|\langle \theta^*, \theta \rangle| \rightarrow 0$ .



# Full batch GF - truncated phase retrieval

Full batch Gradient flow

$$\frac{d\theta_t}{dt} = - \nabla_{\mathbb{S}^{d-1}} \mathcal{L}(\theta_{t-1}, \mathcal{D})$$

Theorem (informal)[ours]

For  $\sigma(z) = \min\{z^2, C\}$  full-batch GF with  $n \sim d$ ,  $|\langle \theta^*, \theta \rangle| \rightarrow c > 0$ .



Theorem (informal)[Ben Arous et. al., 2021]

For  $\sigma(z) = \min\{z^2, C\}$

|                               |                                                        |
|-------------------------------|--------------------------------------------------------|
| SGD with $n \ll d \log d$ ,   | $ \langle \theta^*, \theta_n \rangle  \rightarrow 0$ . |
| SGD with $n \gg d \log^2 d$ , | $ \langle \theta^*, \theta_n \rangle  \rightarrow 1$ . |



# Full batch GD - truncated phase retrieval

Hypothesis class

$$\{\sigma(\langle x, \theta \rangle) \mid \theta \in \mathbb{R}^d\}$$

Empirical loss

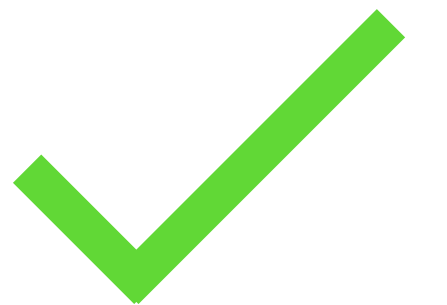
$$\mathcal{L}(\theta, \mathcal{D}) = \frac{1}{n} \sum_{i=1}^n (y_i - \sigma(\langle x, \theta \rangle))^2$$

Full batch Gradient Descent

$$\theta_t = \theta_{t-1} - \eta \nabla_{\mathbb{R}^d} \mathcal{L}(\theta_{t-1}, \mathcal{D})$$

Theorem (informal)[ours]

For  $\sigma(z) = \min\{z^2, C\}$  full-batch GD with  $n \sim d$ ,  $\|\hat{\theta} - \theta^*\| \rightarrow 0$ .  
Moreover, strong recovery is achieved after  $\sim \log d$  steps.

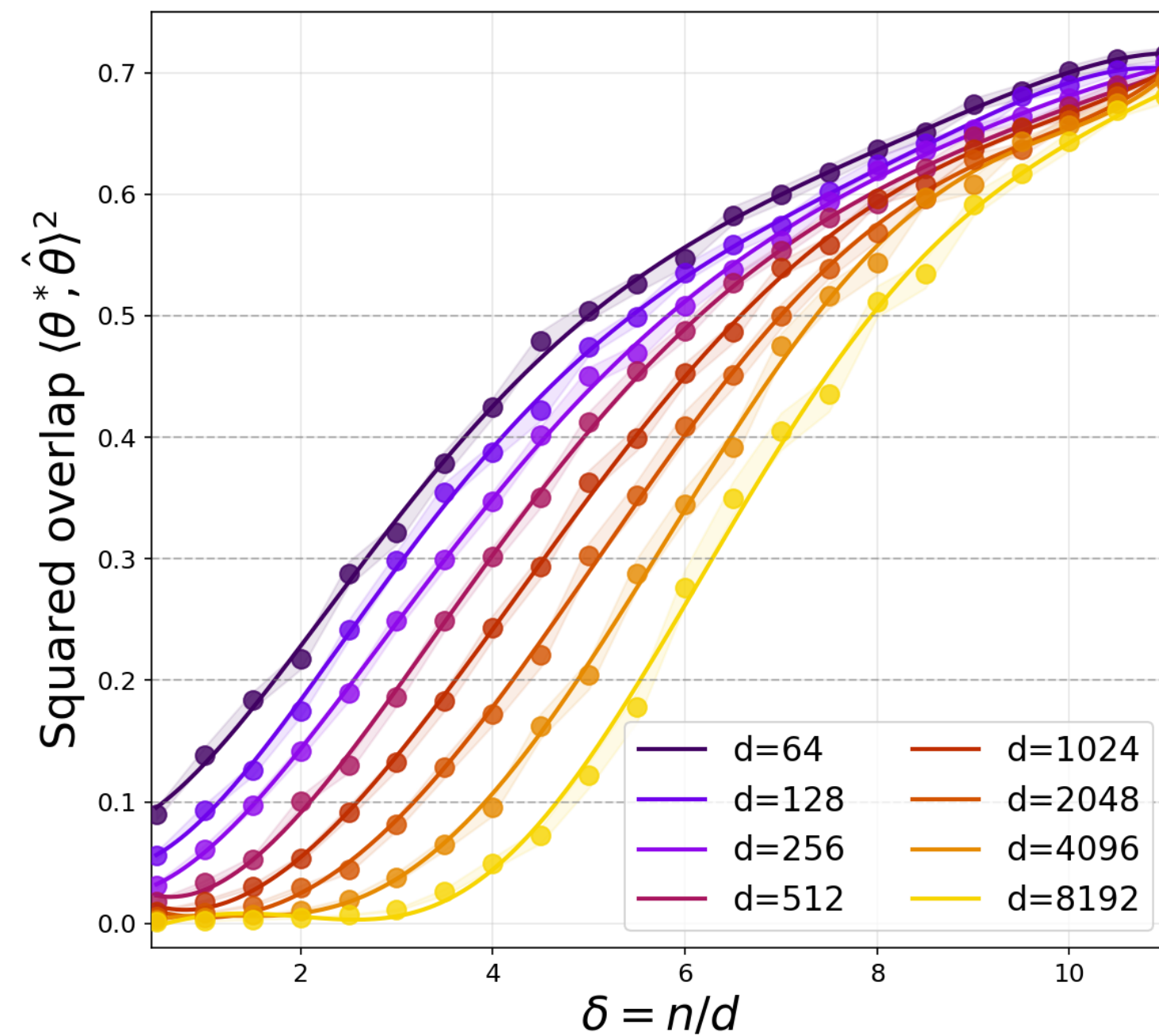


**Thank you for your time!**

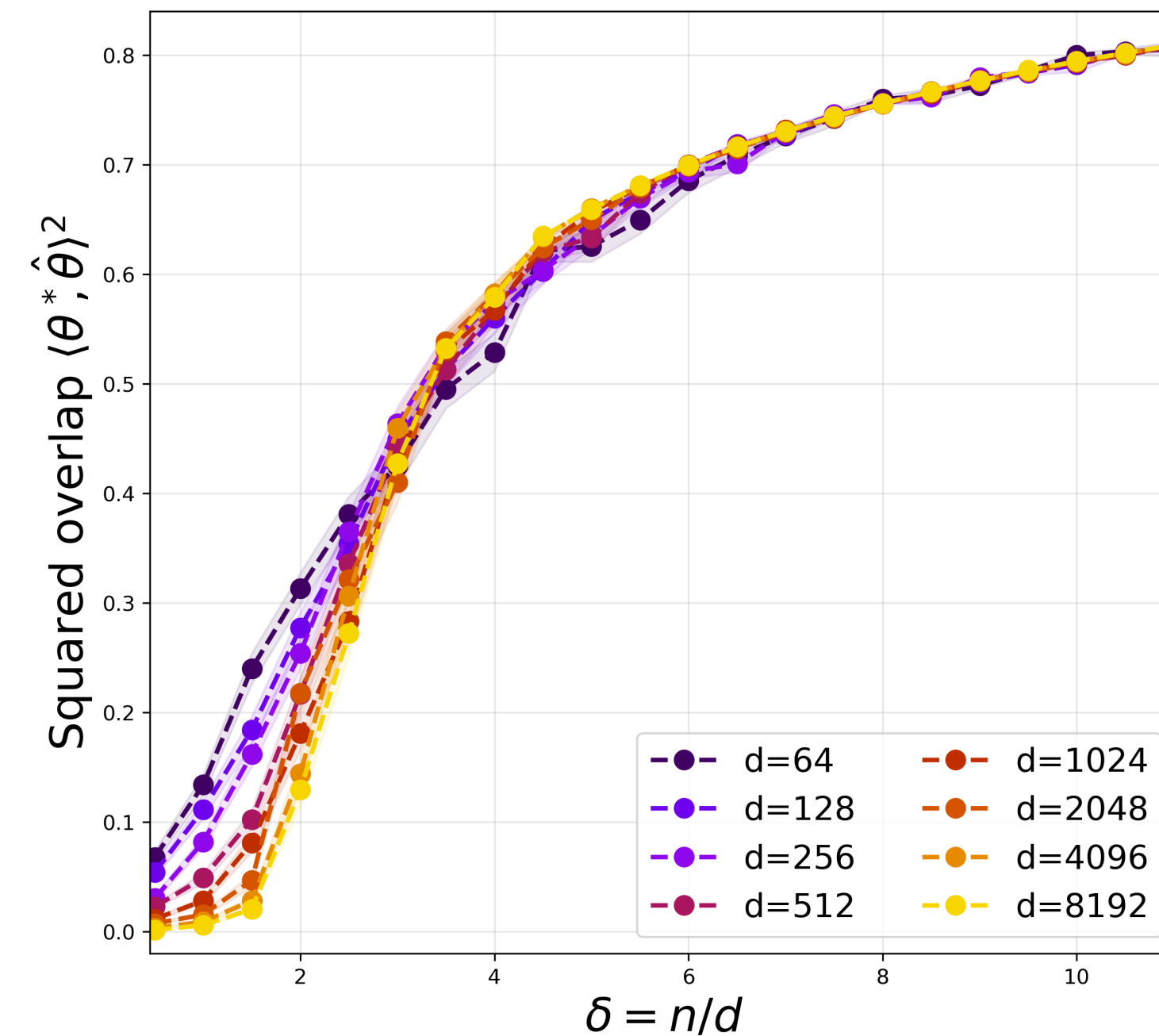


# Experiments - Full batch GD

$$\sigma(z) = z^2$$



$$\sigma(z) = \min\{z^2, C\}$$



Spherical GD with learning rate  $\eta = 0.1$  for  $1000 \log^2 d$  steps.