

Linearly Converging Error Compensated SGD

Eduard Gorbunov
MIPT, Yandex and Sirius

Dmitry Kovalev
KAUST

Dmitry Makarenko
MIPT

Peter Richtárik
KAUST

The talk is based on our paper «Linearly Converging Error Compensated SGD» accepted to
NeurIPS 2020

Yandex



Университет
Сириус



Dmitry Kovalev
PhD student
KAUST



Dmitry Makarenko
PhD student
MIPT

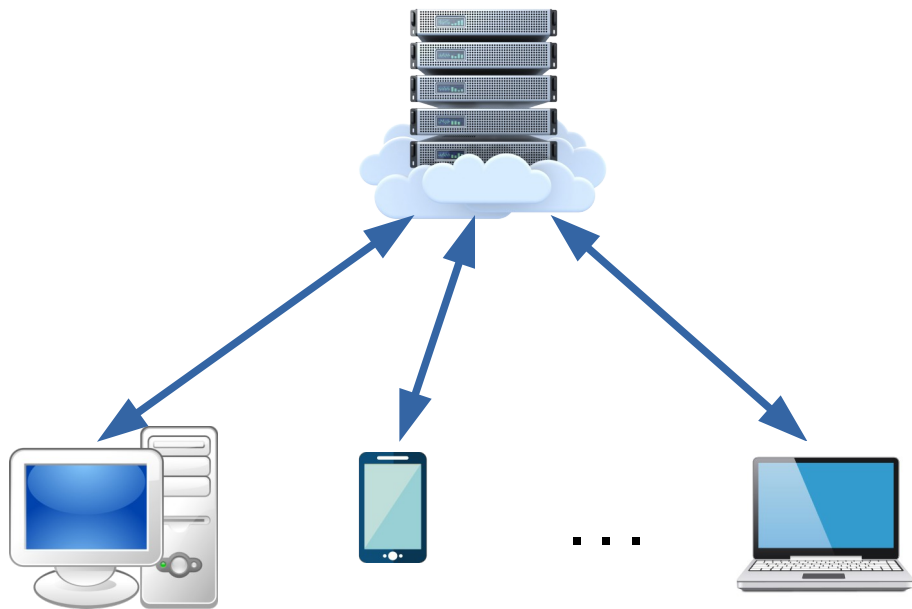


Peter Richtárik
Professor of Computer Science
KAUST

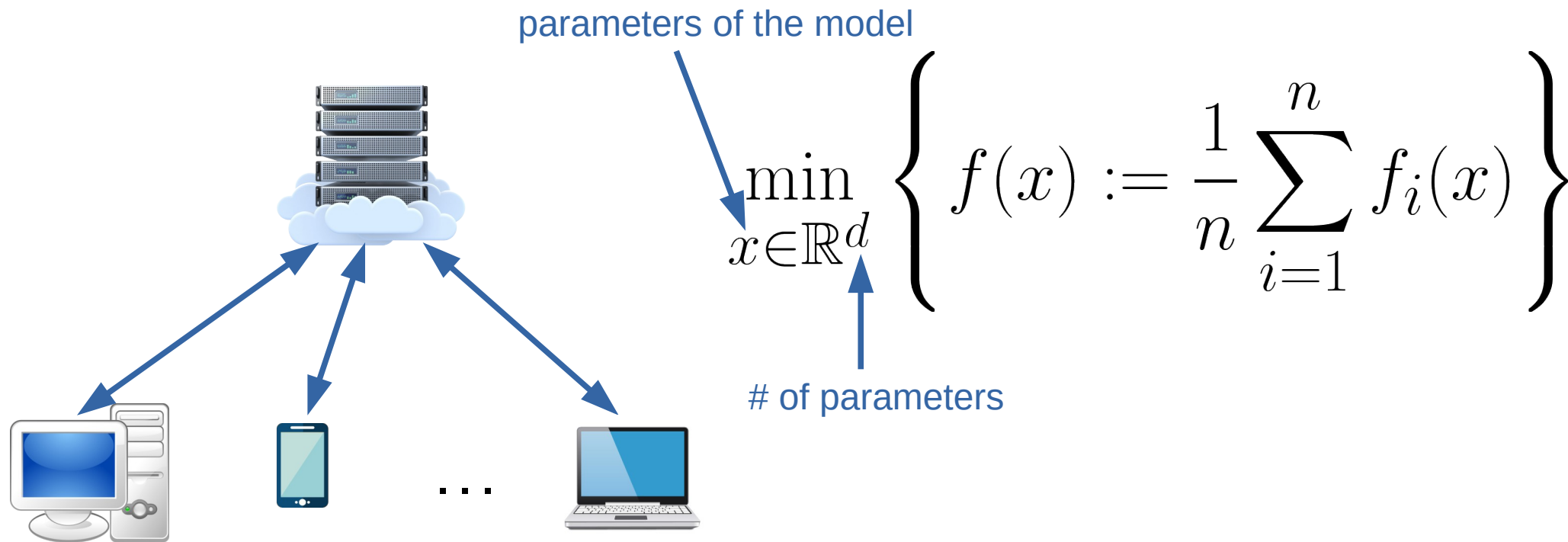
Outline

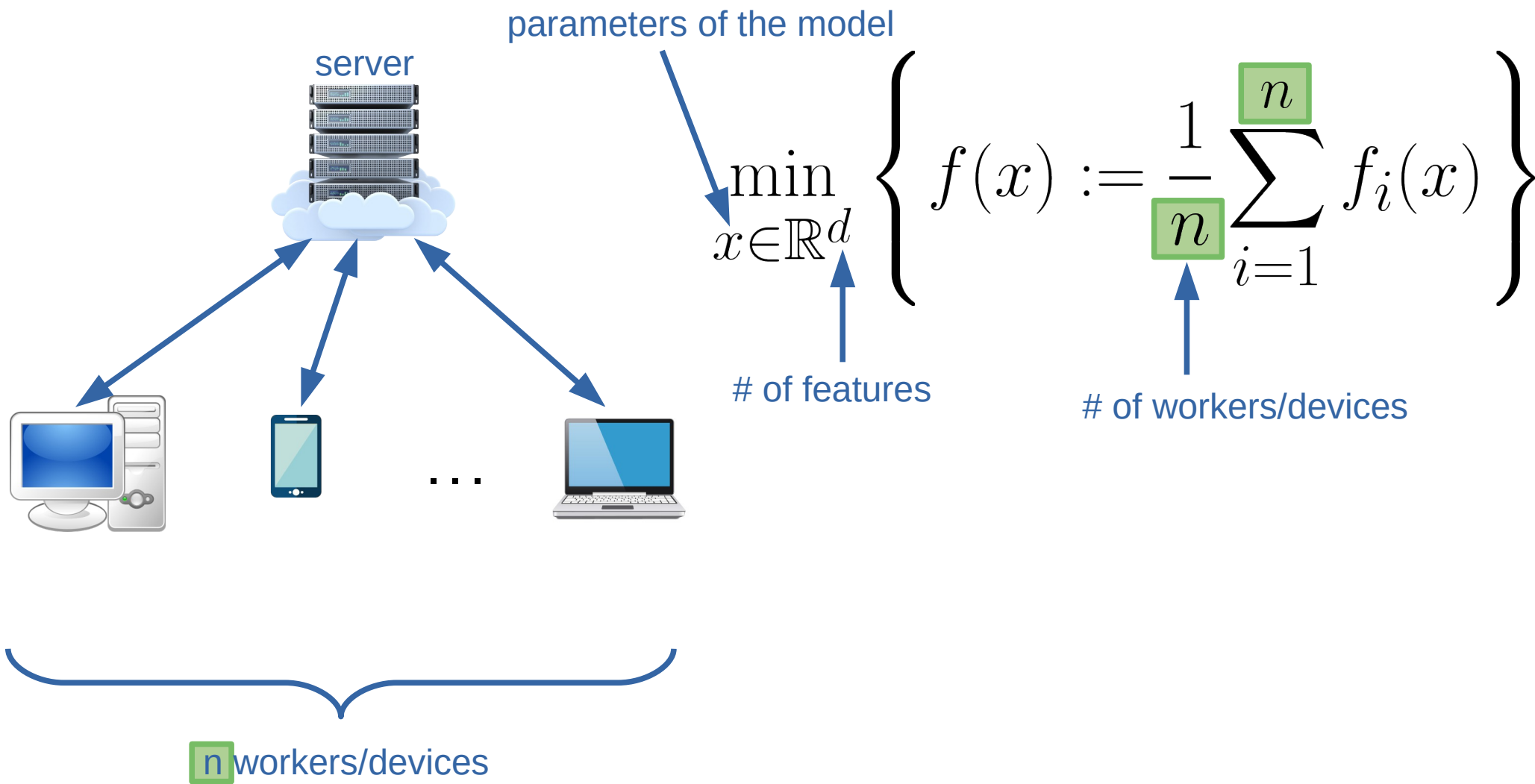
-
- 1 The problem
 - 2 Parallel SGD
 - 3 Communication bottleneck
 - 4 Error-compensated SGD
 - 5 New methods: EC-GDstar, EC-SGD-DIANA and EC-L-SVRG-DIANA
 - 6 Unified convergence analysis
 - 7 Experiments
 - 8 Other applications of the framework: methods without error-compensation and methods with delayed updates
- 15 minutes
- 10 minutes
- 10 minutes
- 5 minutes
- 5 minutes

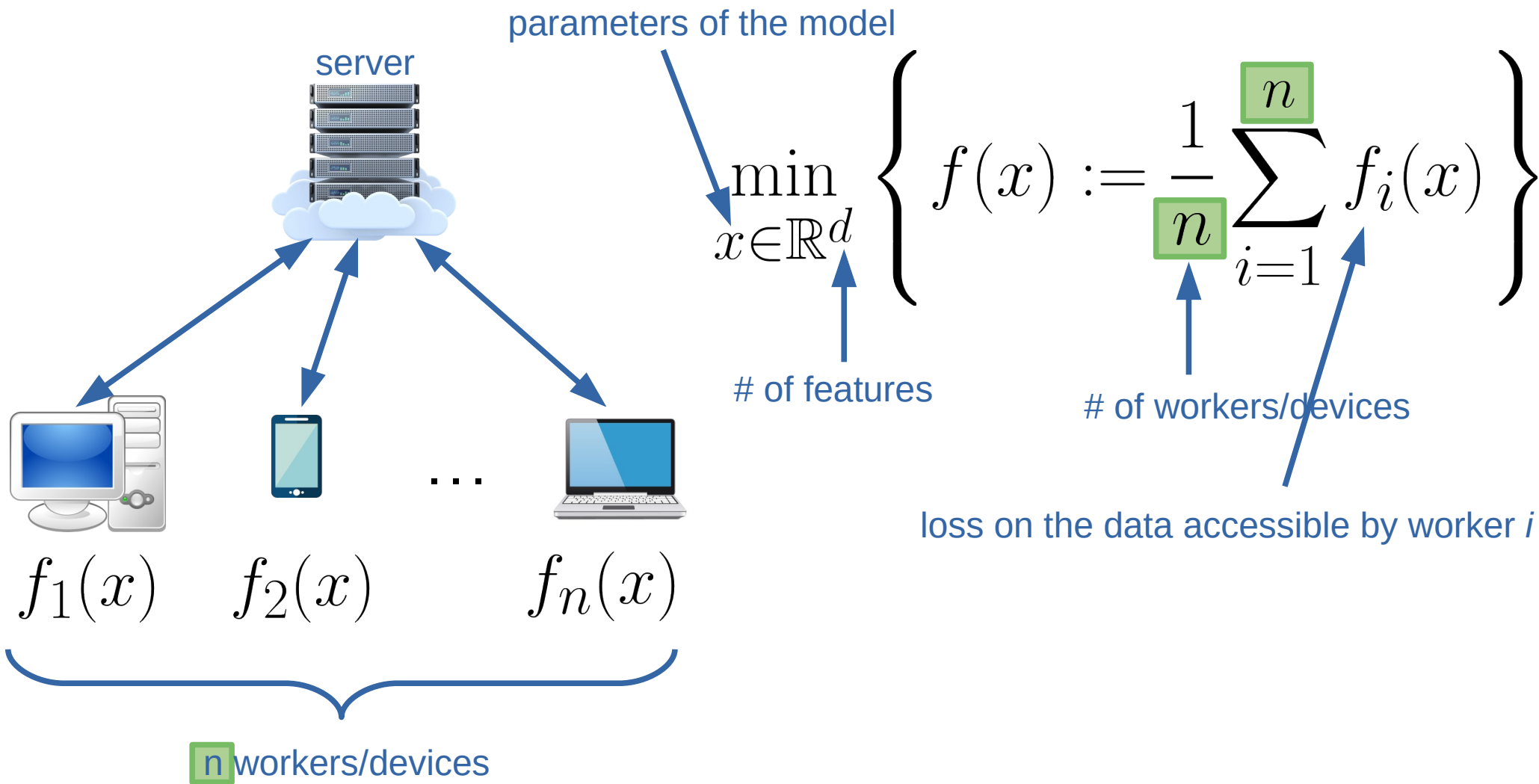
1. The Problem

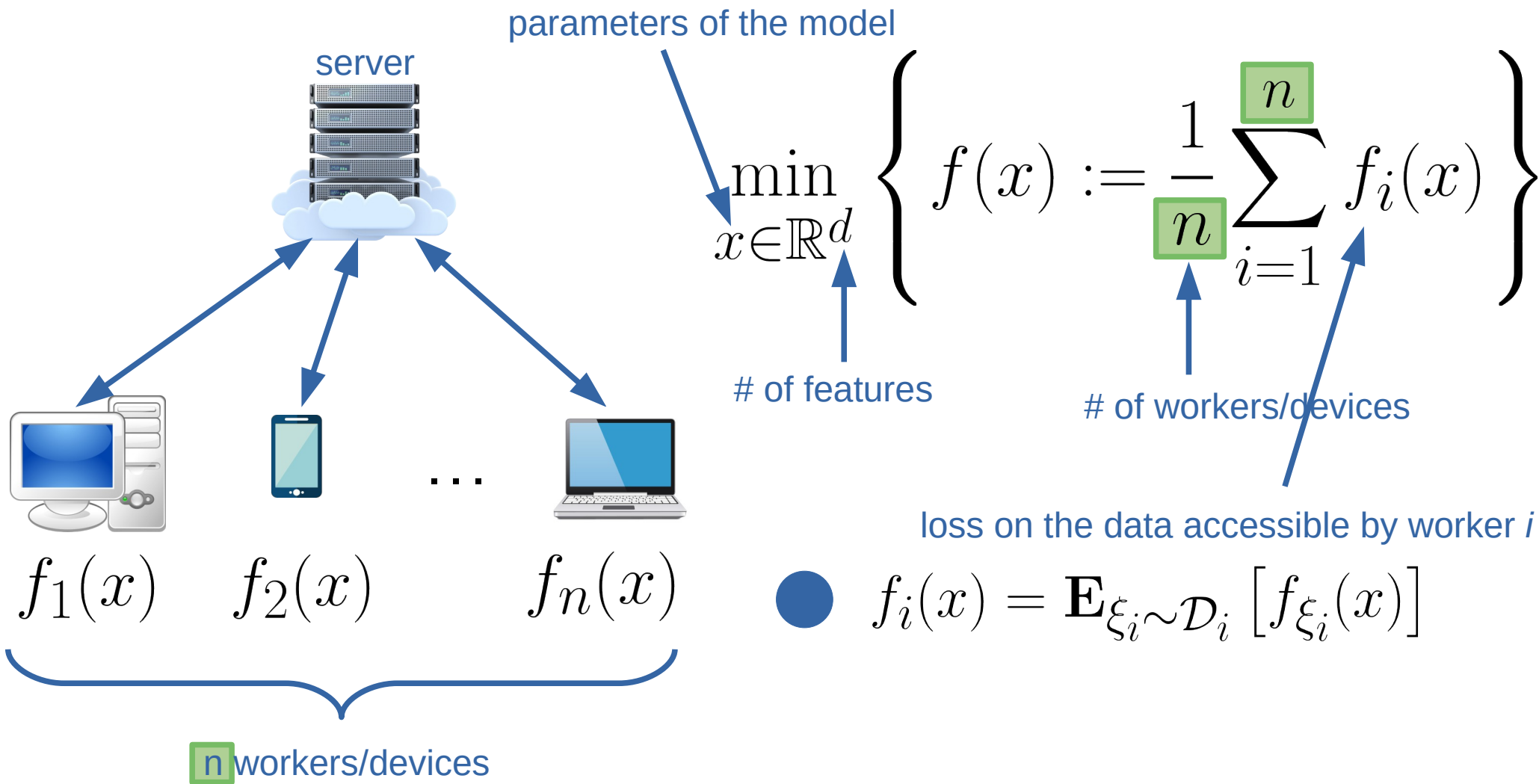


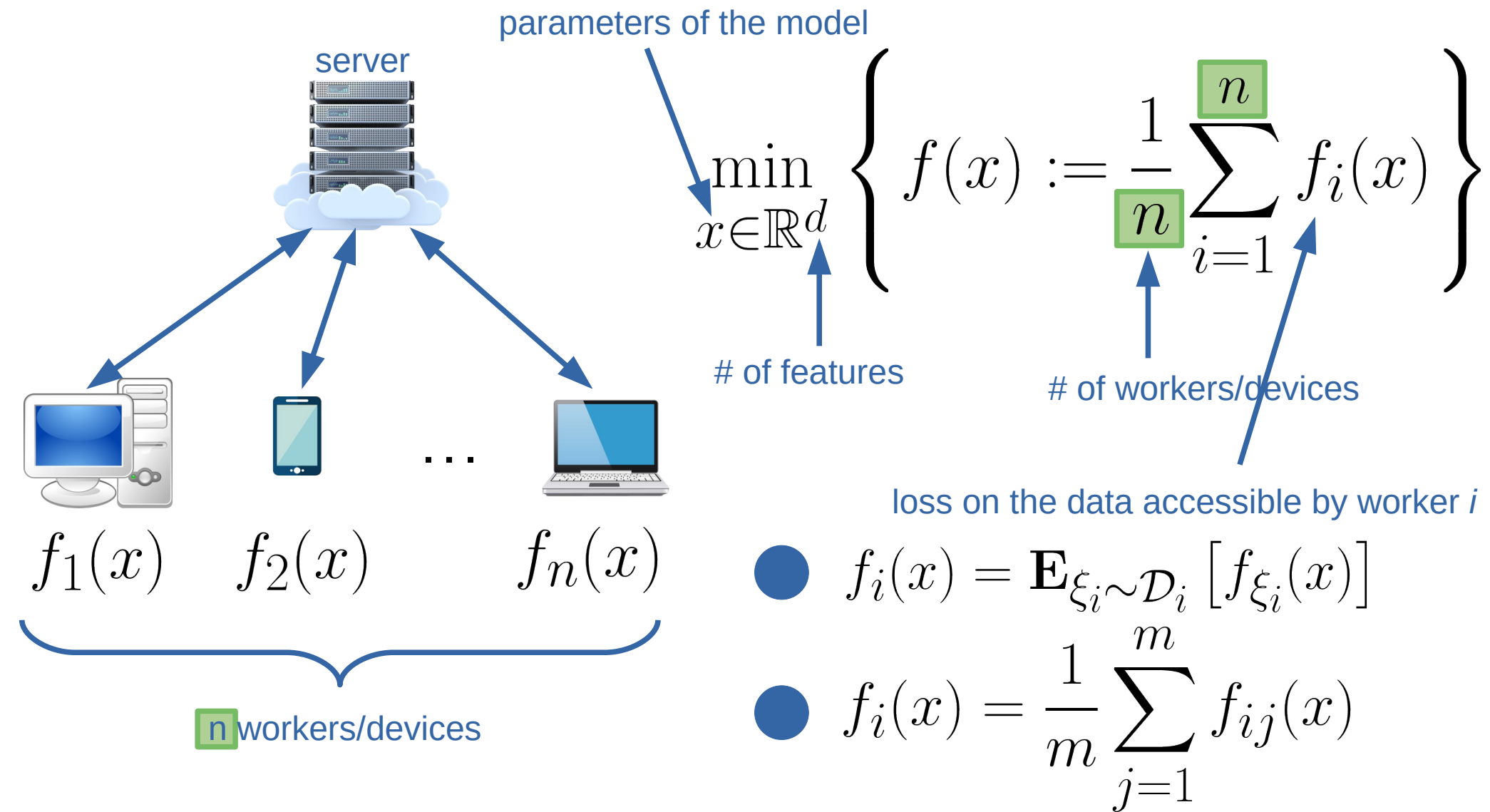
$$\min_{x \in \mathbb{R}^d} \left\{ f(x) := \frac{1}{n} \sum_{i=1}^n f_i(x) \right\}$$











Assumptions

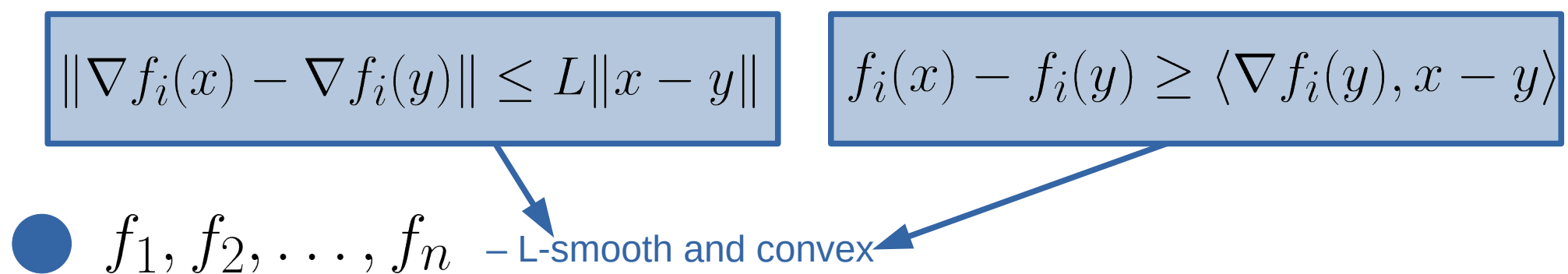
- $f_1, f_2, \dots, f_n$ – L-smooth and convex

Assumptions

$$\|\nabla f_i(x) - \nabla f_i(y)\| \leq L\|x - y\|$$

$$f_i(x) - f_i(y) \geq \langle \nabla f_i(y), x - y \rangle$$

● $f_1, f_2, \dots, f_n$ – L-smooth and convex



Assumptions

$$\|\nabla f_i(x) - \nabla f_i(y)\| \leq L\|x - y\|$$

$$f_i(x) - f_i(y) \geq \langle \nabla f_i(y), x - y \rangle$$

● $f_1, f_2, \dots, f_n$ – L-smooth and convex

● f – strongly quasi-convex

$$f(x^*) \geq f(x) + \langle \nabla f(x), x^* - x \rangle + \frac{\mu}{2} \|x^* - x\|^2$$

Assumptions

$$\|\nabla f_i(x) - \nabla f_i(y)\| \leq L\|x - y\|$$

$$f_i(x) - f_i(y) \geq \langle \nabla f_i(y), x - y \rangle$$

● $f_1, f_2, \dots, f_n$ – L-smooth and convex

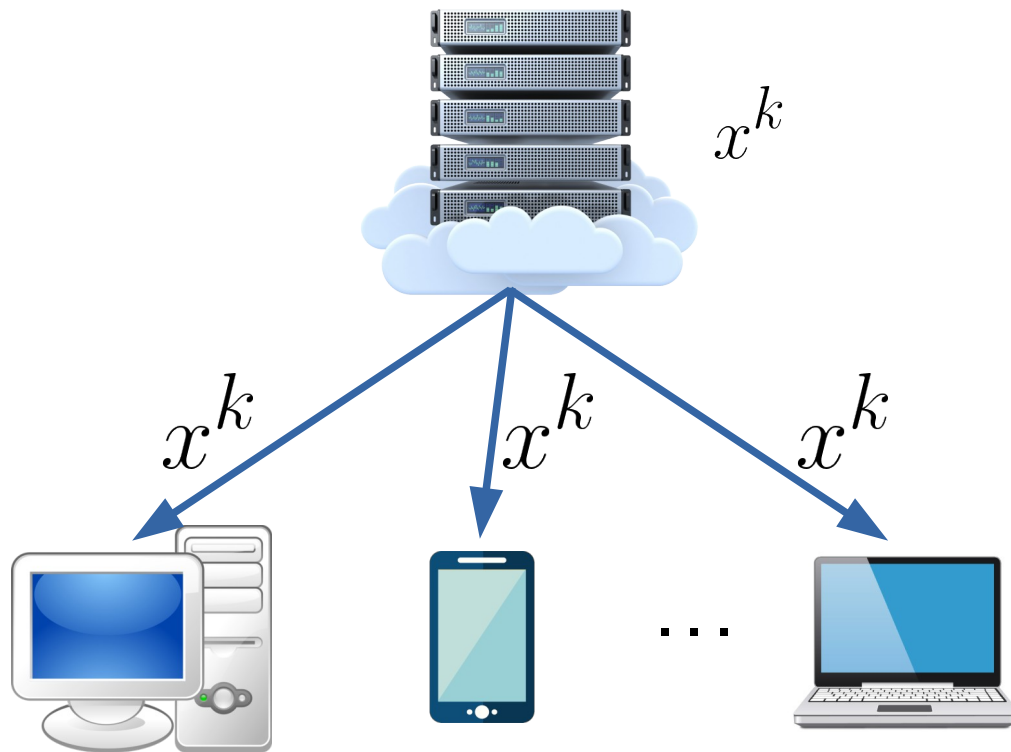
● f – strongly quasi-convex

$$f(x^*) \geq f(x) + \langle \nabla f(x), x^* - x \rangle + \frac{\mu}{2} \|x^* - x\|^2$$

the solution of the problem

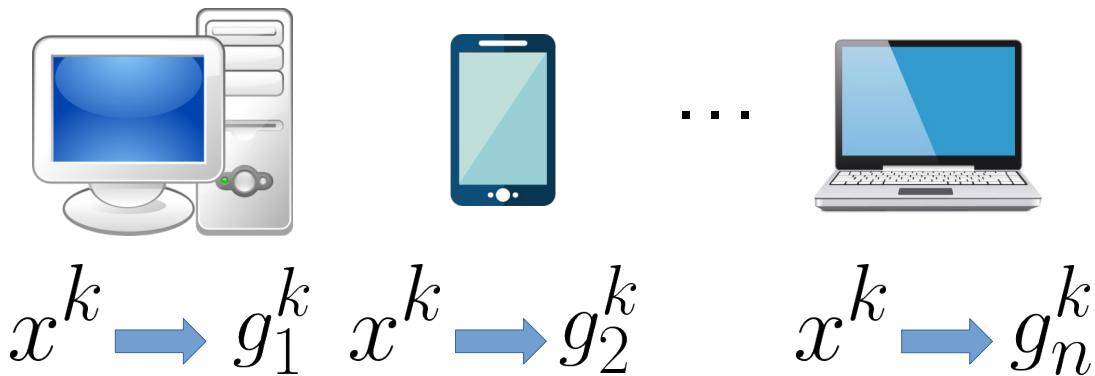
2. Parallel SGD

1 Server broadcasts the parameters

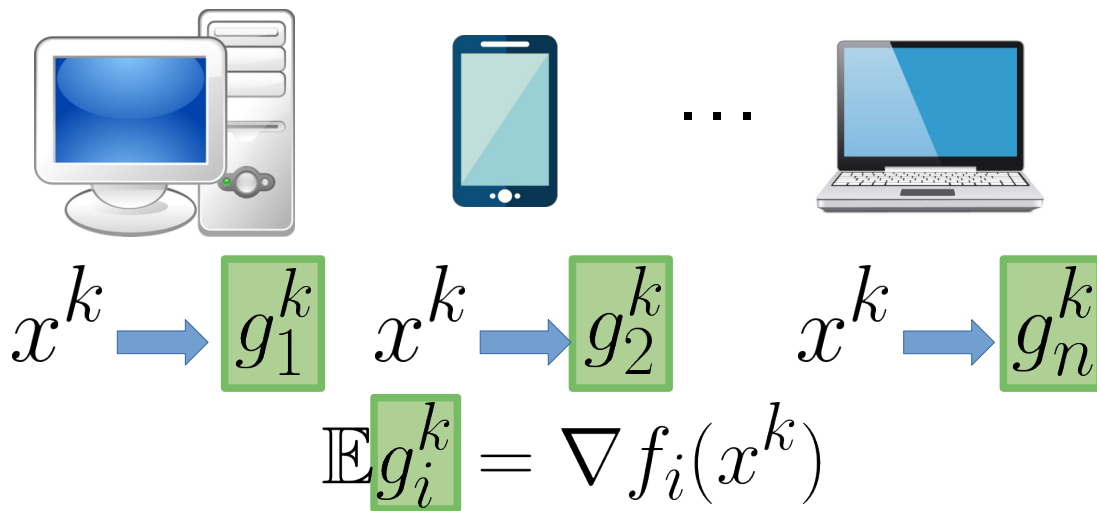


1 Server broadcasts the parameters

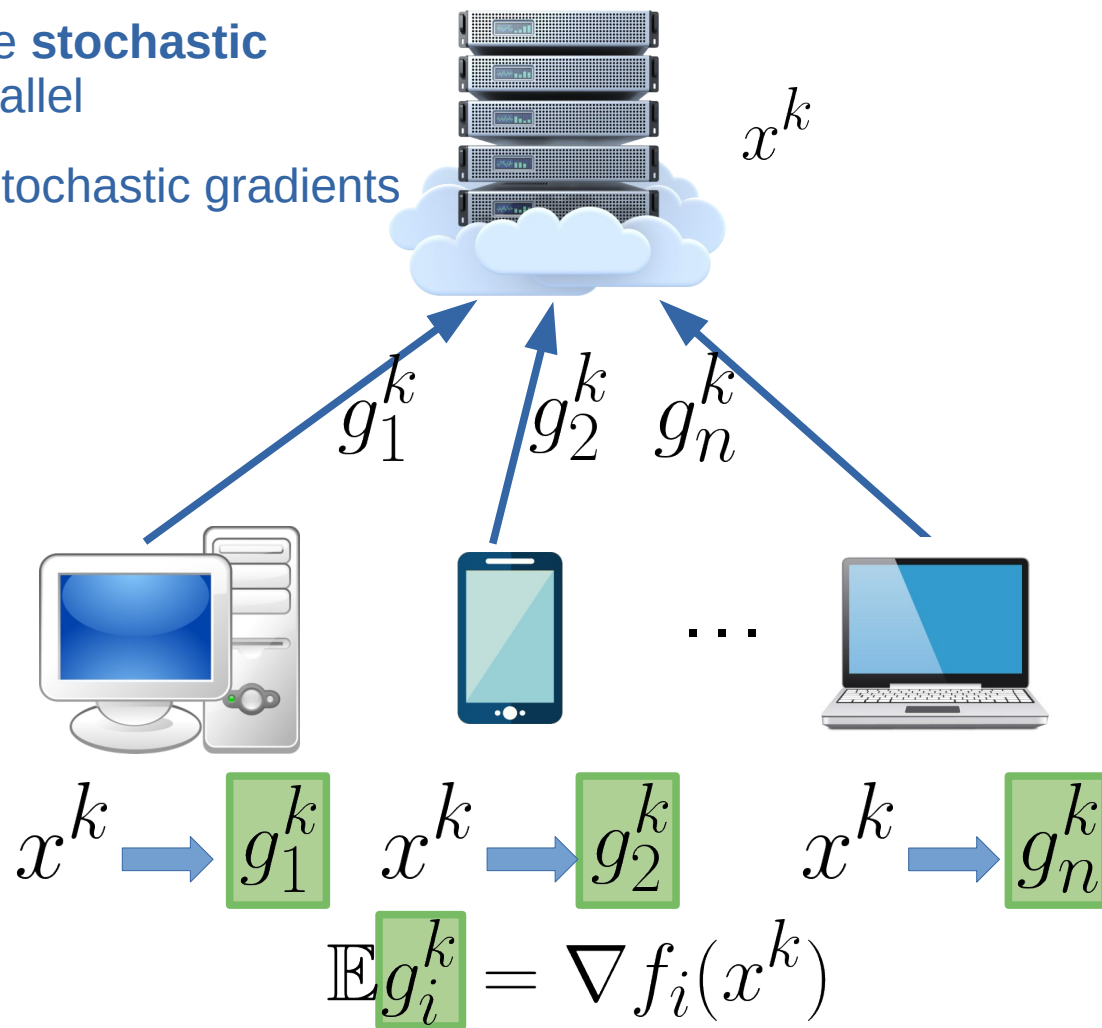
2 Devices compute **stochastic gradients** in parallel



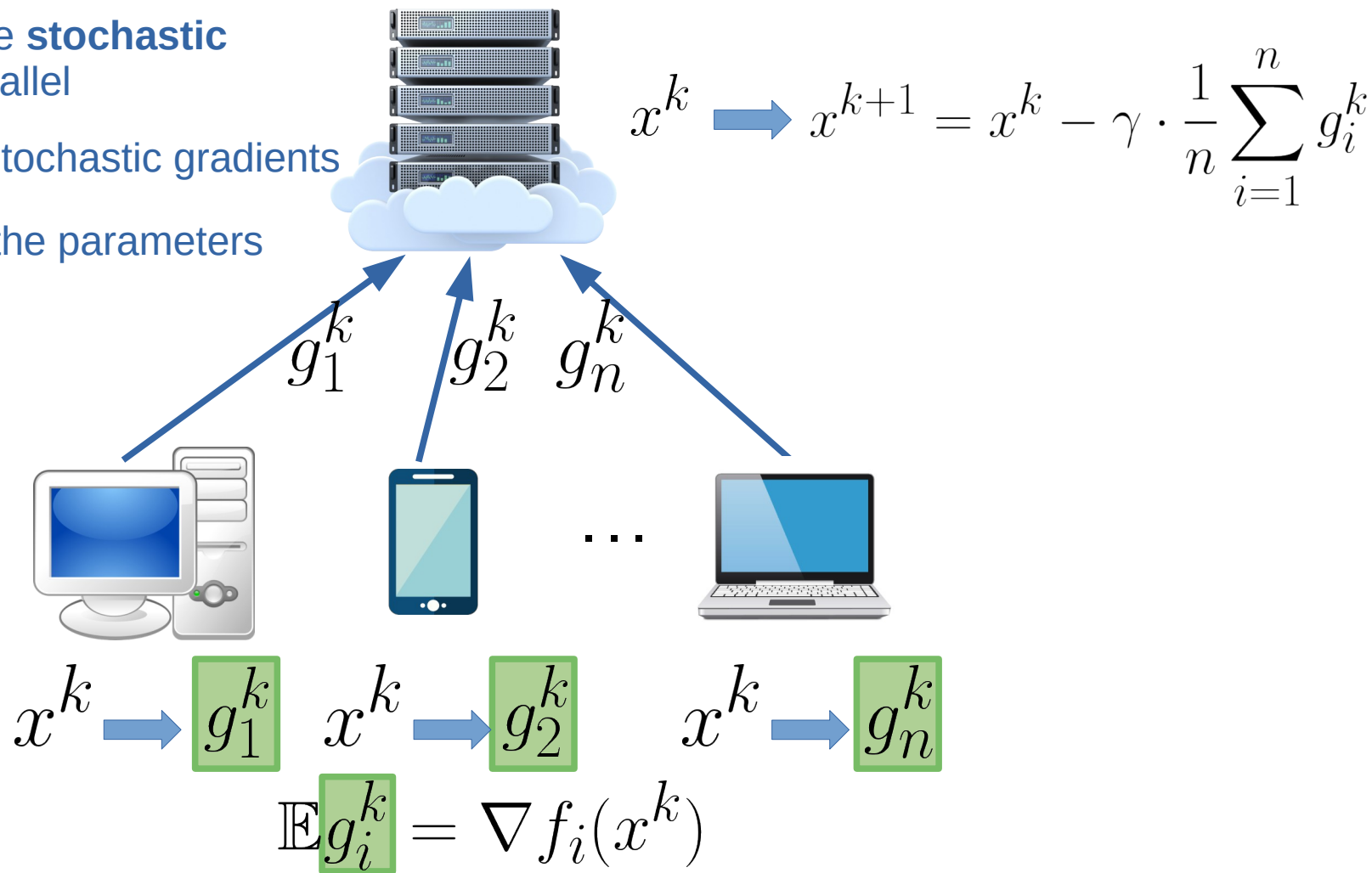
- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel



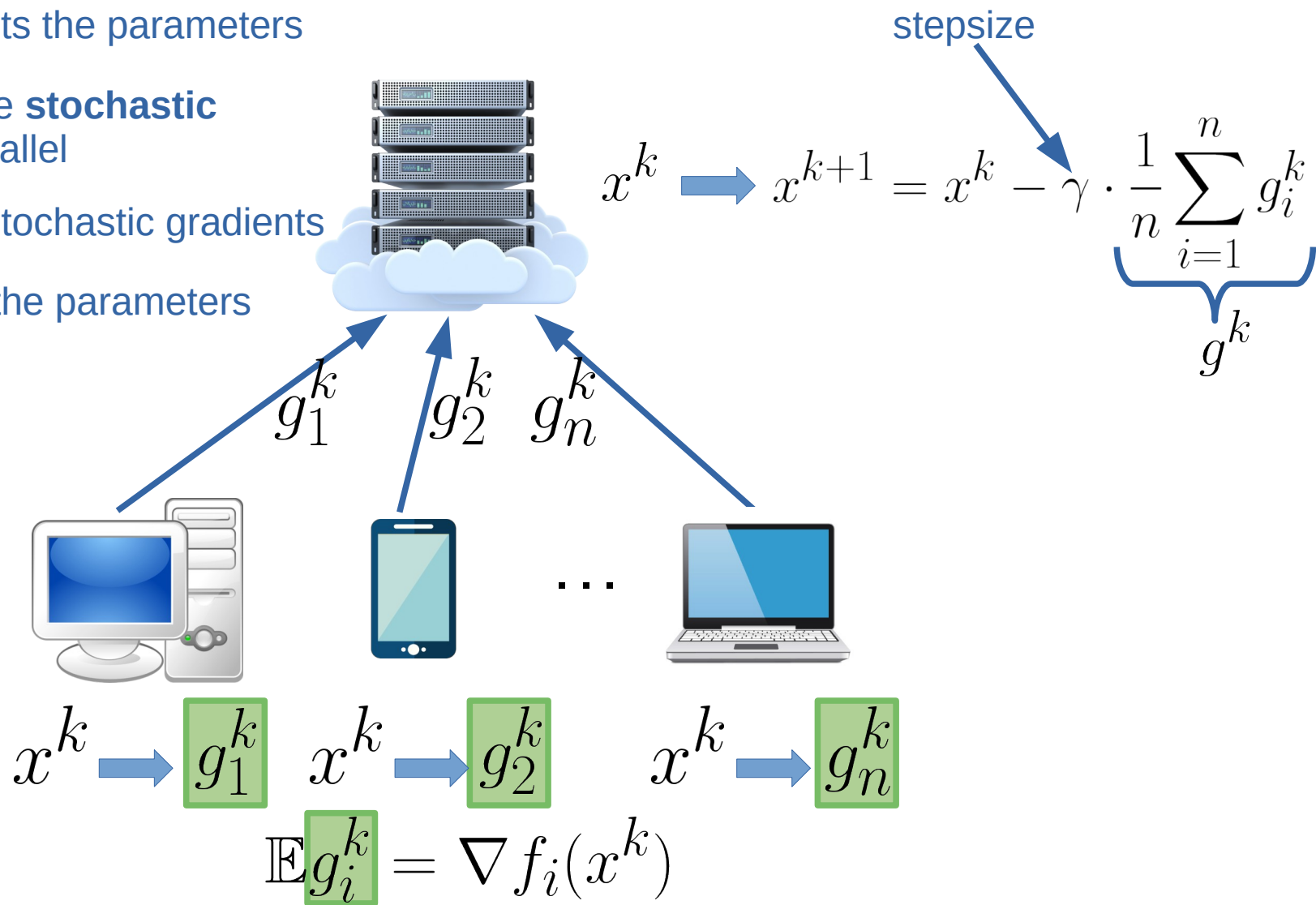
- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Server gathers stochastic gradients



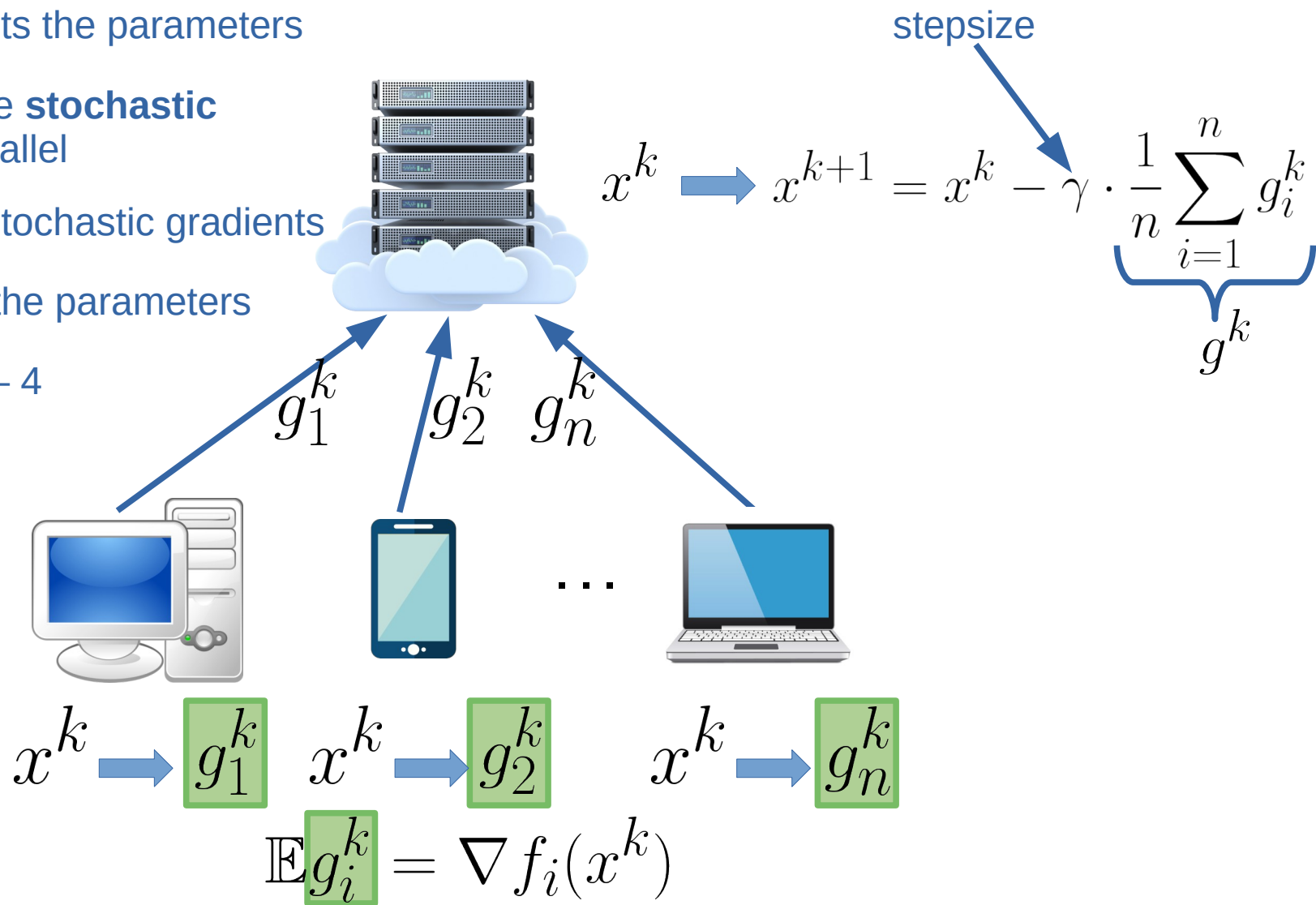
- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Server gathers stochastic gradients
- 4 Server updates the parameters



- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Server gathers stochastic gradients
- 4 Server updates the parameters



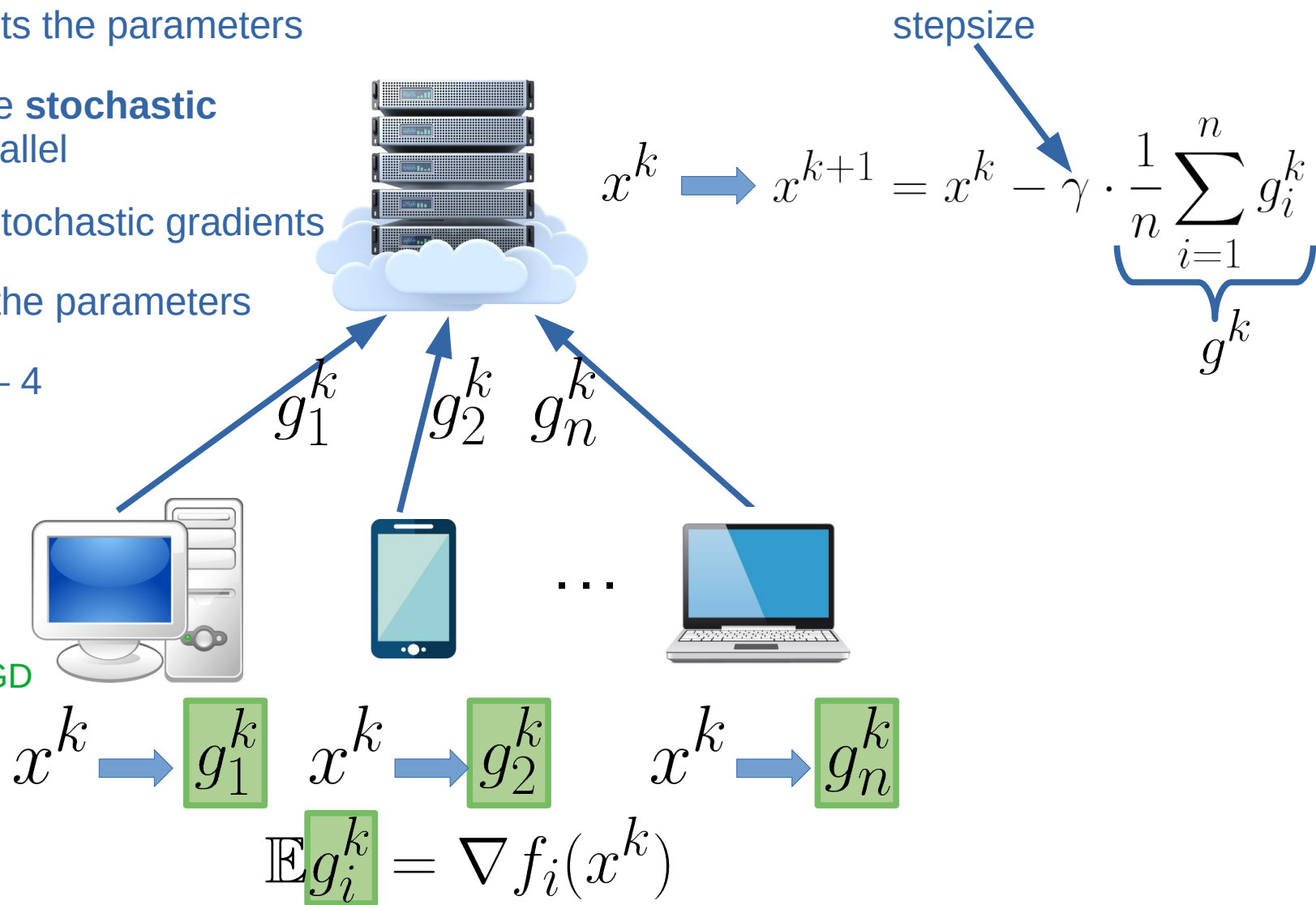
- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Server gathers stochastic gradients
- 4 Server updates the parameters
- 5 Repeat steps 1 – 4



- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Server gathers stochastic gradients
- 4 Server updates the parameters
- 5 Repeat steps 1 – 4

Good news:

- ✓ Very simple algorithm
- ✓ Can be much faster than non-parallel SGD



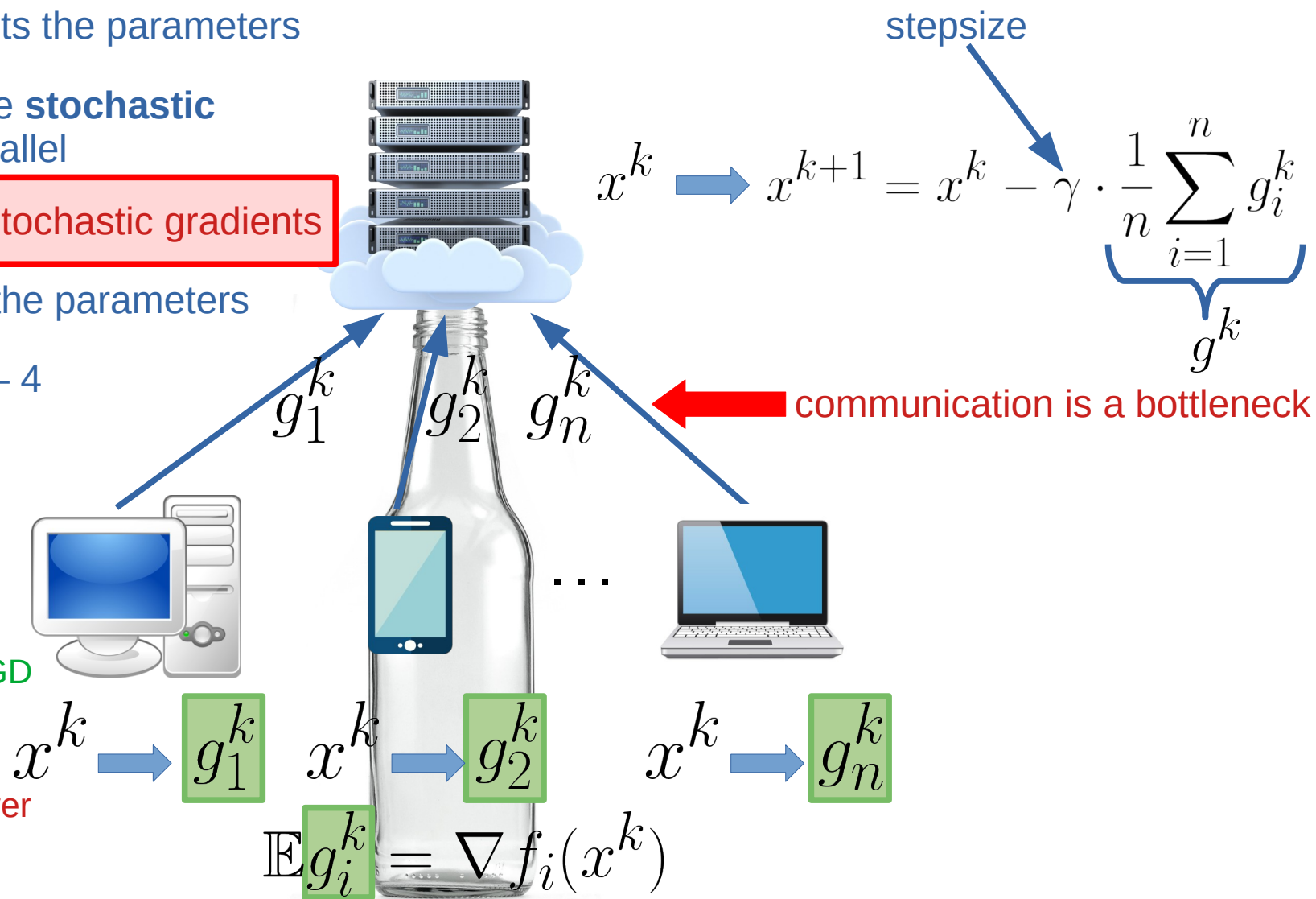
- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 **Server gathers stochastic gradients**
- 4 Server updates the parameters
- 5 Repeat steps 1 – 4

Good news:

- ✓ Very simple algorithm
- ✓ Can be much faster than non-parallel SGD

Issues:

- ✗ Overload of the server



3. Communication Bottleneck

How to Handle Communication Bottleneck?



How to Handle Communication Bottleneck?

● Change the topology of the network → Decentralized optimization

● Do more work on each worker and communicate less → Local-SGD/Federated Averaging

How to Handle Communication Bottleneck?

- Change the topology of the network  Decentralized optimization
- Do more work on each worker and communicate less  Local-SGD/Federated Averaging
- Send less information to reduce the communication cost

How to Handle Communication Bottleneck?

- Change the topology of the network  Decentralized optimization
- Do more work on each worker and communicate less  Local-SGD/Federated Averaging
- Send less information to reduce the communication cost

Example:

$$\text{Send } g_i^k = \begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \quad \img alt="blue arrow" data-bbox="495 780 585 845"/> \quad \text{Send } \mathcal{C}(g_i^k) = \begin{pmatrix} 0 \\ -15 \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

How to Handle Communication Bottleneck?

- Change the topology of the network → Decentralized optimization
 - Do more work on each worker and communicate less → Local-SGD/Federated Averaging
 - Send less information to reduce the communication cost
- We focus on this approach

Example:

Send $g_i^k = \begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix}$

Compression operator

Send $\mathcal{C}(g_i^k) = \begin{pmatrix} 0 \\ -15 \\ 0 \\ 0 \\ 0 \end{pmatrix}$

How to Handle Communication Bottleneck?

- Change the topology of the network → Decentralized optimization
 - Do more work on each worker and communicate less → Local-SGD/Federated Averaging
 - Send less information to reduce the communication cost
- We focus on this approach

Example:

Send $g_i^k = \begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix}$

Compression operator

Send $\mathcal{C}(g_i^k) = \begin{pmatrix} 0 \\ -15 \\ 0 \\ 0 \\ 0 \end{pmatrix}$

What are the options for choosing this?

Compression Operators



Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

Compression Operators



```
graph TD; A[Compression Operators] --> B[Unbiased compressors<br/>(quantizations)]; A --> C[Biased compressors];
```

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E} \|\mathcal{Q}(x) - x\|^2 \leq \omega \|x\|^2$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

Compression Operators



```
graph TD; A[Compression Operators] --> B[Unbiased compressors<br/>(quantizations)]; A --> C[Biased compressors];
```

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E}\|\mathcal{Q}(x) - x\|^2 \leq \omega \|x\|^2$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

Compression Operators

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E} \|\mathcal{Q}(x) - x\|^2 \leq \omega \|x\|^2$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E} \|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

Example: RandK (for $K = 2$)

$$\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \rightarrow$$

Pick $K = 2$ components uniformly at random

Compression Operators

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E}\|\mathcal{Q}(x) - x\|^2 \leq \omega \|x\|^2$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta)\|x\|^2$$

Example: RandK (for $K = 2$)

$$\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \longrightarrow \frac{5}{2} \cdot \begin{pmatrix} 1 \\ 0 \\ 0 \\ -7 \\ 0 \end{pmatrix}$$

Pick $K = 2$ components uniformly at random

Compression Operators

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

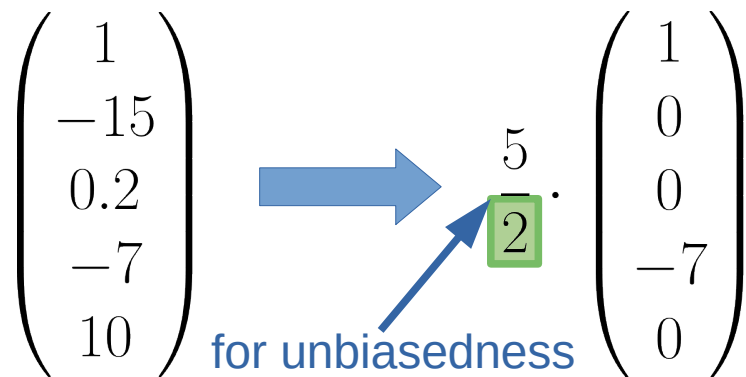
$$\mathbb{E}\|\mathcal{Q}(x) - x\|^2 \leq \omega\|x\|^2$$

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta)\|x\|^2$$

Example: RandK (for $K = 2$)



$$\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \rightarrow \begin{pmatrix} 1 \\ 0 \\ 0 \\ -7 \\ 0 \end{pmatrix}$$

for unbiasedness

Pick $K = 2$ components uniformly at random

Compression Operators

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E}\|\mathcal{Q}(x) - x\|^2 \leq \omega\|x\|^2$$

Example: RandK (for $K = 2$)

$$\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \xrightarrow{\text{RandK}} \begin{pmatrix} 1 \\ 0 \\ 0 \\ -7 \\ 0 \end{pmatrix}$$

for unbiasedness

Pick $K = 2$ components uniformly at random

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta)\|x\|^2$$

Example: TopK (for $K = 2$)

$$\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix} \xrightarrow{\text{TopK}} \begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix}$$

Pick $K = 2$ components with largest absolute value

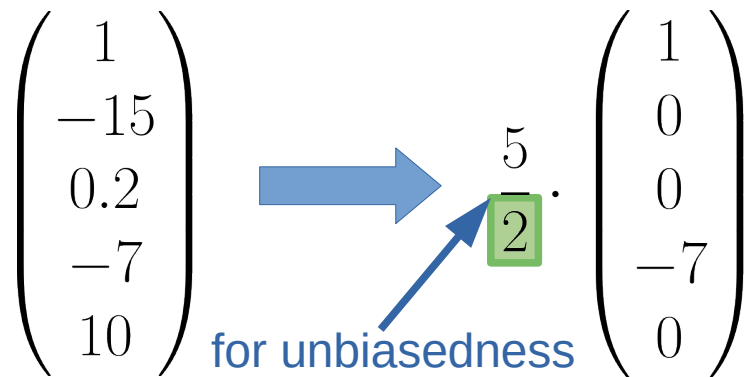
Compression Operators

Unbiased compressors
(quantizations)

$$x \rightarrow \mathcal{Q}(x) \quad \mathbb{E}[\mathcal{Q}(x)] = x$$

$$\mathbb{E}\|\mathcal{Q}(x) - x\|^2 \leq \omega \|x\|^2$$

Example: RandK (for $K = 2$)



The diagram illustrates the RandK quantization process for $K=2$. It starts with a vector $\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix}$. A thick blue arrow points to a second vector $\begin{pmatrix} 1 \\ 0 \\ 0 \\ -7 \\ 0 \end{pmatrix}$. A thin blue arrow points from the text "for unbiasedness" to a green box containing the number 2, which is positioned next to the value 5 in the second vector, indicating that 2 components were randomly selected.

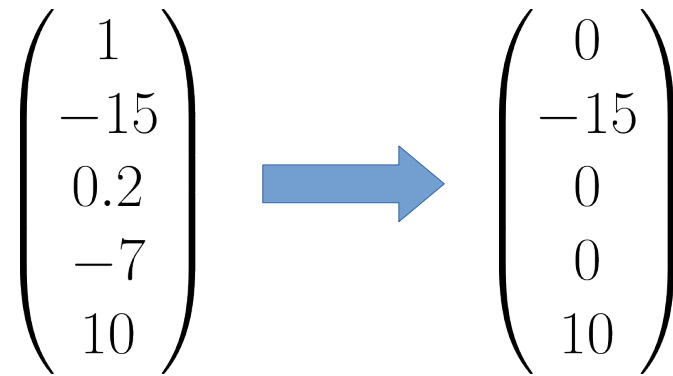
Pick $K = 2$ components uniformly at random

Biased compressors

$$x \rightarrow \mathcal{C}(x)$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

Example: TopK (for $K = 2$)



The diagram illustrates the TopK quantization process for $K=2$. It starts with a vector $\begin{pmatrix} 1 \\ -15 \\ 0.2 \\ -7 \\ 10 \end{pmatrix}$. A thick blue arrow points to a second vector $\begin{pmatrix} 0 \\ -15 \\ 0 \\ 0 \\ 10 \end{pmatrix}$, where the components with the largest absolute values (10 and -15) are preserved, and the others are set to zero.

Pick $K = 2$ components with largest absolute value

Methods with Unbiased Compressors

QSGD



Alistarh, Dan, Demjan Grubic, Jerry Li, Ryota Tomioka, and Milan Vojnovic. **"QSGD: Communication-efficient SGD via gradient quantization and encoding."** In *Advances in Neural Information Processing Systems*, pp. 1709-1720. 2017.

TernGrad



Wen, Wei, Cong Xu, Feng Yan, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. **"Terngrad: Ternary gradients to reduce communication in distributed deep learning."** In *Advances in neural information processing systems*, pp. 1509-1519. 2017.

DQGD



Khirirat, Sarit, Hamid Reza Feyzmahdavian, and Mikael Johansson. **"Distributed learning with compressed gradients."** arXiv preprint arXiv:1806.06573 (2018).

DIANA



Mishchenko, Konstantin, Eduard Gorbunov, Martin Takáč, and Peter Richtárik. **"Distributed learning with compressed gradient differences."** arXiv preprint arXiv:1901.09269 (2019).



Horváth, Samuel, Dmitry Kovalev, Konstantin Mishchenko, Sebastian Stich, and Peter Richtárik. **"Stochastic distributed learning with gradient quantization and variance reduction."** arXiv preprint arXiv:1904.05115 (2019).

Methods with Unbiased Compressors

QSGD



Alistarh, Dan, Demjan Grubic, Jerry Li, Ryota Tomioka, and Milan Vojnovic. **"QSGD: Communication-efficient SGD via gradient quantization and encoding."** In *Advances in Neural Information Processing Systems*, pp. 1709-1720. 2017.

TernGrad



Wen, Wei, Cong Xu, Feng Yan, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. **"Terngrad: Ternary gradients to reduce communication in distributed deep learning."** In *Advances in neural information processing systems*, pp. 1509-1519. 2017.

DQGD



Khirirat, Sarit, Hamid Reza Feyzmahdavian, and Mikael Johansson. **"Distributed learning with compressed gradients."** arXiv preprint arXiv:1806.06573 (2018).

DIANA



Mishchenko, Konstantin, Eduard Gorbunov, Martin Takáč, and Peter Richtárik. **"Distributed learning with compressed gradient differences."** arXiv preprint arXiv:1901.09269 (2019).



Horváth, Samuel, Dmitry Kovalev, Konstantin Mishchenko, Sebastian Stich, and Peter Richtárik. **"Stochastic distributed learning with gradient quantization and variance reduction."** arXiv preprint arXiv:1904.05115 (2019).

Sublinear convergence rates even in the case when workers quantize full gradients

Converges **linearly** when workers quantize full gradients

Parallel SGD with Biased Compressor Can Diverge at Exponential Rate



Beznosikov, Aleksandr, Samuel Horváth, Peter Richtárik, and Mher Safaryan. "On Biased Compression for Distributed Learning." arXiv preprint arXiv:2002.12410 (2020).

$$n = d = 3$$

$$f_1(x) = \langle a, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_2(x) = \langle b, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_3(x) = \langle c, x \rangle^2 + \frac{1}{4} \|x\|^2$$

$$a = (-3, 2, 2)^\top$$

$$b = (2, -3, 2)^\top$$

$$c = (2, 2, -3)^\top$$

$$x^0 = (t, t, t)^\top$$

Parallel SGD with Biased Compressor Can Diverge at Exponential Rate



Beznosikov, Aleksandr, Samuel Horváth, Peter Richtárik, and Mher Safaryan. "On Biased Compression for Distributed Learning." arXiv preprint arXiv:2002.12410 (2020).

$$n = d = 3$$

$$f_1(x) = \langle a, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_2(x) = \langle b, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_3(x) = \langle c, x \rangle^2 + \frac{1}{4} \|x\|^2$$

$$a = (-3, 2, 2)^\top \quad b = (2, -3, 2)^\top \quad c = (2, 2, -3)^\top$$

$$x^0 = (t, t, t)^\top$$

In this case Parallel SGD with Top1 compression operator satisfies

$$x^k = \left(1 + \frac{11\gamma}{6}\right)^k x^0$$

Parallel SGD with Biased Compressor Can Diverge at Exponential Rate



Beznosikov, Aleksandr, Samuel Horváth, Peter Richtárik, and Mher Safaryan. "On Biased Compression for Distributed Learning." arXiv preprint arXiv:2002.12410 (2020).

$$n = d = 3$$

$$f_1(x) = \langle a, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_2(x) = \langle b, x \rangle^2 + \frac{1}{4} \|x\|^2 \quad f_3(x) = \langle c, x \rangle^2 + \frac{1}{4} \|x\|^2$$

$$a = (-3, 2, 2)^\top \quad b = (2, -3, 2)^\top \quad c = (2, 2, -3)^\top$$

$$x^0 = (t, t, t)^\top$$

In this case Parallel SGD with Top1 compression operator satisfies

$$x^k = \left(1 + \frac{11\gamma}{6}\right)^k x^0$$

One can fix this using one special trick called **error-compensation**

4. Error-Compensated SGD



Seide, Frank, Hao Fu, Jasha Droppo, Gang Li, and Dong Yu. "**1-bit stochastic gradient descent and its application to data-parallel distributed training of speech dnns.**" *In Fifteenth Annual Conference of the International Speech Communication Association*. 2014.



Stich, Sebastian U., Jean-Baptiste Cordonnier, and Martin Jaggi. "**Sparsified SGD with memory.**" *In Advances in Neural Information Processing Systems*, pp. 4447-4458. 2018.



Karimireddy, Sai Praneeth, Quentin Rebjock, Sebastian Stich, and Martin Jaggi. "**Error Feedback Fixes SignSGD and other Gradient Compression Schemes.**" *In International Conference on Machine Learning*, pp. 3252-3261. 2019.



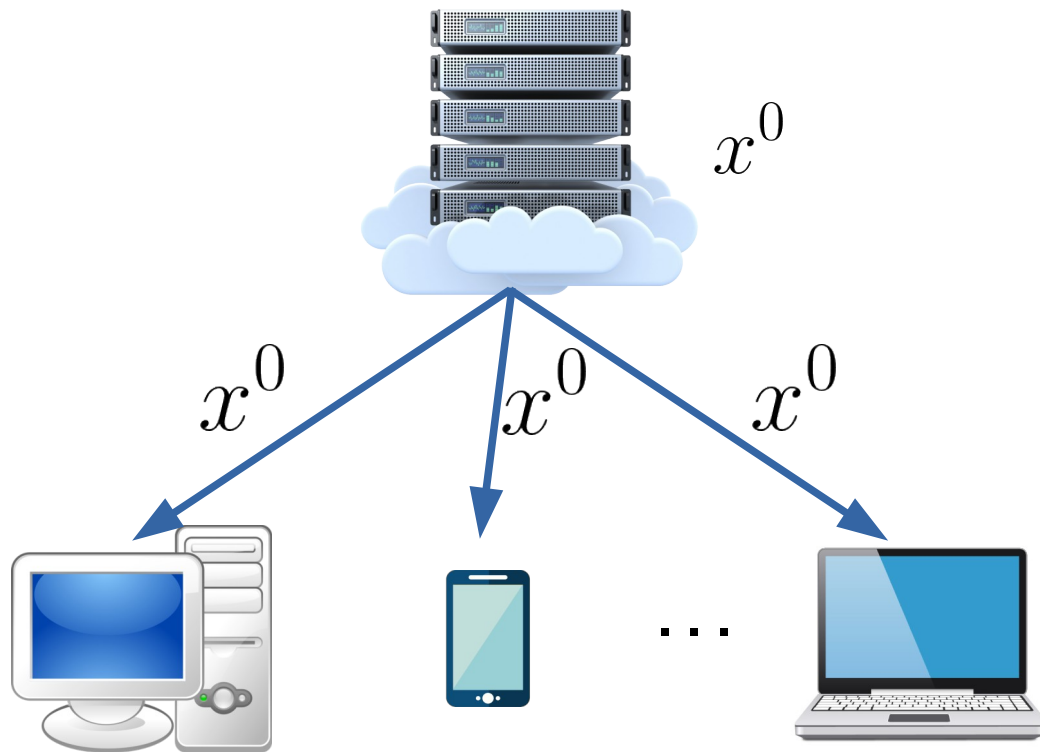
Stich, Sebastian U., and Sai Praneeth Karimireddy. "**The error-feedback framework: Better rates for SGD with delayed gradients and compressed communication.**" arXiv preprint arXiv:1909.05350 (2019).



Beznosikov, Aleksandr, Samuel Horváth, Peter Richtárik, and Mher Safaryan. "**On Biased Compression for Distributed Learning.**" arXiv preprint arXiv:2002.12410 (2020).

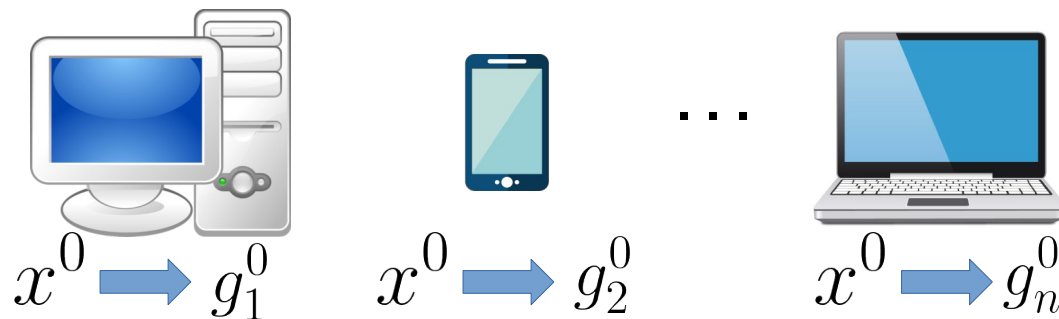
- 1 Server broadcasts the parameters

Step 1



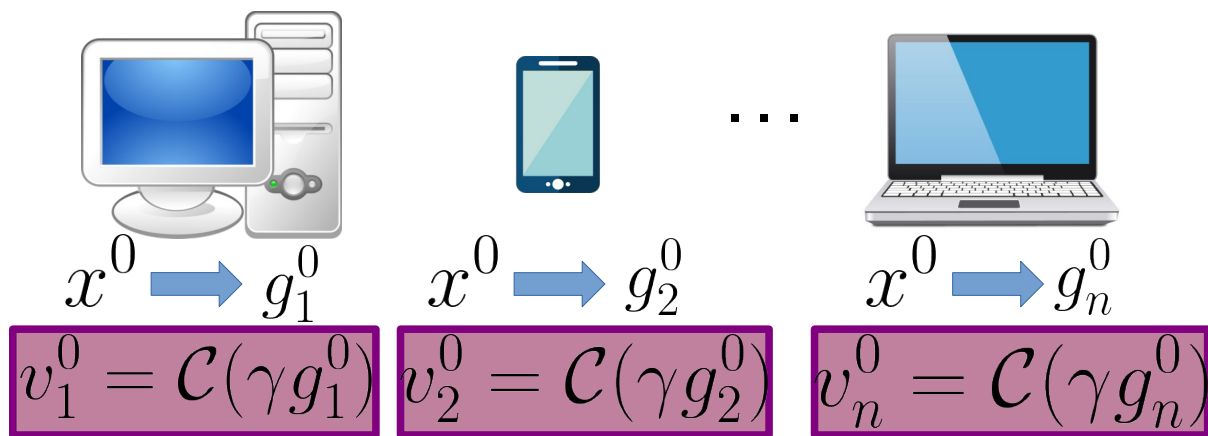
Step 1

- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel



Step 1

- 1 Server broadcasts the parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Compression



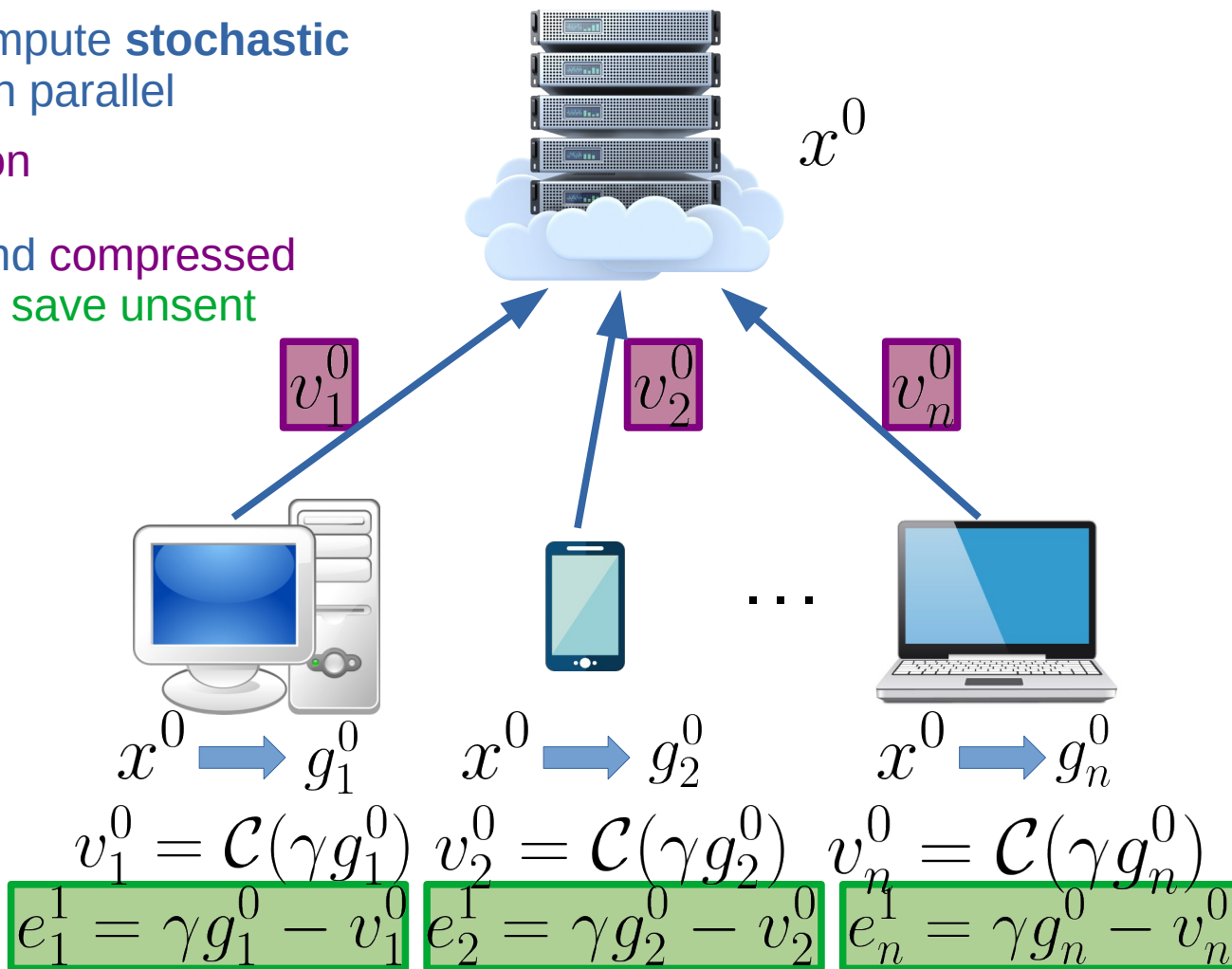
Step 1

1 Server broadcasts the parameters

2 Devices compute **stochastic gradients** in parallel

3 Compression

4 Devices send **compressed vectors** and **save unsent information**



Step 1

1 Server broadcasts the parameters

2 Devices compute **stochastic gradients** in parallel

3 **Compression**

4 Devices send **compressed vectors** and **save unsent information**

5 Server gathers the information and updates the parameters



$$x^0 \rightarrow x^1 = x^0 - \frac{1}{n} \sum_{i=1}^n v_i^0$$



$$x^0 \rightarrow g_1^0$$



$$x^0 \rightarrow g_2^0$$

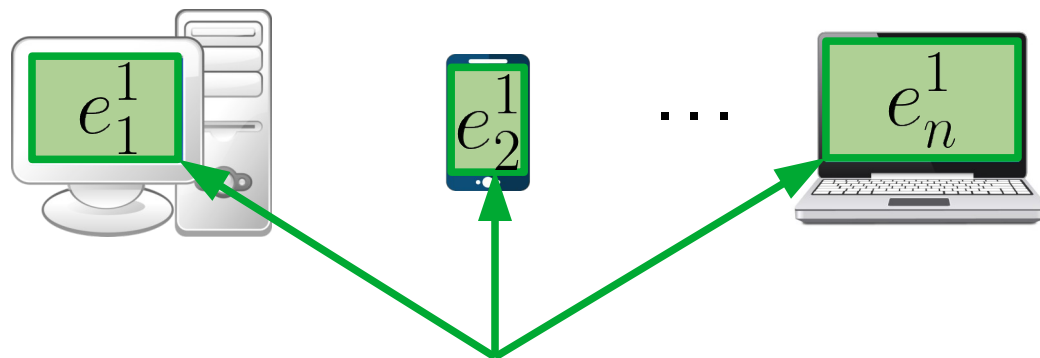
...



$$x^0 \rightarrow g_n^0$$

$$\begin{aligned} v_1^0 &= \mathcal{C}(\gamma g_1^0) & v_2^0 &= \mathcal{C}(\gamma g_2^0) & v_n^0 &= \mathcal{C}(\gamma g_n^0) \\ e_1^1 &= \gamma g_1^0 - v_1^0 & e_2^1 &= \gamma g_2^0 - v_2^0 & e_n^1 &= \gamma g_n^0 - v_n^0 \end{aligned}$$

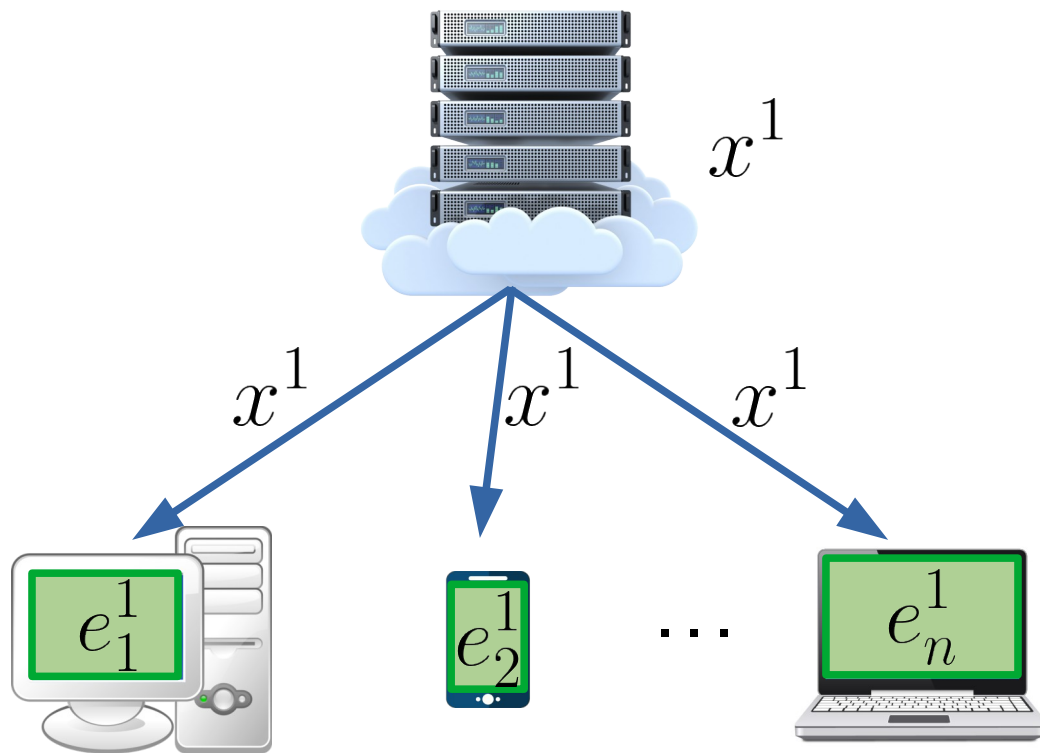
Step 1



devices keep these vectors for the next iterations
to *partially* send them later

- 1 Server broadcasts new parameters

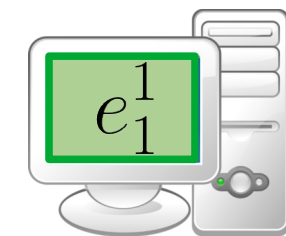
Step 2



Step 2

1 Server broadcasts new parameters

2 Devices compute **stochastic gradients** in parallel

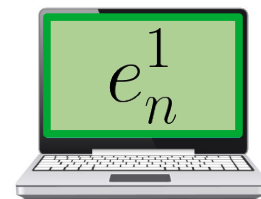


$$x^1 \rightarrow g_1^1$$



$$x^1 \rightarrow g_2^1$$

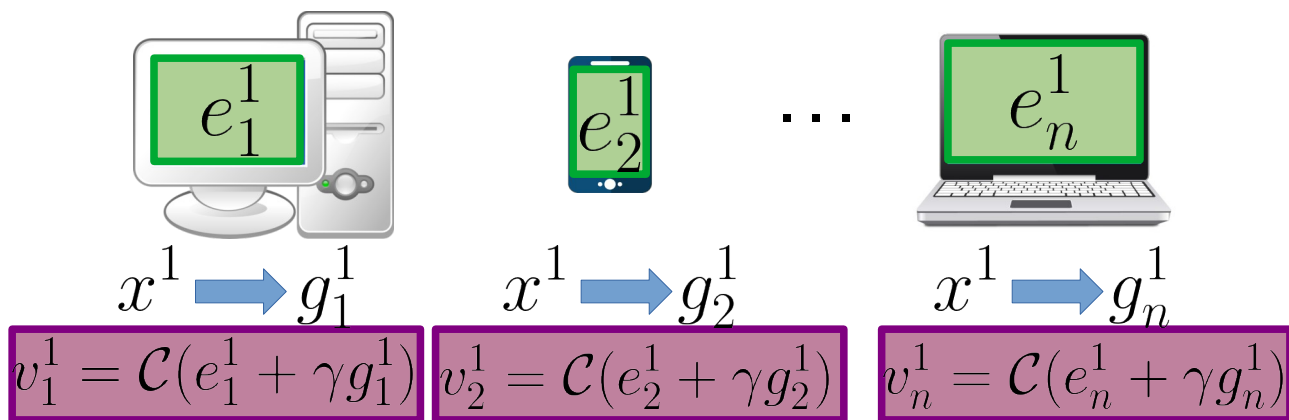
...



$$x^1 \rightarrow g_n^1$$

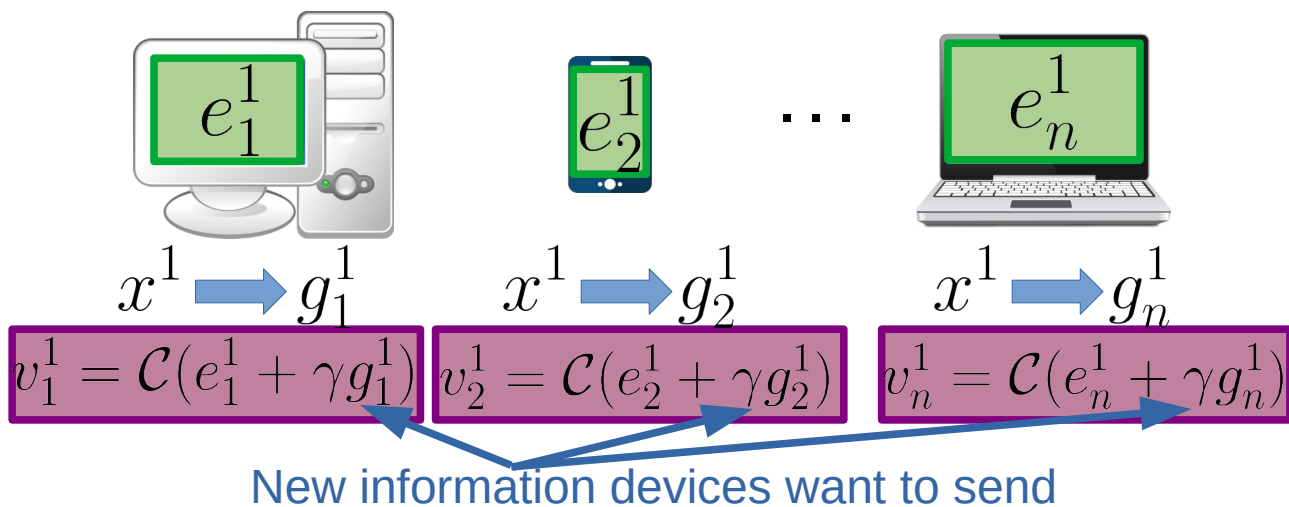
Step 2

- 1 Server broadcasts new parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Compression



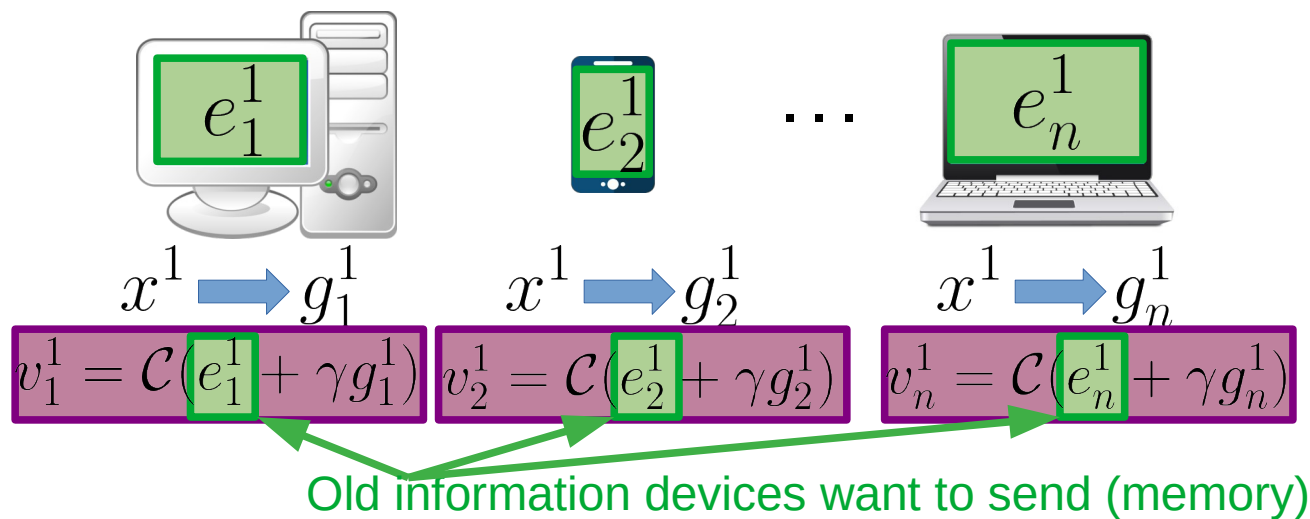
Step 2

- 1 Server broadcasts new parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Compression



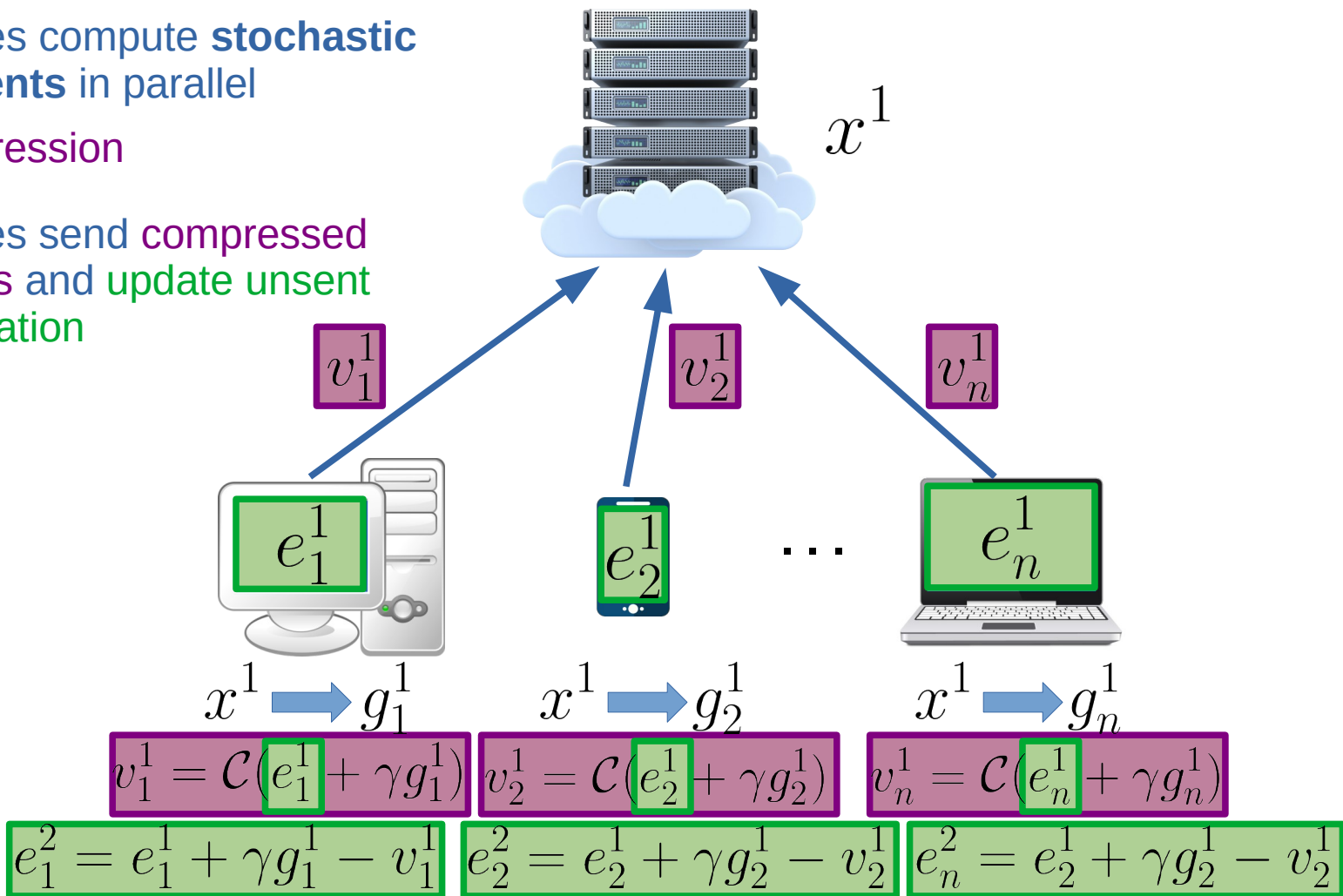
Step 2

- 1 Server broadcasts new parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 Compression



Step 2

- 1 Server broadcasts new parameters
- 2 Devices compute **stochastic gradients** in parallel
- 3 **Compression**
- 4 Devices send **compressed vectors** and **update unsent information**



Step 2

1 Server broadcasts new parameters

2 Devices compute **stochastic gradients** in parallel

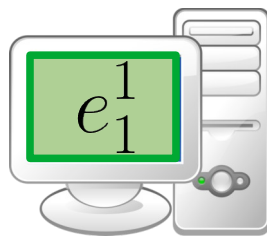
3 **Compression**

4 Devices send **compressed vectors** and **update unsent information**

5 Server gathers the information and updates the parameters



$$x^1 \rightarrow x^2 = x^1 - \frac{1}{n} \sum_{i=1}^n v_i^1$$

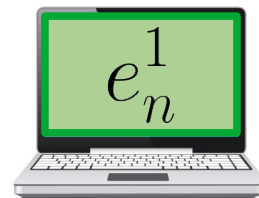


$$x^1 \rightarrow g_1^1$$



$$x^1 \rightarrow g_2^1$$

...



$$x^1 \rightarrow g_n^1$$

$$v_1^1 = \mathcal{C}(e_1^1 + \gamma g_1^1) \quad v_2^1 = \mathcal{C}(e_2^1 + \gamma g_2^1) \quad v_n^1 = \mathcal{C}(e_n^1 + \gamma g_n^1)$$

$$e_1^2 = e_1^1 + \gamma g_1^1 - v_1^1 \quad e_2^2 = e_2^1 + \gamma g_2^1 - v_2^1 \quad e_n^2 = e_n^1 + \gamma g_n^1 - v_n^1$$

Step 2

1 Server broadcasts new parameters

2 Devices compute **stochastic gradients** in parallel

3 **Compression**

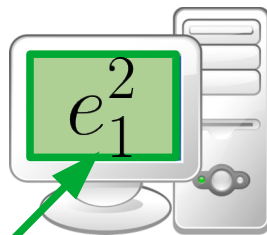
4 Devices send **compressed vectors** and **update unsent information**

5 Server gathers the information and updates the parameters

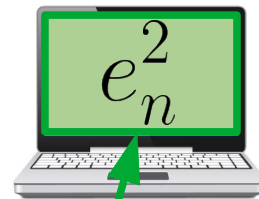
6 **Devices update their memory**



$$x^1 \rightarrow x^2 = x^1 - \frac{1}{n} \sum_{i=1}^n v_i^1$$



...



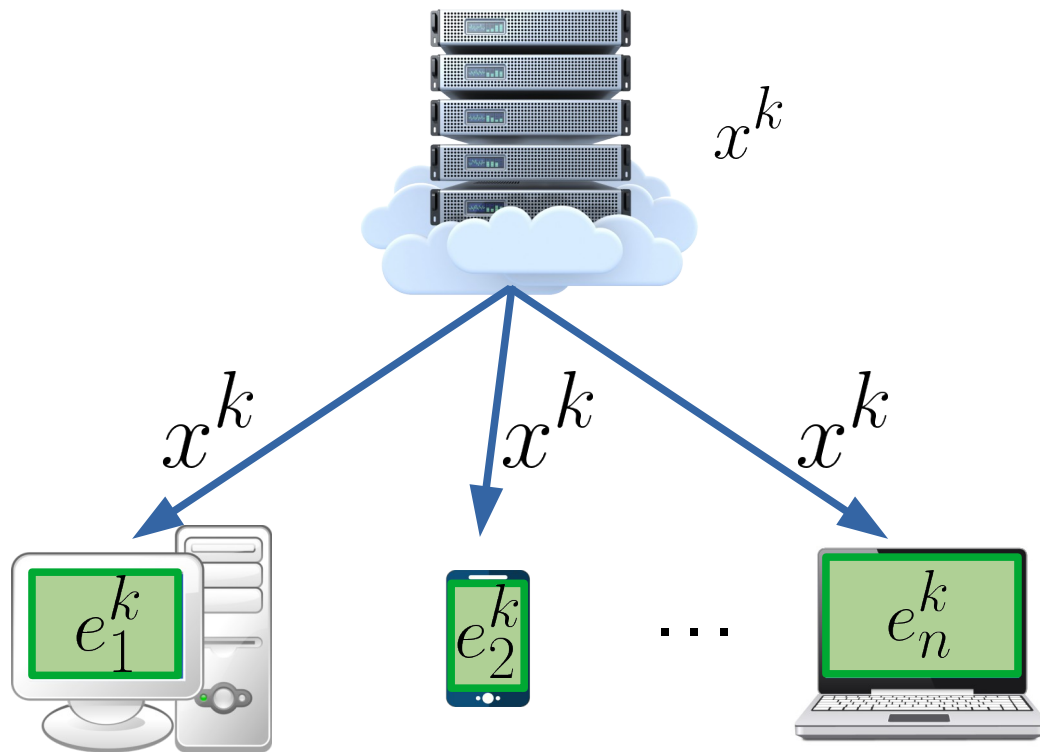
$$e_1^2 = e_1^1 + \gamma g_1^1 - v_1^1$$

$$e_2^2 = e_2^1 + \gamma g_2^1 - v_2^1$$

$$e_n^2 = e_n^1 + \gamma g_n^1 - v_n^1$$

1 Server broadcasts new parameters

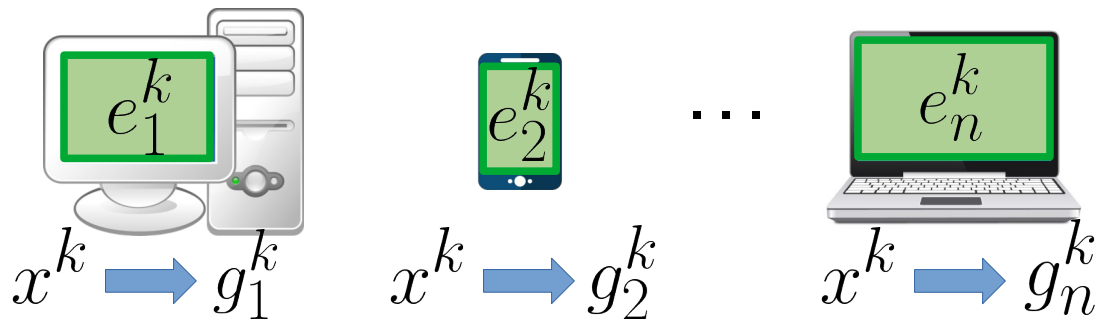
Step $k+1$



Step $k+1$

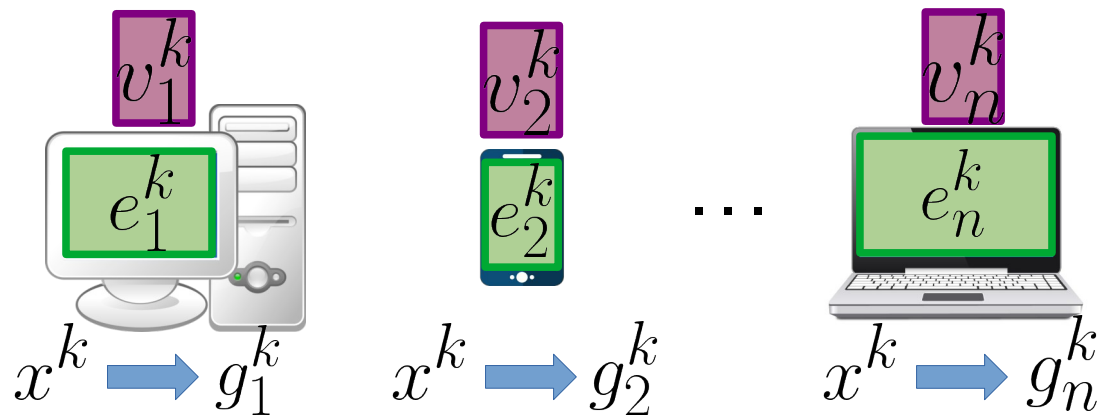
1 Server broadcasts new parameters

2 Workers compute **stochastic gradients** in parallel



Step $k+1$

- 1 Server broadcasts new parameters
- 2 Workers compute **stochastic gradients** in parallel
- 3 Compression



$$v_i^k = \mathcal{C} \left(e_i^k + \gamma g_i^k \right)$$

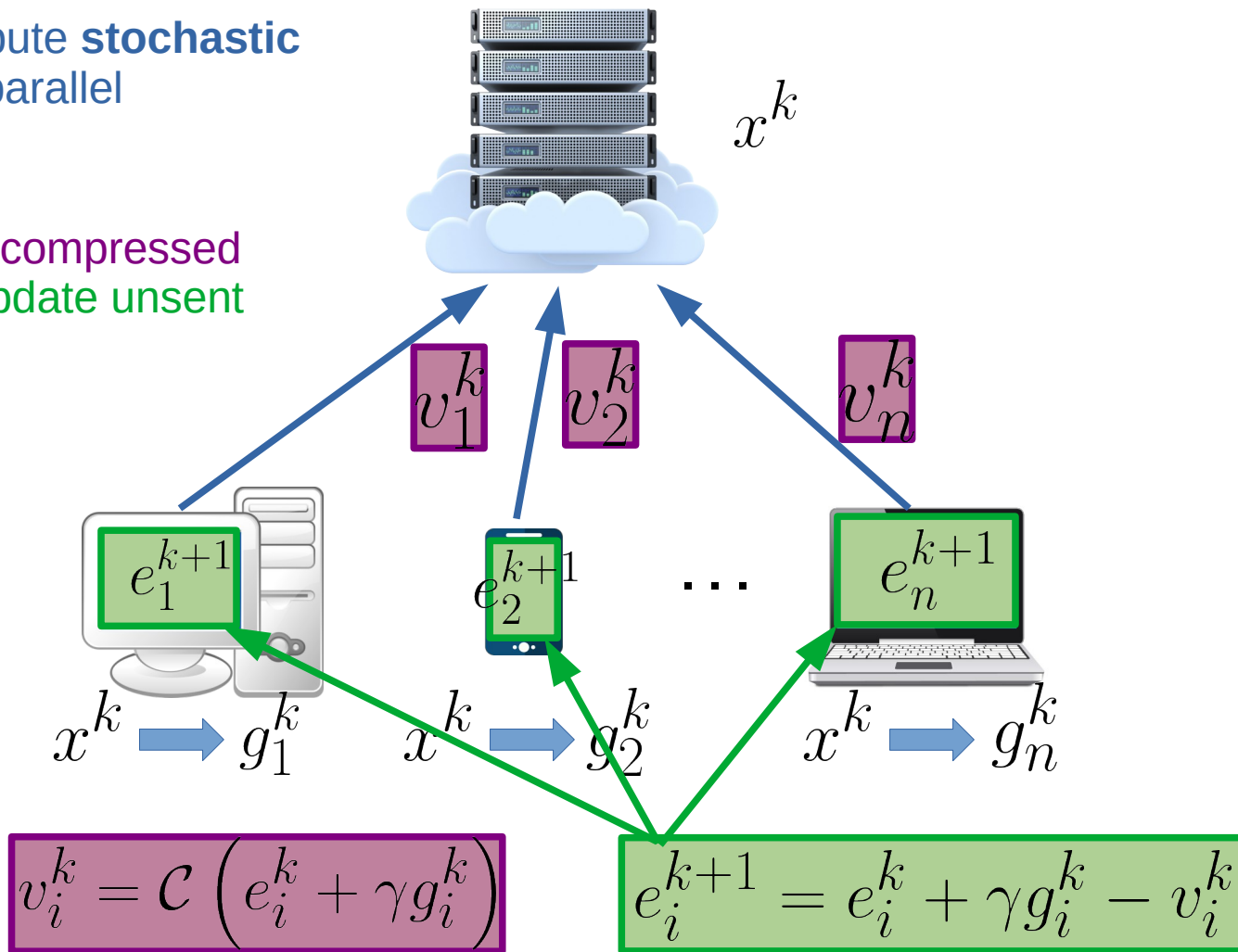
Step $k+1$

1 Server broadcasts new parameters

2 Workers compute **stochastic gradients** in parallel

3 Compression

4 Devices send **compressed vectors** and **update unsent information**



Step $k+1$

1 Server broadcasts new parameters

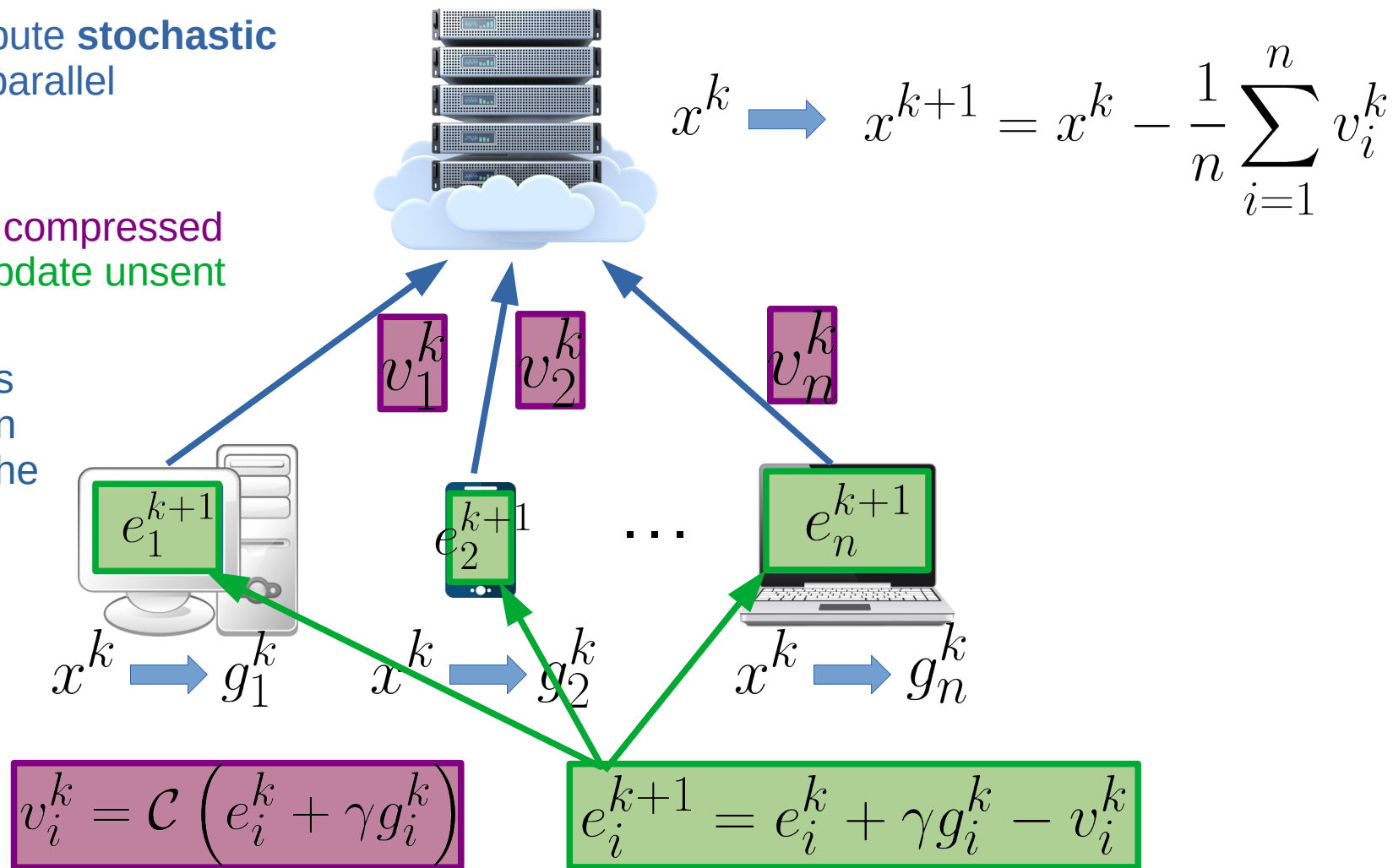
2 Workers compute **stochastic gradients** in parallel

3 **Compression**

4 Devices send **compressed vectors** and **update unsent information**

5 Server gathers the information and updates the parameters

6 Repeat steps 1 – 5



Error-Compensated SGD



Converges even with biased compression operators

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

$$\tilde{\mathcal{O}} \left(\frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right) \text{ iterations}$$

Error-Compensated SGD



Converges even with biased compression operators

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical
factors

$$\rightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right) \text{ iterations}$$

Error-Compensated SGD



Converges even with biased compression operators

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical
factors

$$\rightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta \mu} + \frac{\sigma^2}{n \mu \varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2 / \delta)}}{\mu \sqrt{\delta} \varepsilon} \right) \text{ iterations}$$

$$\mathbb{E} \|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

Error-Compensated SGD



Converges even with biased compression operators

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical
factors

$$\rightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta \mu} + \frac{\sigma^2}{n \mu \varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2 / \delta)}}{\mu \sqrt{\delta \varepsilon}} \right) \text{ iterations}$$

$$\mathbb{E} \|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

$$\mathbb{E} \left[\|g_i^k - \nabla f_i(x^k)\|^2 \mid x^k \right] \leq \sigma^2$$

Error-Compensated SGD



Converges even with biased compression operators

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical
factors

$$\tilde{\mathcal{O}} \left(\frac{L}{\delta \mu} + \frac{\sigma^2}{n \mu \varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2 / \delta)}}{\mu \sqrt{\delta \varepsilon}} \right) \text{ iterations}$$

$$\mathbb{E} \|\mathcal{C}(x) - x\|^2 \leq (1 - \delta) \|x\|^2$$

$$\mathbb{E} \left[\|g_i^k - \nabla f_i(x^k)\|^2 \mid x^k \right] \leq \sigma^2$$

$$\zeta_*^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(x^*)\|^2$$

Error-Compensated SGD



Converges even with biased compression operators



Fails to converge with **linear rate** even when workers compute full gradients

EC-SGD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical
factors

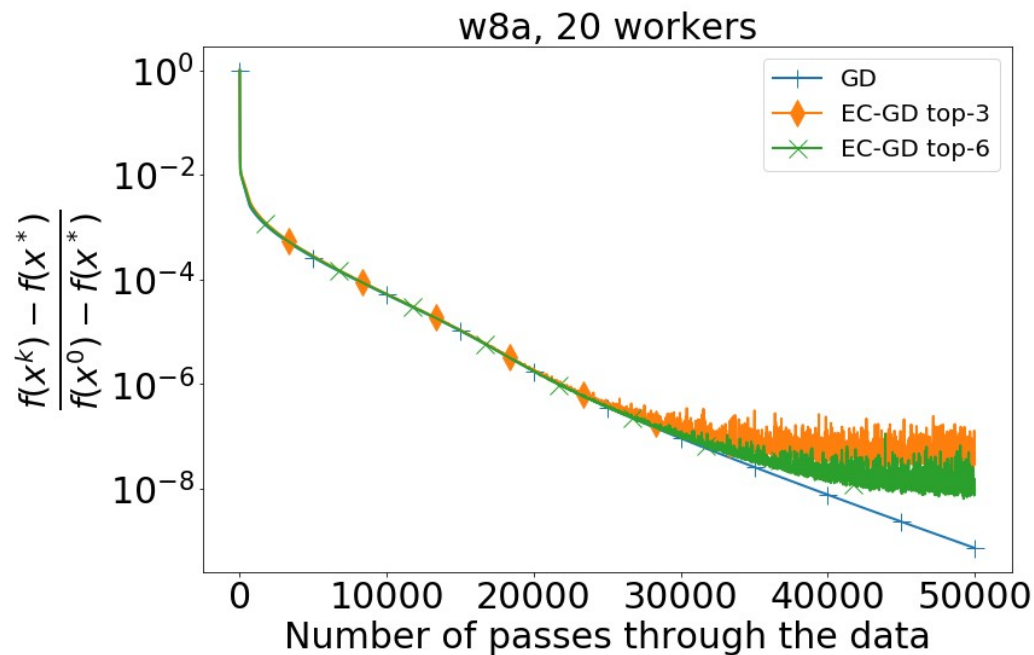
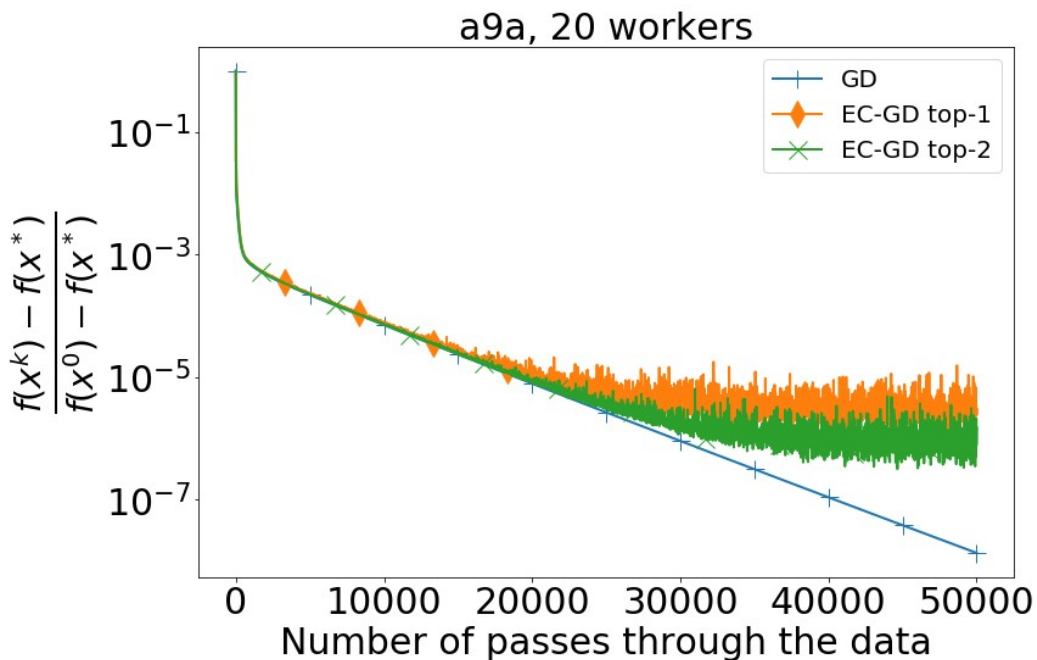
$$\rightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L(\sigma^2 + \zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right) \text{ iterations}$$

$$\mathbb{E}\|\mathcal{C}(x) - x\|^2 \leq (1 - \delta)\|x\|^2 \quad \mathbb{E}\left[\|g_i^k - \nabla f_i(x^k)\|^2 \mid x^k\right] \leq \sigma^2$$

$$\zeta_*^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(x^*)\|^2$$

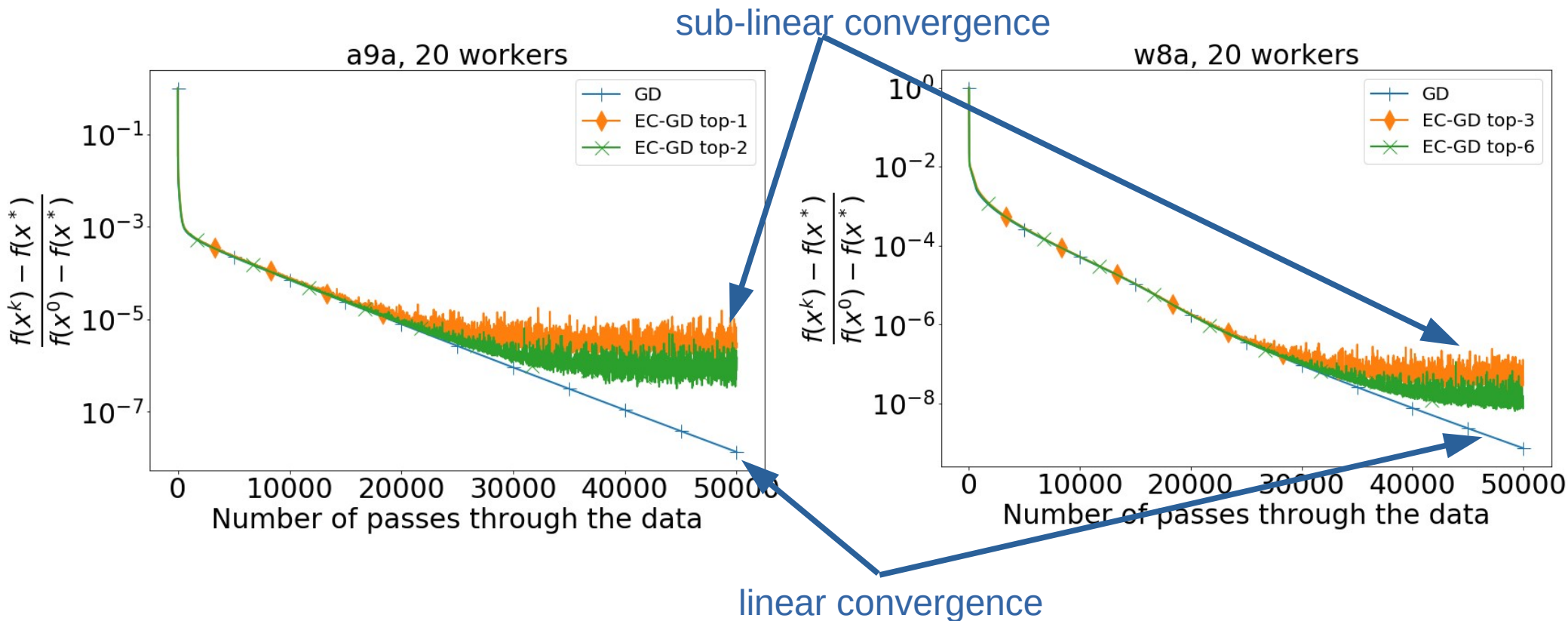
EC-GD and Logistic Regression

$$\min_{x \in \mathbb{R}^d} \left\{ f(x) = \frac{1}{N} \sum_{i=1}^N \log(1 + \exp(-y_i \cdot (Ax)_i)) + \frac{\mu}{2} \|x\|^2 \right\}$$



EC-GD and Logistic Regression

$$\min_{x \in \mathbb{R}^d} \left\{ f(x) = \frac{1}{N} \sum_{i=1}^N \log(1 + \exp(-y_i \cdot (Ax)_i)) + \frac{\mu}{2} \|x\|^2 \right\}$$



Error-Compensated SGD



Converges even with biased compression operators



Fails to converge with **linear rate** even when workers compute full gradients

Questions:

1

*Is it possible to design **linearly converging** SGD with error compensation when workers compute full gradients, i.e., linearly converging **EC-GD**?*

2

*Is it possible to design **linearly converging SGD** with error compensation when **the local loss functions have a finite-sum form**?*

Error-Compensated SGD



Converges even with biased compression operators



Fails to converge with **linear rate** even when workers compute full gradients

Questions:

1

*Is it possible to design **linearly converging** SGD with error compensation when workers compute full gradients, i.e., linearly converging **EC-GD**?*

2

*Is it possible to design **linearly converging SGD** with error compensation when **the local loss functions have a finite-sum form**?*

The answer is Yes for both questions

5.1. New method: EC-GDstar

Error-Compensated GD

EC-GD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical factors $\rightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta\mu} + \frac{\sqrt{L\zeta_*^2}}{\mu\delta\sqrt{\varepsilon}} \right)$ iterations

$$\zeta_*^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(x^*)\|^2$$

Error-Compensated GD

EC-GD finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

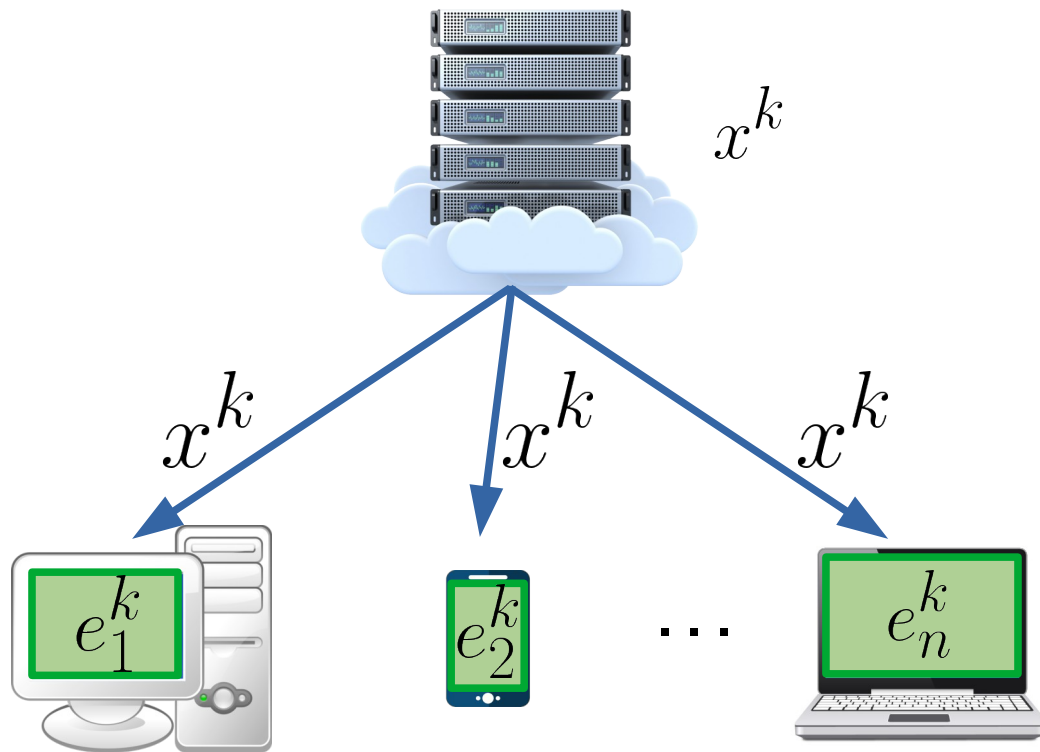
Hides logarithmical factors $\longrightarrow \tilde{\mathcal{O}} \left(\frac{L}{\delta\mu} + \frac{\sqrt{L\zeta_*^2}}{\mu\delta\sqrt{\varepsilon}} \right)$ iterations

$$\zeta_*^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(x^*)\|^2$$

What if devices know these vectors from the beginning?

EC-GDstar

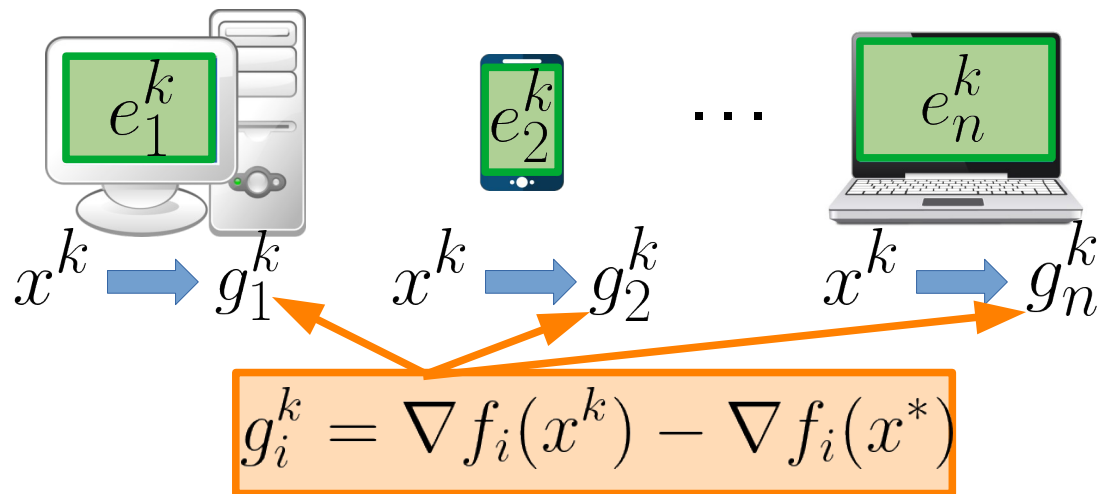
- 1 Server broadcasts new parameters



EC-GDstar

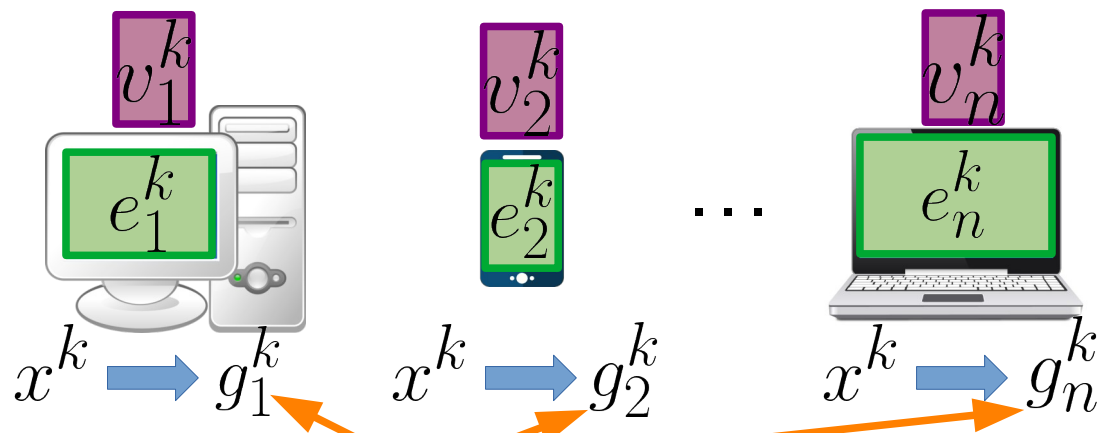
1 Server broadcasts new parameters

2 Workers compute **shifted** gradients in parallel



EC-GDstar

- 1 Server broadcasts new parameters
- 2 Workers compute **shifted gradients** in parallel
- 3 Compression

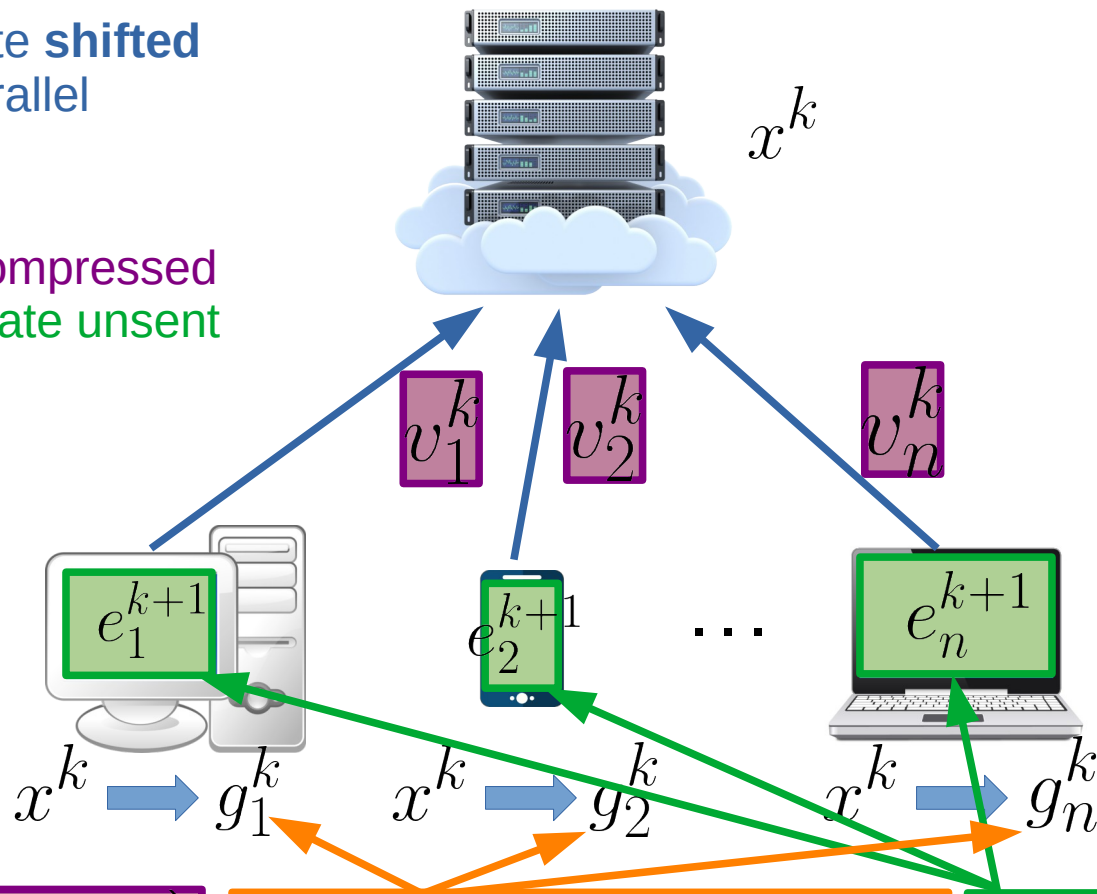


$$v_i^k = \mathcal{C} \left(e_i^k + \gamma g_i^k \right)$$

$$g_i^k = \nabla f_i(x^k) - \nabla f_i(x^*)$$

EC-GDstar

- 1 Server broadcasts new parameters
- 2 Workers compute **shifted gradients** in parallel
- 3 Compression
- 4 Devices send **compressed vectors** and **update unsent information**



$$v_i^k = \mathcal{C} \left(e_i^k + \gamma g_i^k \right)$$

$$g_i^k = \nabla f_i(x^k) - \nabla f_i(x^*)$$

$$e_i^{k+1} = e_i^k + \gamma g_i^k - v_i^k$$

Step $k+1$

1 Server broadcasts new parameters

2 Workers compute **shifted gradients** in parallel

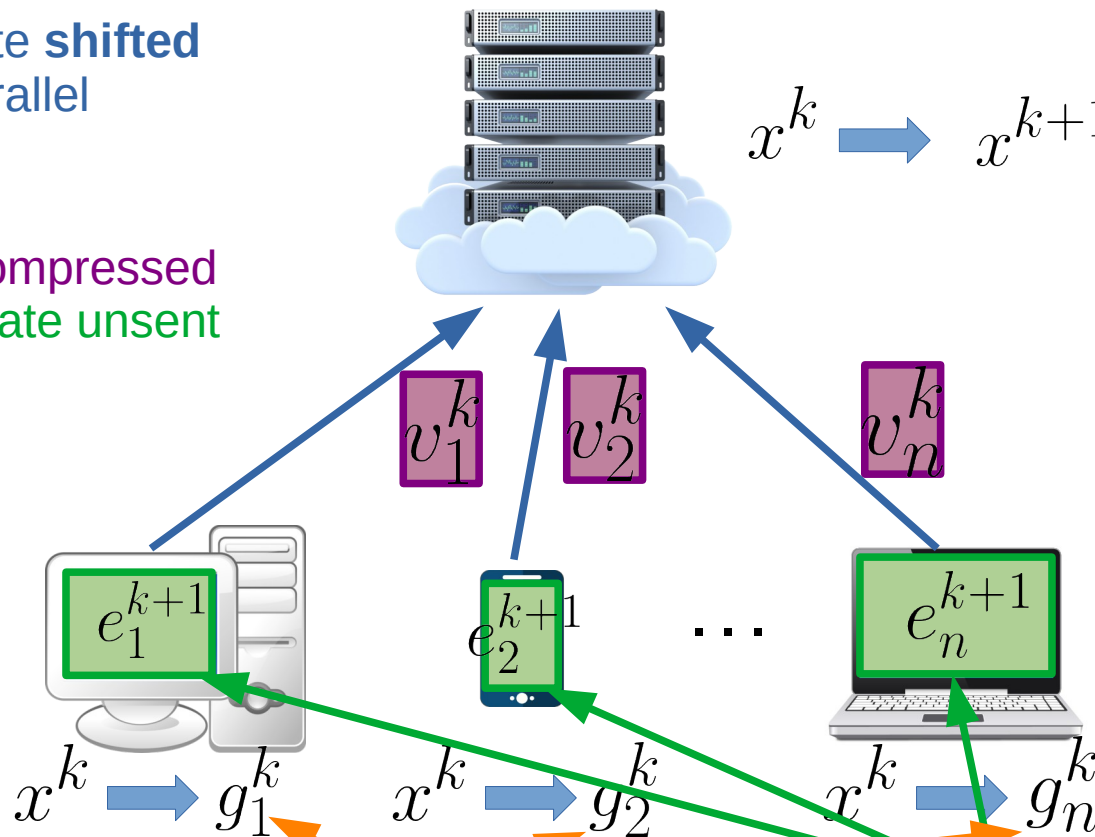
3 **Compression**

4 Devices send **compressed vectors** and **update unsent information**

5 Server gathers the information and updates the parameters

6 Repeat steps 1 – 5

$$x^k \rightarrow x^{k+1} = x^k - \frac{1}{n} \sum_{i=1}^n v_i^k$$



$$v_i^k = \mathcal{C} \left(e_i^k + \gamma g_i^k \right)$$

$$g_i^k = \nabla f_i(x^k) - \nabla f_i(x^*)$$

$$e_i^{k+1} = e_i^k + \gamma g_i^k - v_i^k$$

EC-GDstar: Rate of Convergence

EC-GDstar finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

$$\mathcal{O}\left(\frac{L}{\delta\mu} \ln \frac{1}{\varepsilon}\right) \text{ iterations}$$

EC-GDstar: Rate of Convergence

EC-GDstar finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

$$\mathcal{O}\left(\frac{L}{\delta\mu} \ln \frac{1}{\varepsilon}\right) \text{ iterations}$$



Linear rate



The method is impractical: it uses the gradients at the solution

EC-GDstar: Rate of Convergence

EC-GDstar finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

$$\mathcal{O}\left(\frac{L}{\delta\mu} \ln \frac{1}{\varepsilon}\right) \text{ iterations}$$



Linear rate



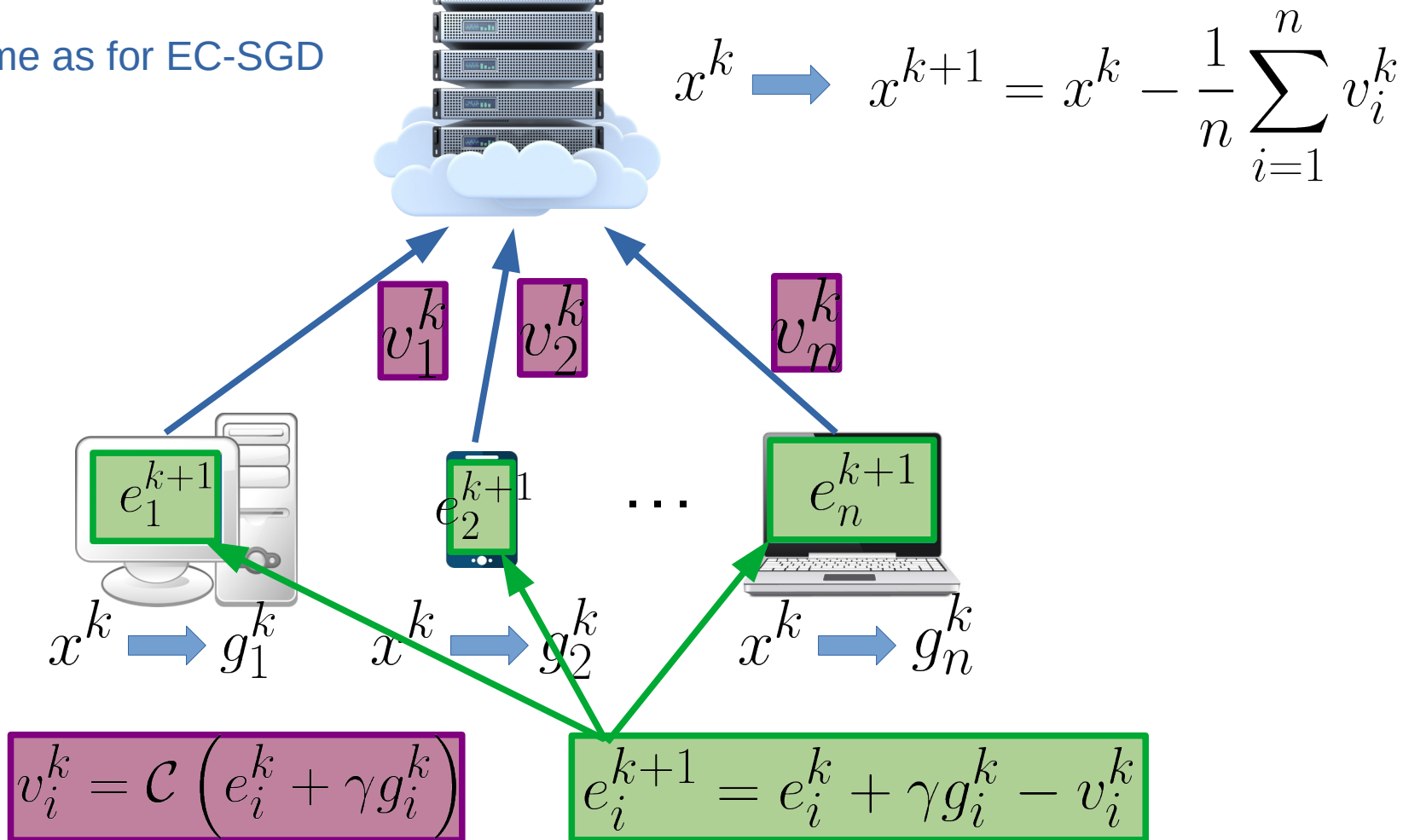
The method is impractical: it uses the gradients at the solution

Can we develop a practical analog?

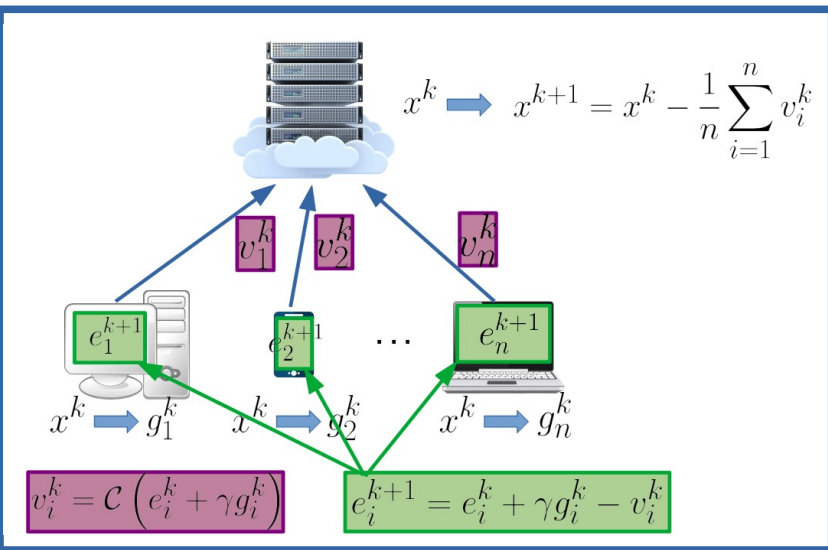
5.2. New method: EC-SGD-DIANA

EC-SGD-DIANA

The same scheme as for EC-SGD



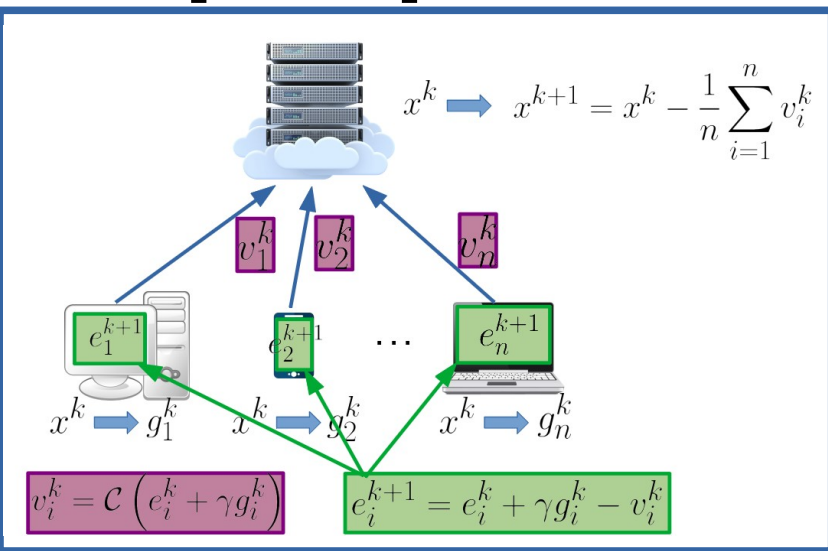
$$g_i^k = \hat{g}_i^k - h_i^k + h^k$$



EC-SGD-DIANA

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k$$

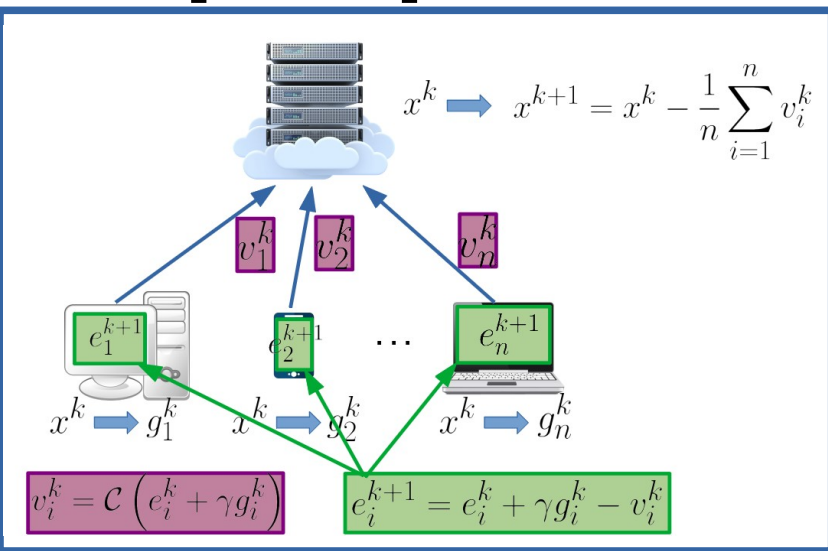
$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$



EC-SGD-DIANA

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k$$

$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$



Works for both cases:

- $\bullet$ $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)]$ $\boxed{\hat{g}_i^k} = \nabla f_{\xi_i}(x^k)$
 - $\bullet$ $f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$ $\boxed{\hat{g}_i^k} = \nabla f_{il}(x^k)$
- $l \sim [m]$ uniformly at random

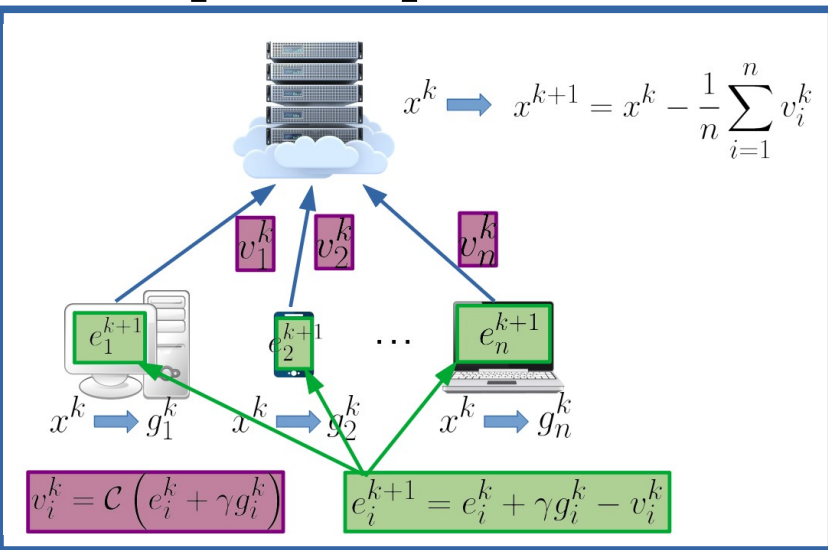
EC-SGD-DIANA

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k$$

stepsize

$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$

$$h_i^{k+1} = h_i^k + \alpha \mathcal{Q} \left(\hat{g}_i^k - h_i^k \right)$$



Works for both cases:

- $\bullet$ $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)]$ $\boxed{\hat{g}_i^k} = \nabla f_{\xi_i}(x^k)$
 - $\bullet$ $f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$ $\boxed{\hat{g}_i^k} = \nabla f_{il}(x^k)$
- $l \sim [m]$ uniformly at random

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k$$

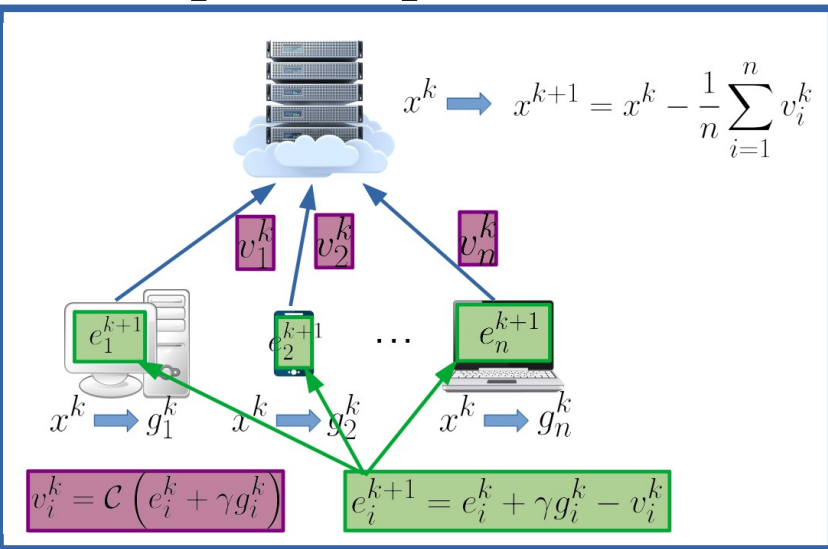
$$\mathbb{E} \left[\hat{g}_i^k \mid x^k \right] = \nabla f_i(x^k)$$

The diagram shows the update rule $h_i^{k+1} = h_i^k + \alpha \mathcal{Q}(\hat{g}_i^k - h_i^k)$. A blue arrow labeled "stepsize" points to the α term. An orange box highlights the term $\mathcal{Q}(\hat{g}_i^k - h_i^k)$, with an orange arrow labeled "Workers send" pointing to it.

Works for both cases:

● $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)]$ $\hat{g}_i^k$ $= \nabla f_{\xi_i}(x^k)$

$$\bullet \quad f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x) \quad \boxed{\hat{g}_i^k} = \nabla f_{il}(x^k) \\ l \sim [m] \text{ uniformly at random}$$



EC-SGD-DIANA

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k \leftarrow h^k = \frac{1}{n} \sum_{i=1}^n h_i^k$$

stepsize

Server broadcasts
this vector to the
workers

$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$

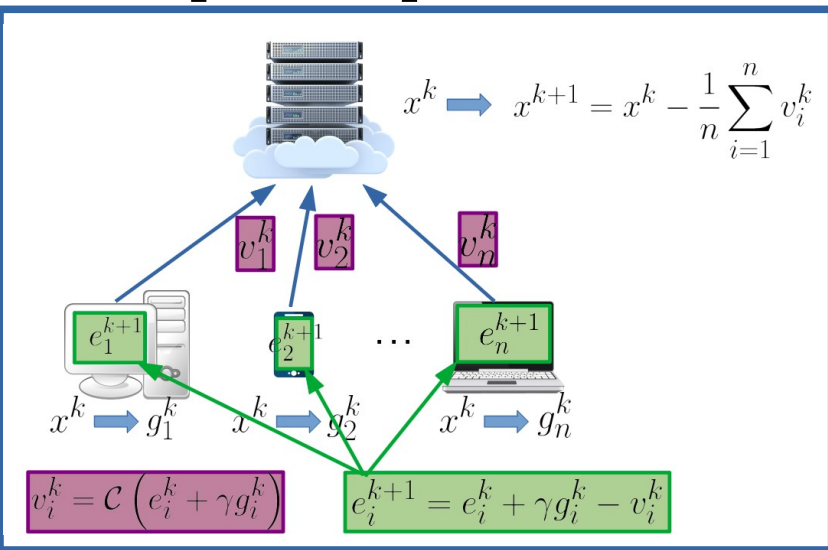
$$h_i^{k+1} = h_i^k + \alpha \boxed{\mathcal{Q} \left(\hat{g}_i^k - h_i^k \right)}$$

Workers send
these vectors
to the server

Works for both cases:

● $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)] \quad \boxed{\hat{g}_i^k} = \nabla f_{\xi_i}(x^k)$

● $f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x) \quad \boxed{\hat{g}_i^k} = \nabla f_{il}(x^k)$
 $l \sim [m]$ uniformly at random



EC-SGD-DIANA

The key insight why do we need $\{h_i^k\}_{i=1}^n$:

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k \leftarrow h^k = \frac{1}{n} \sum_{i=1}^n h_i^k$$

stepsize

Server broadcasts this vector to the workers

Workers send these vectors to the server

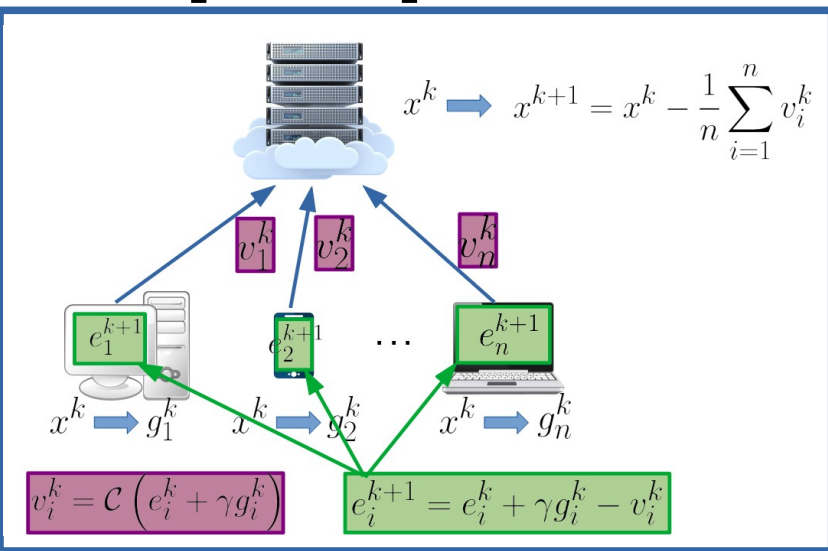
$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$

$$h_i^{k+1} = h_i^k + \alpha \mathcal{Q} \left(\boxed{\hat{g}_i^k} - h_i^k \right)$$

Works for both cases:

● $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)] \quad \boxed{\hat{g}_i^k} = \nabla f_{\xi_i}(x^k)$

● $f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x) \quad \boxed{\hat{g}_i^k} = \nabla f_{il}(x^k)$
 $l \sim [m] \text{ uniformly at random}$



EC-SGD-DIANA

The key insight why do we need $\{h_i^k\}_{i=1}^n$:

it reduces the variance coming from compressions via learning the gradients at the solution!

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k \leftarrow h^k = \frac{1}{n} \sum_{i=1}^n h_i^k$$

stepsize

Server broadcasts this vector to the workers

Workers send these vectors to the server

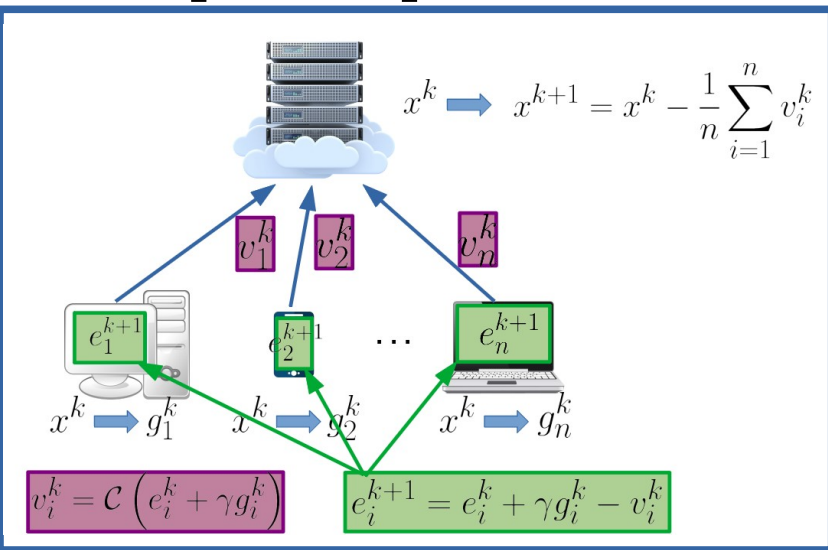
$$\mathbb{E} \left[\boxed{\hat{g}_i^k} \mid x^k \right] = \nabla f_i(x^k)$$

$$h_i^{k+1} = h_i^k + \alpha \boxed{\mathcal{Q} \left(\hat{g}_i^k - h_i^k \right)}$$

Works for both cases:

● $f_i(x) = \mathbf{E}_{\xi_i \sim \mathcal{D}_i} [f_{\xi_i}(x)] \quad \boxed{\hat{g}_i^k} = \nabla f_{\xi_i}(x^k)$

● $f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x) \quad \boxed{\hat{g}_i^k} = \nabla f_{il}(x^k)$
 $l \sim [m]$ uniformly at random



EC-SGD-DIANA: Rate of Convergence

EC-SGD-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

EC-SGD-DIANA: Rate of Convergence

EC-SGD-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ iterations

Option II: $\tilde{\mathcal{O}} \left(\frac{1 + \omega}{\delta} + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$

EC-SGD-DIANA: Rate of Convergence

EC-SGD-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical factors

Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ iterations

Option II: $\tilde{\mathcal{O}} \left(\frac{1 + \omega}{\delta} + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$

EC-SGD-DIANA: Rate of Convergence

EC-SGD-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

Hides logarithmical factors

Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ iterations

Option II: $\tilde{\mathcal{O}} \left(\frac{1 + \omega}{\delta} + \frac{L}{\delta\mu} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$

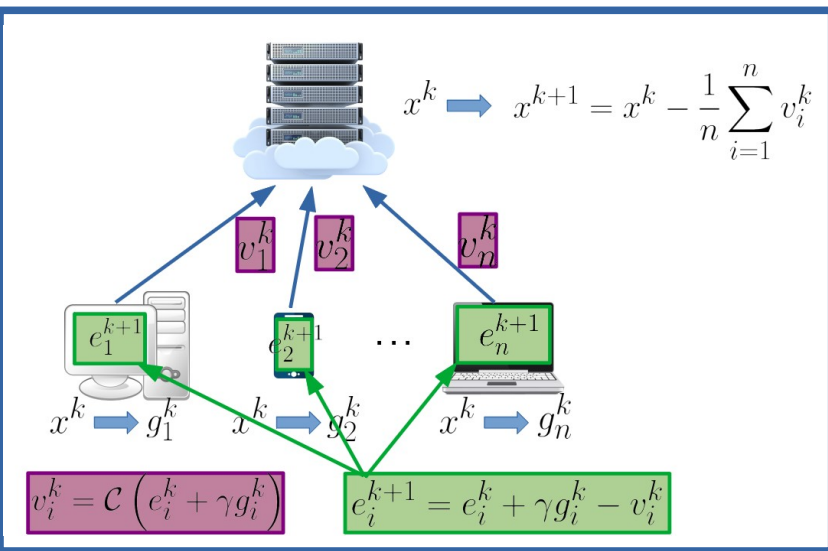
Moreover, if workers compute **full gradients**, then the rate of convergence is linear

$$\mathcal{O} \left(\left(\omega + \frac{L}{\delta\mu} \right) \log \frac{1}{\varepsilon} \right)$$

5.3. New method: EC-LSVRG-DIANA

EC-LSVRG-DIANA

$$g_i^k = \hat{g}_i^k - h_i^k + h^k$$



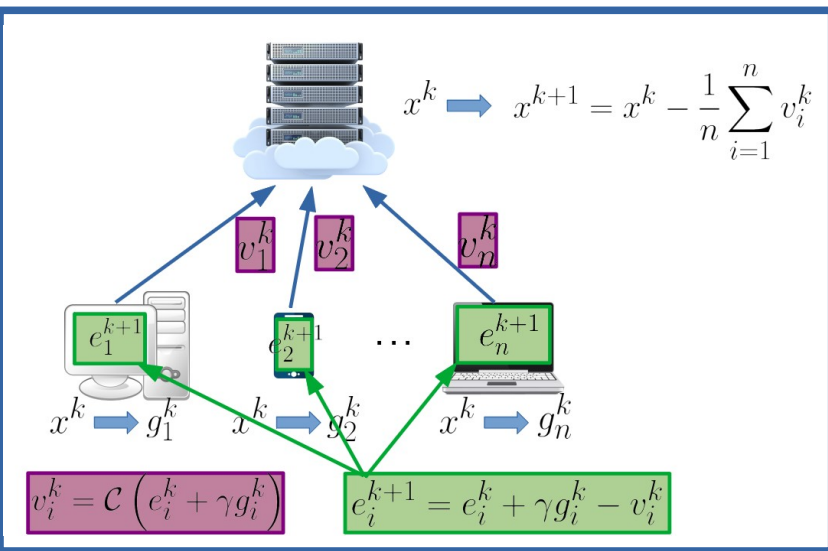
Works for the case:

$$\bullet f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$$

EC-LSVRG-DIANA

$$g_i^k = \hat{g}_i^k - h_i^k + h^k$$

$$\hat{g}_i^k = \nabla f_{il} \left(x^k \right) - \nabla f_{il} \left(w_i^k \right) + \nabla f_i \left(w_i^k \right)$$



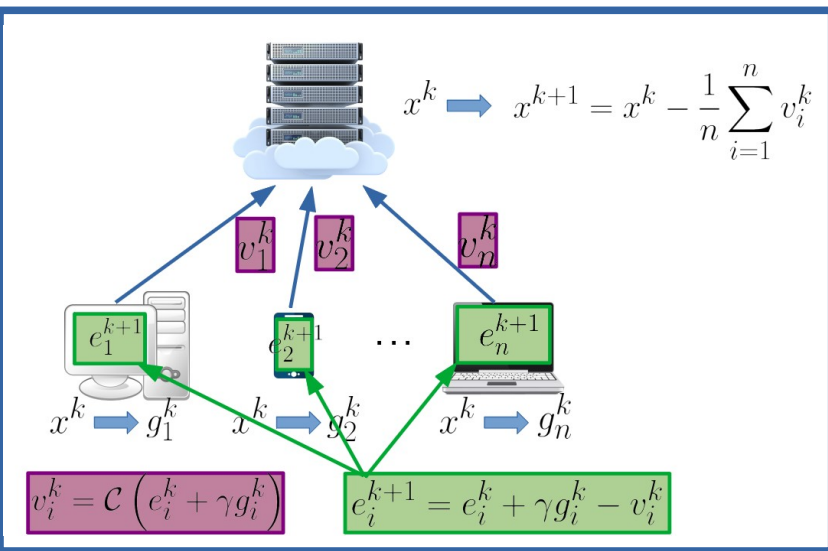
Works for the case:

$$\bullet \quad f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$$

EC-LSVRG-DIANA

$$g_i^k = \hat{g}_i^k - h_i^k + h^k \quad l \sim [m] \text{ uniformly at random}$$

$$\hat{g}_i^k = \nabla f_{il} \left(x^k \right) - \nabla f_{il} \left(w_i^k \right) + \nabla f_i \left(w_i^k \right)$$



Works for the case:

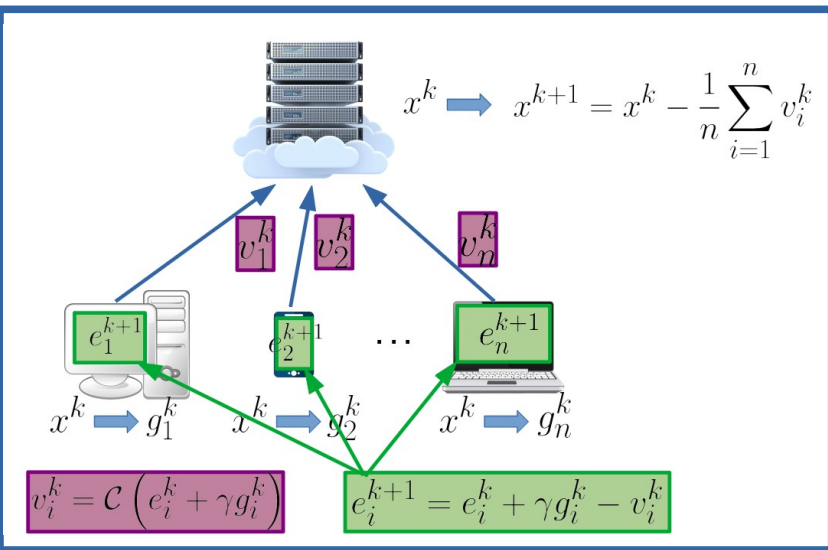
$$\bullet \quad f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$$

EC-LSVRG-DIANA

$$g_i^k = \hat{g}_i^k - h_i^k + h^k \quad l \sim [m] \text{ uniformly at random}$$

$$\hat{g}_i^k = \nabla f_{il} \left(x^k \right) - \nabla f_{il} \left(w_i^k \right) + \nabla f_i \left(w_i^k \right)$$

$$w_i^{k+1} = \begin{cases} x^k, & \text{with probability } p \\ w_i^k, & \text{with probability } 1 - p \end{cases}$$



Works for the case:

$$\bullet f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$$

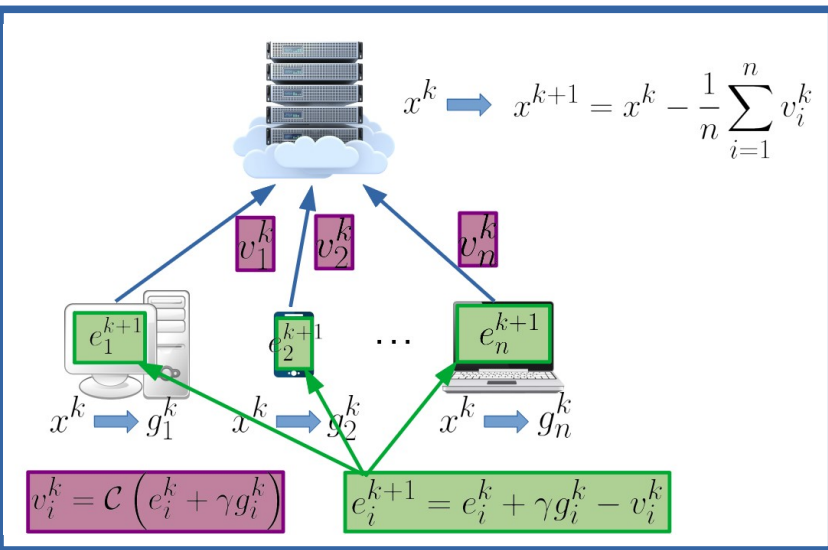
EC-LSVRG-DIANA

$$g_i^k = \boxed{\hat{g}_i^k} - h_i^k + h^k \quad \boxed{l} \sim [m] \text{ uniformly at random}$$

$$\boxed{\hat{g}_i^k} = \nabla f_{i\boxed{l}}(x^k) - \nabla f_{i\boxed{l}}(\boxed{w_i^k}) + \nabla f_i(\boxed{w_i^k})$$

Reduction of the variance introduced due to the stochasticity of the gradients

$$\boxed{w_i^{k+1}} = \begin{cases} x^k, & \text{with probability } p \\ \boxed{w_i^k}, & \text{with probability } 1 - p \end{cases}$$



Works for the case:

$$\bullet f_i(x) = \frac{1}{m} \sum_{j=1}^m f_{ij}(x)$$

EC-LSVRG-DIANA: Rate of Convergence

EC-LSVRG-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

EC-LSVRG-DIANA: Rate of Convergence

EC-LSVRG-DIANA finds such $\hat{x}$ that $\mathbb{E}[f(\hat{x})] - f(x^*) \leq \varepsilon$ after

$$\mathcal{O} \left(\left(\omega + m + \frac{L}{\delta\mu} \right) \log \frac{1}{\epsilon} \right) \text{ iterations}$$

6.1. Unified Theory of SGD



Gorbunov, Eduard, Filip Hanzely, and Peter Richtárik. **"A Unified Theory of SGD: Variance reduction, Sampling, Quantization and Coordinate Descent."** In International Conference on Artificial Intelligence and Statistics, pp. 680-690. 2020.

Unified Theory of SGD

The Problem:

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

Unified Theory of SGD

The Problem:

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \begin{cases} \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \\ \frac{1}{n} \sum_{i=1}^n f_i(x) \end{cases}$$

Unified Theory of SGD

The Problem:

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

$$f(x) = \begin{cases} \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \\ \frac{1}{n} \sum_{i=1}^n f_i(x) \end{cases}$$

The Method:

$$x^{k+1} = x^k - \gamma g^k$$

Unified Theory of SGD

The Problem:

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \begin{cases} \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \\ \frac{1}{n} \sum_{i=1}^n f_i(x) \end{cases}$$

The Method:

$$x^{k+1} = x^k - \gamma g^k$$

The Assumption:

$$\mathbb{E} [g^k \mid x^k] = \nabla f(x^k)$$

Unified Theory of SGD

The Problem:

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \begin{cases} \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \\ \frac{1}{n} \sum_{i=1}^n f_i(x) \end{cases}$$

The Method:

$$x^{k+1} = x^k - \gamma g^k$$

The Assumption:

$$\mathbb{E} [g^k \mid x^k] = \nabla f(x^k)$$

$$\mathbb{E} [\|g^k\|^2 \mid x^k] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$

Reflects smoothness properties of the problem and noises introduced by stochastic gradients

Unified Theory of SGD

The Problem:

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \begin{cases} \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \\ \frac{1}{n} \sum_{i=1}^n f_i(x) \end{cases}$$

The Method:

$$x^{k+1} = x^k - \gamma g^k$$

The Assumption:

$$\mathbb{E} [g^k \mid x^k] = \nabla f(x^k)$$

$$\mathbb{E} [\|g^k\|^2 \mid x^k] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$

$$\mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

Reflects smoothness properties of the problem and noises introduced by stochastic gradients

Describes the process of variance reduction

Method	A	B	ρ	C	D_1	D_2
SGD	$2L$	0	1	0	$2\sigma^2$	0
SGD-SR	$2\mathcal{L}$	0	1	0	$2\sigma^2$	0
SGD-MB	$\frac{A'+L(\tau-1)}{\tau}$	0	1	0	$\frac{D'}{\tau}$	0
SGD-star	$2\mathcal{L}$	0	1	0	0	0
SAGA	$2L$	2	$1/n$	L/n	0	0
N-SAGA	$2L$	2	$1/n$	L/n	$2\sigma^2$	$\frac{\sigma^2}{n}$
SEGA	$2dL$	$2d$	$1/d$	L/d	0	0
N-SEGA	$2dL$	$2d$	$1/d$	L/d	$2d\sigma^2$	$\frac{\sigma^2}{d}$
SVRG ^a	$2L$	2	0	0	0	0
L-SVRG	$2L$	2	p	Lp	0	0
DIANA	$(1 + \frac{2\omega}{n}) L$	$\frac{2\omega}{n}$	α	$L\alpha$	$\frac{(1+\omega)\sigma^2}{n}$	$\alpha\sigma^2$
DIANA ^b	$(1 + 2\omega) L$	2ω	α	$L\alpha$	0	0
Q-SGD-SR	$2(1 + \omega)\mathcal{L}$	0	1	0	$2(1 + \omega)\sigma^2$	0
VR-DIANA	$(1 + \frac{4\omega+2}{n}) L$	$\frac{2(\omega+1)}{n}$	α	$(\frac{1}{m} + 4\alpha) L$	0	0
JacSketch	$2\mathcal{L}_1$	$\frac{2\lambda_{\max}}{n}$	$\lambda_{\min}$	$\frac{\mathcal{L}_2}{n}$	0	0

Main Theorem

If the stepsize satisfies

$$0 < \gamma \leq \min \left\{ \frac{1}{\mu}, \frac{1}{A + CM} \right\}, \quad \text{where} \quad M > \frac{B}{\rho}$$

Main Theorem

If the stepsize satisfies

$$0 < \gamma \leq \min \left\{ \frac{1}{\mu}, \frac{1}{A + CM} \right\}, \quad \text{where} \quad M > \frac{B}{\rho}$$

then the iterates of SGD satisfy

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho \right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \left\{ \gamma\mu, \rho - \frac{B}{M} \right\}}$$

$$\text{where} \quad V^k \stackrel{\text{def}}{=} \|x^k - x^*\|^2 + M\gamma^2\sigma_k^2$$

Main Theorem

If the stepsize satisfies

$$0 < \gamma \leq \min \left\{ \frac{1}{\mu}, \frac{1}{A + CM} \right\}, \quad \text{where} \quad M > \frac{B}{\rho}$$

then the iterates of SGD satisfy

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho \right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \left\{ \gamma\mu, \rho - \frac{B}{M} \right\}}$$

$$\text{where} \quad V^k \stackrel{\text{def}}{=} \|x^k - x^*\|^2 + M\gamma^2\sigma_k^2$$

$$\|x^k - x^*\|^2 \leq V^k$$

Table 1: List of specific existing (in some cases generalized) and new methods which fit our general analysis framework. VR = variance reduced method, AS = arbitrary sampling, Quant = supports gradient quantization, RCD = randomized coordinate descent type method. ^a Special case of SVRG with 1 outer loop only; ^b Special case of DIANA with 1 node and quantization of exact gradient.

Problem	Method	Alg #	Citation	VR?	AS?	Quant?	RCD?	Section	Result
(1)+(2)	SGD	Alg 1	Nguyen et al. (2018)	✗	✗	✗	✗	A.1	Cor A.1
(1)+(3)	SGD-SR	Alg 2	Gower et al. (2019)	✗	✓	✗	✗	A.2	Cor A.2
(1)+(3)	SGD-MB	Alg 3	NEW	✗	✗	✗	✗	A.3	Cor A.3
(1)+(3)	SGD-star	Alg 4	NEW	✓	✓	✗	✗	A.4	Cor A.4
(1)+(3)	SAGA	Alg 5	Defazio et al. (2014)	✓	✗	✗	✗	A.5	Cor A.5
(1)+(3)	N-SAGA	Alg 6	NEW	✗	✗	✗	✗	A.6	Cor A.6
(1)	SEGA	Alg 7	Hanzely et al. (2018)	✓	✗	✗	✓	A.7	Cor A.7
(1)	N-SEGA	Alg 8	NEW	✗	✗	✗	✓	A.8	Cor A.8
(1)+(3)	SVRG ^a	Alg 9	Johnson and Zhang (2013)	✓	✗	✗	✗	A.9	Cor A.9
(1)+(3)	L-SVRG	Alg 10	Hofmann et al. (2015)	✓	✗	✗	✗	A.10	Cor A.10
(1)+(3)	DIANA	Alg 11	Mishchenko et al. (2019a)	✗	✗	✓	✗	A.11	Cor A.11
(1)+(3)	DIANA ^b	Alg 12	Mishchenko et al. (2019a)	✓	✗	✓	✗	A.11	Cor A.12
(1)+(3)	Q-SGD-SR	Alg 13	NEW	✗	✓	✓	✗	A.12	Cor A.13
(1)+(3)+(4)	VR-DIANA	Alg 14	Horváth et al. (2019)	✓	✗	✓	✗	A.13	Cor A.15
(1)+(3)	JacSketch	Alg 15	Gower et al. (2018)	✓	✓✗	✗	✗	A.14	Cor A.16

6.2. Unified Convergence Analysis of Methods with Error Compensation

Key Assumption

$$g^k = \frac{1}{n} \sum_{i=1}^n g_i^k, \quad \mathbb{E} [g^k \mid x^k] = \nabla f (x^k) \quad \bar{g}_i^k = \mathbb{E} [g_i^k \mid x^k]$$

$$\frac{1}{n} \sum_{i=1}^n \|\bar{g}_i^k\|^2 \leq 2A (f (x^k) - f (x^*)) + B_1 \sigma_{1,k}^2 + B_2 \sigma_{2,k}^2 + D_1$$

$$\frac{1}{n} \sum_{i=1}^n \mathbb{E} [\|g_i^k - \bar{g}_i^k\|^2 \mid x^k] \leq 2\tilde{A} (f (x^k) - f (x^*)) + \tilde{B}_1 \sigma_{1,k}^2 + \tilde{B}_2 \sigma_{2,k}^2 + \tilde{D}_1$$

$$\mathbb{E} [\|g^k\|^2 \mid x^k] \leq 2A' (f (x^k) - f (x^*)) + B'_1 \sigma_{1,k}^2 + B'_2 \sigma_{2,k}^2 + D'_1$$

$$\mathbb{E} [\sigma_{1,k+1}^2 \mid \sigma_{1,k}^2, \sigma_{2,k}^2] \leq (1 - \rho_1) \sigma_{1,k}^2 + 2C_1 (f (x^k) - f (x^*)) + G\rho_1 \sigma_{2,k}^2 + D_2$$

$$\mathbb{E} [\sigma_{2,k+1}^2 \mid \sigma_{2,k}^2] \leq (1 - \rho_2) \sigma_{2,k}^2 + 2C_2 (f (x^k) - f (x^*))$$

Key Assumption

$$g^k = \frac{1}{n} \sum_{i=1}^n g_i^k, \quad \mathbb{E} [g^k \mid x^k] = \nabla f(x^k) \quad \bar{g}_i^k = \mathbb{E} [g_i^k \mid x^k]$$

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n \|\bar{g}_i^k\|^2 &\leq 2A (f(x^k) - f(x^*)) + B_1 \sigma_{1,k}^2 + B_2 \sigma_{2,k}^2 + D_1 \\ \frac{1}{n} \sum_{i=1}^n \mathbb{E} [\|g_i^k - \bar{g}_i^k\|^2 \mid x^k] &\leq 2\tilde{A} (f(x^k) - f(x^*)) + \tilde{B}_1 \sigma_{1,k}^2 + \tilde{B}_2 \sigma_{2,k}^2 + \tilde{D}_1 \\ \mathbb{E} [\|g^k\|^2 \mid x^k] &\leq 2A' (f(x^k) - f(x^*)) + B'_1 \sigma_{1,k}^2 + B'_2 \sigma_{2,k}^2 + D'_1 \end{aligned}$$

$$\mathbb{E} [\sigma_{1,k+1}^2 \mid \sigma_{1,k}^2, \sigma_{2,k}^2] \leq (1 - \rho_1) \sigma_{1,k}^2 + 2C_1 (f(x^k) - f(x^*)) + G\rho_1 \sigma_{2,k}^2 + D_2$$

$$\mathbb{E} [\sigma_{2,k+1}^2 \mid \sigma_{2,k}^2] \leq (1 - \rho_2) \sigma_{2,k}^2 + 2C_2 (f(x^k) - f(x^*))$$

 Reflects smoothness properties of the problem and noises introduced by compressions and stochastic gradients

Key Assumption

$$g^k = \frac{1}{n} \sum_{i=1}^n g_i^k, \quad \mathbb{E} [g^k \mid x^k] = \nabla f(x^k) \quad \bar{g}_i^k = \mathbb{E} [g_i^k \mid x^k]$$

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n \|\bar{g}_i^k\|^2 &\leq 2A (f(x^k) - f(x^*)) + B_1 \sigma_{1,k}^2 + B_2 \sigma_{2,k}^2 + D_1 \\ \frac{1}{n} \sum_{i=1}^n \mathbb{E} [\|g_i^k - \bar{g}_i^k\|^2 \mid x^k] &\leq 2\tilde{A} (f(x^k) - f(x^*)) + \tilde{B}_1 \sigma_{1,k}^2 + \tilde{B}_2 \sigma_{2,k}^2 + \tilde{D}_1 \\ \mathbb{E} [\|g^k\|^2 \mid x^k] &\leq 2A' (f(x^k) - f(x^*)) + B'_1 \sigma_{1,k}^2 + B'_2 \sigma_{2,k}^2 + D'_1 \end{aligned}$$

$$\mathbb{E} [\sigma_{1,k+1}^2 \mid \sigma_{1,k}^2, \sigma_{2,k}^2] \leq (1 - \rho_1) \sigma_{1,k}^2 + 2C_1 (f(x^k) - f(x^*)) + G\rho_1 \sigma_{2,k}^2 + D_2$$

$$\mathbb{E} [\sigma_{2,k+1}^2 \mid \sigma_{2,k}^2] \leq (1 - \rho_2) \sigma_{2,k}^2 + 2C_2 (f(x^k) - f(x^*))$$

Reflects smoothness properties of the problem and noises introduced by compressions and stochastic gradients

Describes the process of variance reduction of the variance coming from compressions

Key Assumption

$$g^k = \frac{1}{n} \sum_{i=1}^n g_i^k, \quad \mathbb{E} [g^k \mid x^k] = \nabla f(x^k) \quad \bar{g}_i^k = \mathbb{E} [g_i^k \mid x^k]$$

$$\frac{1}{n} \sum_{i=1}^n \|\bar{g}_i^k\|^2 \leq 2A (f(x^k) - f(x^*)) + B_1 \sigma_{1,k}^2 + B_2 \sigma_{2,k}^2 + D_1$$

$$\frac{1}{n} \sum_{i=1}^n \mathbb{E} [\|g_i^k - \bar{g}_i^k\|^2 \mid x^k] \leq 2\tilde{A} (f(x^k) - f(x^*)) + \tilde{B}_1 \sigma_{1,k}^2 + \tilde{B}_2 \sigma_{2,k}^2 + \tilde{D}_1$$

$$\mathbb{E} [\|g^k\|^2 \mid x^k] \leq 2A' (f(x^k) - f(x^*)) + B'_1 \sigma_{1,k}^2 + B'_2 \sigma_{2,k}^2 + D'_1$$

$$\mathbb{E} [\sigma_{1,k+1}^2 \mid \sigma_{1,k}^2, \sigma_{2,k}^2] \leq (1 - \rho_1) \sigma_{1,k}^2 + 2C_1 (f(x^k) - f(x^*)) + G\rho_1 \sigma_{2,k}^2 + D_2$$

$$\mathbb{E} [\sigma_{2,k+1}^2 \mid \sigma_{2,k}^2] \leq (1 - \rho_2) \sigma_{2,k}^2 + 2C_2 (f(x^k) - f(x^*))$$

Reflects smoothness properties of the problem and noises introduced by compressions and stochastic gradients

Describes the process of variance reduction of the variance coming from compressions

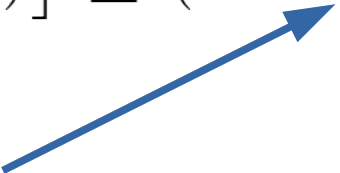
Describes the process of variance reduction of the variance coming from stochastic gradients

Main Theorem

$$\mathbb{E} \left[f \left(\bar{x}^K \right) - f \left(x^* \right) \right] \leq (1 - \eta)^K \frac{\Psi \left(x^0, \gamma \right)}{\gamma} + \gamma \Phi \left(D_1, \tilde{D}_1, D'_1, D_2 \right)$$

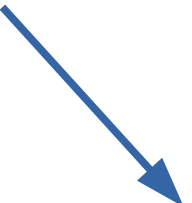
Main Theorem

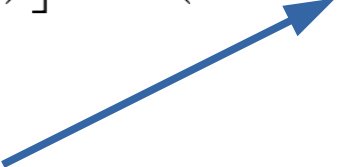
$$\mathbb{E} \left[f \left(\bar{x}^K \right) - f \left(x^* \right) \right] \leq (1 - \eta)^K \frac{\Psi \left(x^0, \gamma \right)}{\gamma} + \gamma \Phi \left(D_1, \tilde{D}_1, D'_1, D_2 \right)$$

$$\eta = \min \left\{ \frac{\gamma \mu}{2}, \frac{\rho_1}{4}, \frac{\rho_2}{4} \right\}$$


Main Theorem

Some quantity depending only on the starting point and stepsize


$$\mathbb{E} \left[f \left(\bar{x}^K \right) - f \left(x^* \right) \right] \leq (1 - \eta)^K \frac{\Psi \left(x^0, \gamma \right)}{\gamma} + \gamma \Phi \left(D_1, \tilde{D}_1, D'_1, D_2 \right)$$


$$\eta = \min \left\{ \frac{\gamma \mu}{2}, \frac{\rho_1}{4}, \frac{\rho_2}{4} \right\}$$

Main Theorem

Some quantity depending only on the starting point and stepsize

$$\mathbb{E} \left[f \left(\bar{x}^K \right) - f \left(x^* \right) \right] \leq (1 - \eta)^K \frac{\Psi \left(x^0, \gamma \right)}{\gamma} + \gamma \Phi \left(D_1, \tilde{D}_1, D'_1, D_2 \right)$$

Linear function

$$\eta = \min \left\{ \frac{\gamma \mu}{2}, \frac{\rho_1}{4}, \frac{\rho_2}{4} \right\}$$

Methods with Error Compensation Covered by Our Framework

Problem	Method	Alg #	Citation	Sec #	Rate (constants ignored)
(1)+(3)	EC-SGDsr	Alg 3	new	H.1	$\tilde{\mathcal{O}} \left(\frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-SGD	Alg 4	[45]	H.2	$\tilde{\mathcal{O}} \left(\frac{\kappa}{\delta} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-GDstar	Alg 5	new	H.3	$\mathcal{O} \left(\frac{\kappa}{\delta} \log \frac{1}{\varepsilon} \right)$
(1)+(2)	EC-SGD-DIANA	Alg 6	new	H.4	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(3)	EC-SGDsr-DIANA	Alg 7	new	H.5	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-GD-DIANA [†]	Alg 6	new	H.4	$\mathcal{O} \left(\left(\omega + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG	Alg 8	new	H.6	$\tilde{\mathcal{O}} \left(m + \frac{\kappa}{\delta} + \frac{\sqrt{L\zeta_*^2}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-LSVRGstar	Alg 9	new	H.7	$\mathcal{O} \left(\left(m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG-DIANA	Alg 10	new	H.8	$\mathcal{O} \left(\left(\omega + m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$

Methods with Error Compensation Covered by Our Framework

Problem	Method	Alg #	Citation	Sec #	Rate (constants ignored)
(1)+(3)	EC-SGDsr	Alg 3	new	H.1	$\tilde{\mathcal{O}} \left(\frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-SGD	Alg 4	[45]	H.2	$\tilde{\mathcal{O}} \left(\frac{\kappa}{\delta} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-GDstar	Alg 5	new	H.3	$\mathcal{O} \left(\frac{\kappa}{\delta} \log \frac{1}{\varepsilon} \right)$
(1)+(2)	EC-SGD-DIANA	Alg 6	new	H.4	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(3)	EC-SGDsr-DIANA	Alg 7	new	H.5	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-GD-DIANA [†]	Alg 6	new	H.4	$\mathcal{O} \left(\left(\omega + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG	Alg 8	new	H.6	$\tilde{\mathcal{O}} \left(m + \frac{\kappa}{\delta} + \frac{\sqrt{L\zeta_*^2}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-LSVRGstar	Alg 9	new	H.7	$\mathcal{O} \left(\left(m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG-DIANA	Alg 10	new	H.8	$\mathcal{O} \left(\left(\omega + m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$

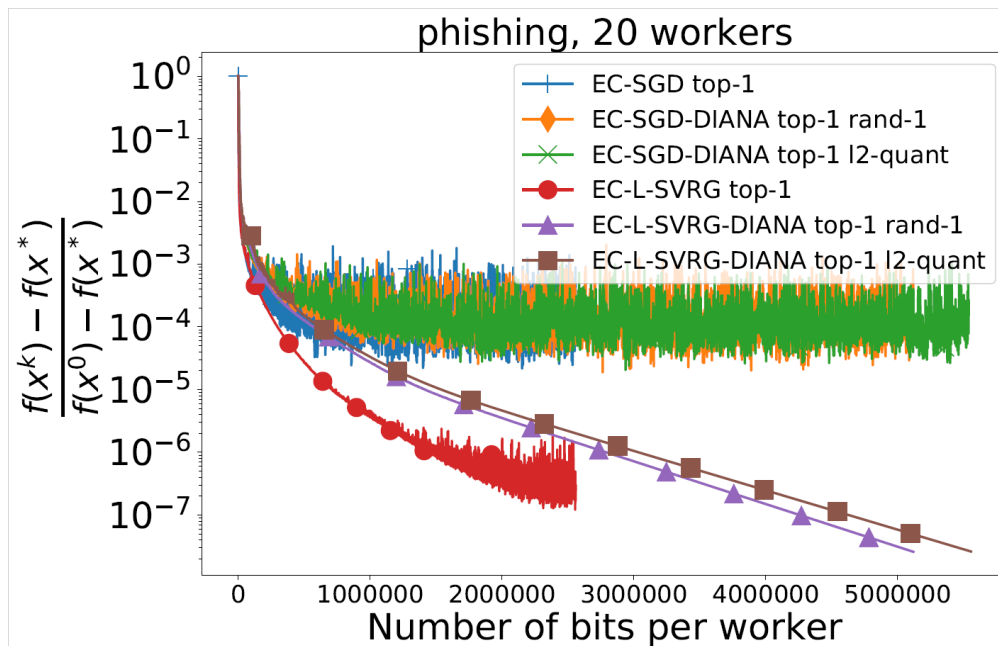
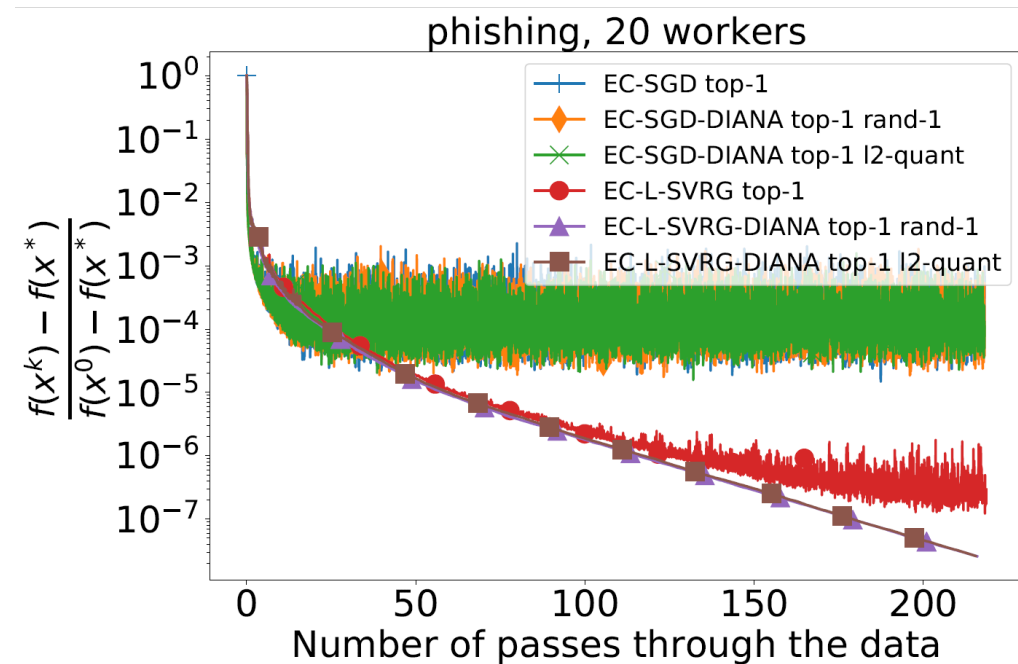
Methods with Error Compensation Covered by Our Framework

Problem	Method	Alg #	Citation	Sec #	Rate (constants ignored)
(1)+(3)	EC-SGDsr	Alg 3	new	H.1	$\tilde{\mathcal{O}} \left(\frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-SGD	Alg 4	[45]	H.2	$\tilde{\mathcal{O}} \left(\frac{\kappa}{\delta} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L(\text{Var}+\zeta_*^2/\delta)}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-GDstar	Alg 5	new	H.3	$\mathcal{O} \left(\frac{\kappa}{\delta} \log \frac{1}{\varepsilon} \right)$
(1)+(2)	EC-SGD-DIANA	Alg 6	new	H.4	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\kappa}{\delta} + \frac{\sigma^2}{n\mu\varepsilon} + \frac{\sqrt{L\sigma^2}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(3)	EC-SGDsr-DIANA	Alg 7	new	H.5	Option I: $\tilde{\mathcal{O}} \left(\omega + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\delta\mu\sqrt{\varepsilon}} \right)$ Option II: $\tilde{\mathcal{O}} \left(\frac{1+\omega}{\delta} + \frac{\mathcal{L}}{\mu} + \frac{\sqrt{L\mathcal{L}}}{\delta\mu} + \frac{\text{Var}}{n\mu\varepsilon} + \frac{\sqrt{L\text{Var}}}{\mu\sqrt{\delta\varepsilon}} \right)$
(1)+(2)	EC-GD-DIANA [†]	Alg 6	new	H.4	$\mathcal{O} \left(\left(\omega + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG	Alg 8	new	H.6	$\tilde{\mathcal{O}} \left(m + \frac{\kappa}{\delta} + \frac{\sqrt{L\zeta_*^2}}{\delta\mu\sqrt{\varepsilon}} \right)$
(1)+(3)	EC-LSVRGstar	Alg 9	new	H.7	$\mathcal{O} \left(\left(m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$
(1)+(3)	EC-LSVRG-DIANA	Alg 10	new	H.8	$\mathcal{O} \left(\left(\omega + m + \frac{\kappa}{\delta} \right) \log \frac{1}{\varepsilon} \right)$

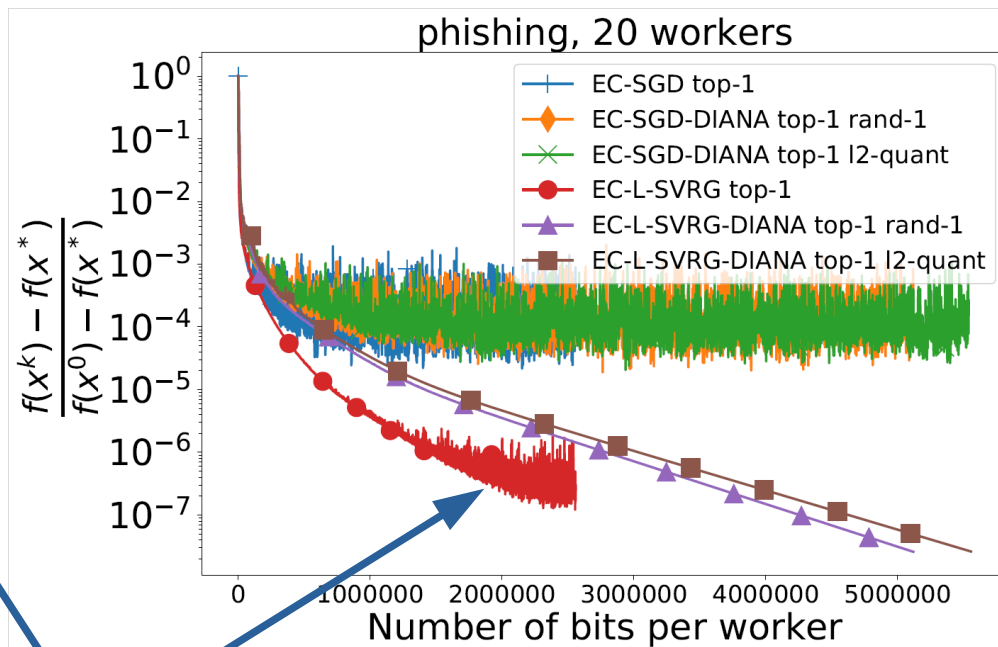
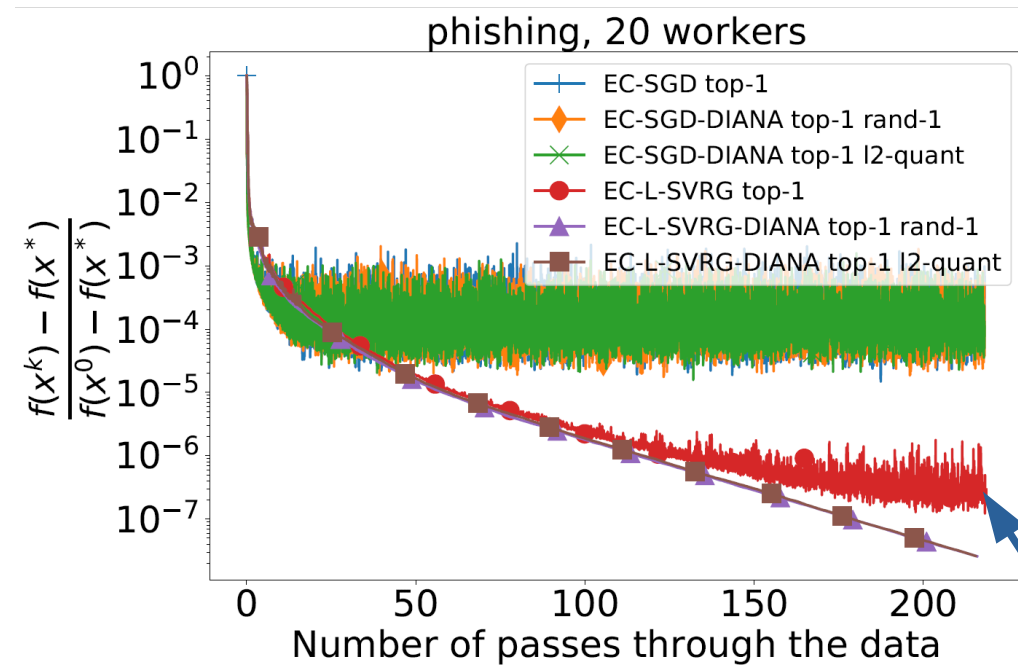
Our framework covers even methods without error compensation and methods with delayed updates

7. Experiments

Logistic Regression with l2-regularization

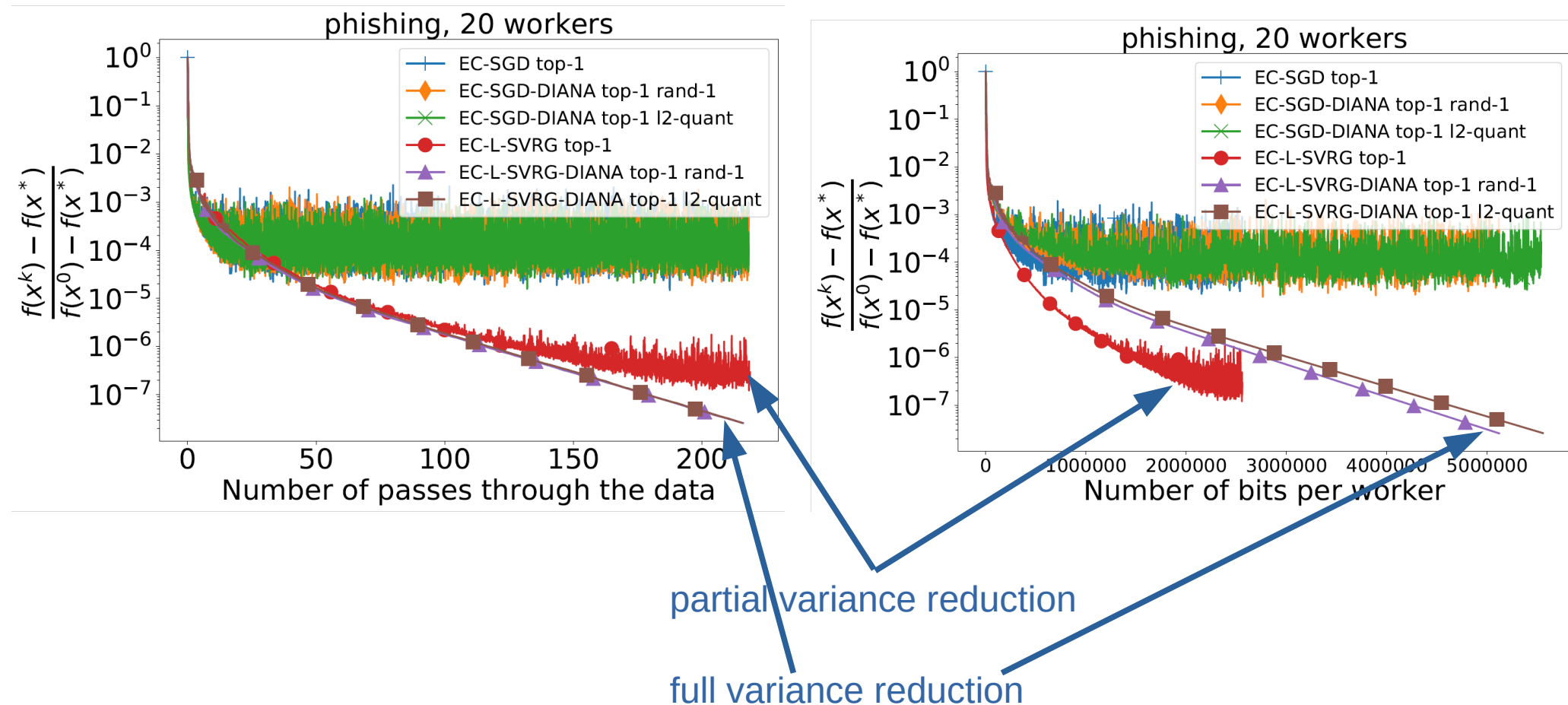


Logistic Regression with l2-regularization



partial variance reduction

Logistic Regression with l2-regularization



8. More Methods

More Methods Fitting our Framework

The generality of our approach helps to obtain convergence guarantees for a big number of different stochastic methods (even without error compensation). Here are some examples.

More Methods Fitting our Framework

The generality of our approach helps to obtain convergence guarantees for a big number of different stochastic methods (even without error compensation). Here are some examples.

- Methods without error feedback: SGD, SGD-SR (arbitrary sampling), SAGA, SVRG, L-SVRG, QSGD, TernGrad, DQGD, DIANA, **DIANA_{sr-DQ}**, VR-DIANA, JacSketch, SEGA

More Methods Fitting our Framework

The generality of our approach helps to obtain convergence guarantees for a big number of different stochastic methods (even without error compensation). Here are some examples.

- Methods without error feedback: SGD, SGD-SR (arbitrary sampling), SAGA, SVRG, L-SVRG, QSGD, TernGrad, DQGD, DIANA, **DIANAsr-DQ**, VR-DIANA, JacSketch, SEGA
- Methods with delayed updates: D-SGD, **D-SGD-SR** (arbitrary sampling), **D-QSGD**, **D-SGD-DIANA**, **D-LSVRG**, **D-QLSVRG**, **D-LSVRG-DIANA**

More Methods Fitting our Framework

The generality of our approach helps to obtain convergence guarantees for a big number of different stochastic methods (even without error compensation). Here are some examples.

- Methods without error feedback: SGD, SGD-SR (arbitrary sampling), SAGA, SVRG, L-SVRG, QSGD, TernGrad, DQGD, DIANA, **DIANAsr-DQ**, VR-DIANA, JacSketch, SEGA
- Methods with delayed updates: D-SGD, **D-SGD-SR** (arbitrary sampling), **D-QSGD**, **D-SGD-DIANA**, **D-LSVRG**, **D-QLSVRG**, **D-LSVRG-DIANA**
- ✓ In one theorem, we recover the sharpest rates for all known special cases
- ✓ Our analysis works for non-strongly convex objectives as well

More Methods Fitting our Framework

The generality of our approach helps to obtain convergence guarantees for a big number of different stochastic methods (even without error compensation). Here are some examples.

- Methods without error feedback: SGD, SGD-SR (arbitrary sampling), SAGA, SVRG, L-SVRG, QSGD, TernGrad, DQGD, DIANA, **DIANAsr-DQ**, VR-DIANA, JacSketch, SEGA
- Methods with delayed updates: D-SGD, **D-SGD-SR** (arbitrary sampling), **D-QSGD**, **D-SGD-DIANA**, **D-LSVRG**, **D-QLSVRG**, **D-LSVRG-DIANA**
- ✓ In one theorem, we recover the sharpest rates for all known special cases
- ✓ Our analysis works for non-strongly convex objectives as well

Thank you for your attention!

9. Extra Examples on Unified SGD



Gorbunov, Eduard, Filip Hanzely, and Peter Richtárik. **"A Unified Theory of SGD: Variance reduction, Sampling, Quantization and Coordinate Descent."** In International Conference on Artificial Intelligence and Statistics, pp. 680-690. 2020.

9.1. SGD-SR

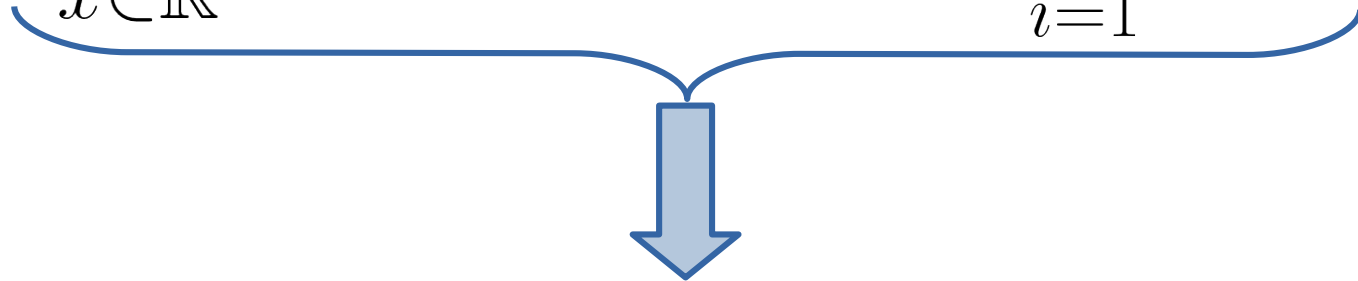


Gower, R. M., Loizou, N., Qian, X., Sailanbayev, A., Shulgin, E., & Richtárik, P. (2019, May).
SGD: General Analysis and Improved Rates. *In International Conference on Machine Learning* (pp. 5200-5209).

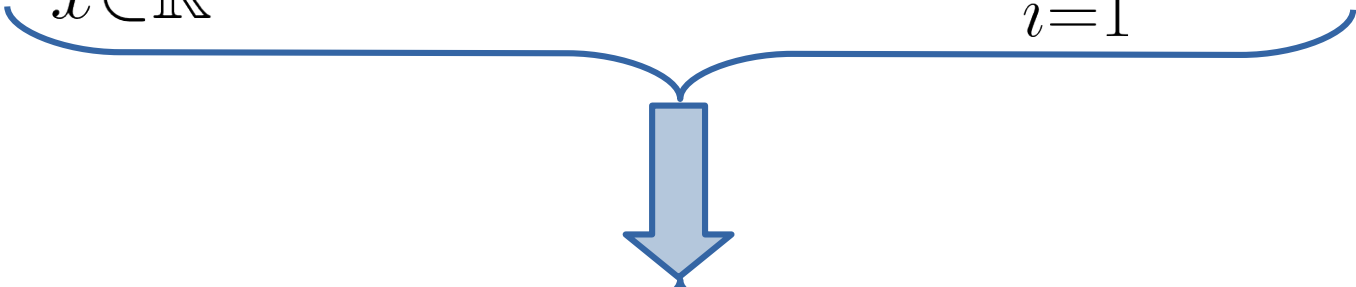
Stochastic Reformulation

$$\min_{x \in \mathbb{R}^d} f(x) \qquad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$

Stochastic Reformulation

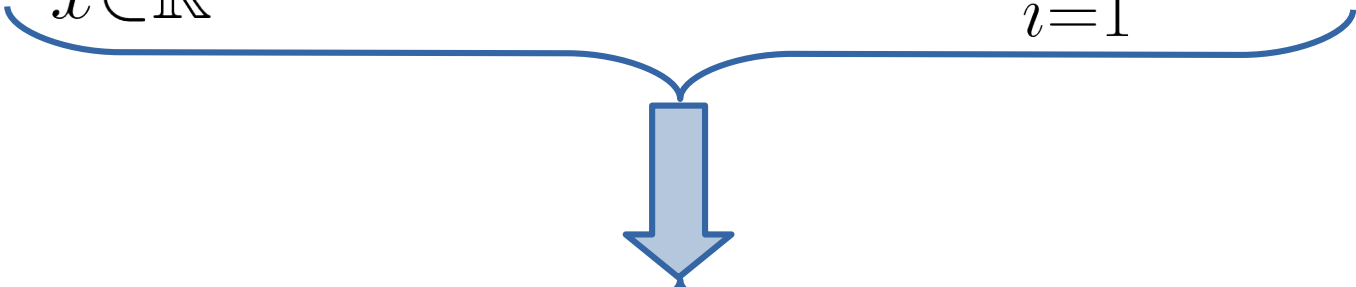
$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$


Stochastic Reformulation

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$


$$f(x) = \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \quad f_{\xi}(x) \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n \xi_i f_i(x)$$

Stochastic Reformulation

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$


$$f(x) = \mathbb{E}_{\xi \sim \mathcal{D}} [f_{\xi}(x)] \quad f_{\xi}(x) \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n \xi_i f_i(x)$$

- $\xi \sim \mathcal{D} : \mathbb{E}_{\xi \sim \mathcal{D}} [\xi_i] = 1$
- $f_i(x)$ is smooth, but can be non-convex

- 1 **for** $k=0,1,2, \dots$ **do**
- 2 Sample $\xi \sim \mathcal{D}$

- 1 **for** $k=0,1,2, \dots$ **do**
- 2 Sample $\xi \sim \mathcal{D}$
- 3 $g^k = \nabla f_{\xi} \left(x^k \right)$

1 **for** $k=0,1,2, \dots$ **do**

2 Sample $\xi \sim \mathcal{D}$

3 $g^k = \nabla f_{\xi} \left(x^k \right)$

4 $x^{k+1} = x^k - \gamma g^k$

5 **end for**

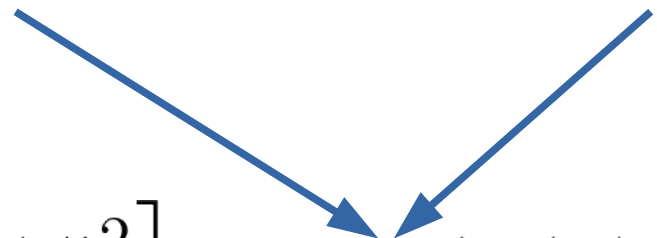
Expected Smoothness

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x) - \nabla f_{\xi}(x^*)\|^2 \right] \leq 2\mathcal{L} (f(x) - f(x^*))$$

Expected Smoothness

Can be significantly smaller than the
maximal smoothness constant of $f_i(x)$

Some positive constant
depending on $\mathcal{D}$


$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x) - \nabla f_{\xi}(x^*)\|^2 \right] \leq 2\mathcal{L} (f(x) - f(x^*))$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$
$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$
$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

Key assumption:

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$
$$\mathbb{E} \left[\sigma_{k+1}^2 \mid \sigma_k^2 \right] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$
$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

That is, SGD-SR satisfies the key assumption with parameters:

Key assumption:

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$
$$\mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$

$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

That is, SGD-SR satisfies the key assumption with parameters:

$$A = 2\mathcal{L}, B = 0, D_1 = 2\sigma_*^2, \sigma_k = 0, \rho = 1, C = 0, D_2 = 0$$

Key assumption:

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$

$$\mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

Main Theorem: Reminder

If the stepsize satisfies

$$0 < \gamma \leq \min \left\{ \frac{1}{\mu}, \frac{1}{A + CM} \right\}, \quad \text{where} \quad M > \frac{B}{\rho}$$

then the iterates of SGD satisfy

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho \right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \left\{ \gamma\mu, \rho - \frac{B}{M} \right\}}$$

where $V^k \stackrel{\text{def}}{=} \|x^k - x^*\|^2 + M\gamma^2\sigma_k^2$

$$\|x^k - x^*\|^2 \leq V^k$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$

$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

That is, SGD-SR satisfies the key assumption with parameters:

$$A = 2\mathcal{L}, B = 0, D_1 = 2\sigma_*^2, \sigma_k = 0, \rho = 1, C = 0, D_2 = 0$$

Main theorem:

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho \right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \{ \gamma\mu, \rho - \frac{B}{M} \}}$$

Key assumption:

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$

$$\mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

Main Result for SGD-SR

Young inequality and expected smoothness imply:

$$\mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x)\|^2 \right] \leq 4\mathcal{L} (f(x) - f(x^*)) + 2\sigma_*^2$$

$$\sigma_*^2 \stackrel{\text{def}}{=} \mathbb{E}_{\xi \sim \mathcal{D}} \left[\|\nabla f_{\xi}(x^*)\|^2 \right]$$

That is, SGD-SR satisfies the key assumption with parameters:

$$A = 2\mathcal{L}, B = 0, D_1 = 2\sigma_*^2, \sigma_k = 0, \rho = 1, C = 0, D_2 = 0$$

$$\mathbb{E} \|x^k - x^*\|^2 \leq (1 - \gamma\mu)^k \|x^0 - x^*\|^2 + \frac{2\gamma\sigma_*^2}{\mu}$$

$$\gamma \leq \frac{1}{2\mathcal{L}}$$

Main theorem:

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho \right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \{ \gamma\mu, \rho - \frac{B}{M} \}}$$

Key assumption:

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1$$

$$\mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] \leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2$$

9.2. SAGA



Defazio, Aaron, Francis Bach, and Simon Lacoste-Julien. **SAGA: A fast incremental gradient method with support for non-strongly convex composite objectives.** *In Advances in neural information processing systems*, pp. 1646-1654. 2014.

The Problem

$$\min_{x \in \mathbb{R}^d} f(x) \qquad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$

The Problem

$$\min_{x \in \mathbb{R}^d} f(x) \quad f(x) = \frac{1}{n} \sum_{i=1}^n f_i(x)$$


$$\begin{aligned} \|\nabla f_i(x) - \nabla f_i(y)\| &\leq L\|x - y\| \\ f_i(y) &\geq f_i(x) + \langle \nabla f_i(x), y - x \rangle \end{aligned}$$

SAGA

1 Set $\phi_i^0 = x^0 \quad \forall i \in [n]$

SAGA

- 1 Set $\phi_i^0 = x^0 \quad \forall i \in [n]$
- 2 **for** $k=0,1,2, \dots$ **do**
- 3 Sample $j \in [n]$ uniformly at random

SAGA

- 1 Set $\phi_i^0 = x^0 \quad \forall i \in [n]$
- 2 **for** $k=0,1,2, \dots$ **do**
- 3 Sample $j \in [n]$ uniformly at random
- 4 $g^k = \nabla f_j(x^k) - \nabla f_j(\phi_j^k) + \frac{1}{n} \sum_{i=1}^n \nabla f_i(\phi_i^k)$

SAGA

- 1 Set $\phi_i^0 = x^0 \quad \forall i \in [n]$
- 2 **for** $k=0,1,2, \dots$ **do**
- 3 Sample $j \in [n]$ uniformly at random
- 4 $g^k = \nabla f_j(x^k) - \nabla f_j(\phi_j^k) + \frac{1}{n} \sum_{i=1}^n \nabla f_i(\phi_i^k)$
- 5 Set $\phi_j^{k+1} = x^k$ and $\phi_i^{k+1} = \phi_i^k$ for $i \neq j$

SAGA

- 1 Set $\phi_i^0 = x^0 \quad \forall i \in [n]$
- 2 **for** $k=0,1,2, \dots$ **do**
- 3 Sample $j \in [n]$ uniformly at random
- 4 $g^k = \nabla f_j(x^k) - \nabla f_j(\phi_j^k) + \frac{1}{n} \sum_{i=1}^n \nabla f_i(\phi_i^k)$
- 5 Set $\phi_j^{k+1} = x^k$ and $\phi_i^{k+1} = \phi_i^k$ for $i \neq j$
- 6 $x^{k+1} = x^k - \gamma g^k$
- 7 **end for**

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L (f(x^k) - f(x^*)) + 2\sigma_k^2$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L \left(f(x^k) - f(x^*) \right) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L \left(f(x^k) - f(x^*) \right) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n} \right) \sigma_k^2 + \frac{2L}{n} \left(f(x^k) - f(x^*) \right)$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L \left(f(x^k) - f(x^*) \right) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n}\right) \sigma_k^2 + \frac{2L}{n} \left(f(x^k) - f(x^*) \right)$$

That is, SAGA satisfies the key assumption with parameters:

Key assumption:

$$\begin{aligned} \mathbb{E} \left[\|g^k\|^2 \mid x^k \right] &\leq 2A \left(f(x^k) - f(x^*) \right) + B\sigma_k^2 + D_1 \\ \mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] &\leq (1 - \rho)\sigma_k^2 + 2C \left(f(x^k) - f(x^*) \right) + D_2 \end{aligned}$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L \left(f(x^k) - f(x^*) \right) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n} \right) \sigma_k^2 + \frac{2L}{n} \left(f(x^k) - f(x^*) \right)$$

That is, SAGA satisfies the key assumption with parameters:

$$A = 2L, B = 2, D_1 = 0, \rho = \frac{1}{n}, C = \frac{L}{n}, D_2 = 0$$

Key assumption:

$$\begin{aligned} \mathbb{E} \left[\|g^k\|^2 \mid x^k \right] &\leq 2A \left(f(x^k) - f(x^*) \right) + B\sigma_k^2 + D_1 \\ \mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] &\leq (1 - \rho)\sigma_k^2 + 2C \left(f(x^k) - f(x^*) \right) + D_2 \end{aligned}$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L (f(x^k) - f(x^*)) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n}\right) \sigma_k^2 + \frac{2L}{n} (f(x^k) - f(x^*))$$

That is, SAGA satisfies the key assumption with parameters:

$$A = 2L, B = 2, D_1 = 0, \rho = \frac{1}{n}, C = \frac{L}{n}, D_2 = 0$$

$$\mathbb{E} V^k \leq \left(1 - \min \left\{ \frac{\mu}{6L}, \frac{1}{2n} \right\} \right)^k V^0$$

Main theorem:

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho\right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \{ \gamma\mu, \rho - \frac{B}{M} \}}$$

Key assumption:

$$\begin{aligned} \mathbb{E} \left[\|g^k\|^2 \mid x^k \right] &\leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1 \\ \mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] &\leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2 \end{aligned}$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L (f(x^k) - f(x^*)) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n}\right) \sigma_k^2 + \frac{2L}{n} (f(x^k) - f(x^*))$$

That is, SAGA satisfies the key assumption with parameters:

$$A = 2L, B = 2, D_1 = 0, \rho = \frac{1}{n}, C = \frac{L}{n}, D_2 = 0$$

$$\mathbb{E} V^k \leq \left(1 - \min \left\{ \frac{\mu}{6L}, \frac{1}{2n} \right\} \right)^k V^0$$

Main theorem:

$$\mathbb{E} [V^k] \leq \max \left\{ (1 - \gamma\mu)^k, \left(1 + \frac{B}{M} - \rho\right)^k \right\} V^0 + \frac{(D_1 + MD_2) \gamma^2}{\min \left\{ \gamma\mu, \rho - \frac{B}{M} \right\}}$$

Key assumption:

$$\begin{aligned} \mathbb{E} \left[\|g^k\|^2 \mid x^k \right] &\leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1 \\ \mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] &\leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2 \end{aligned}$$

Main Result for SAGA

$$\mathbb{E} \left[\|g^k\|^2 \mid x^k \right] \leq 4L (f(x^k) - f(x^*)) + 2\sigma_k^2$$

$$\sigma_k^2 = \frac{1}{n} \sum_{i=1}^n \|\nabla f_i(\phi_i^k) - \nabla f_i(x^*)\|^2 \quad \mathbb{E} [\sigma_{k+1}^2 \mid x^k] \leq \left(1 - \frac{1}{n}\right) \sigma_k^2 + \frac{2L}{n} (f(x^k) - f(x^*))$$

That is, SAGA satisfies the key assumption with parameters:

$$A = 2L, B = 2, D_1 = 0, \rho = \frac{1}{n}, C = \frac{L}{n}, D_2 = 0$$

$$\mathbb{E} V^k \leq \left(1 - \min \left\{ \frac{\mu}{6L}, \frac{1}{2n} \right\}\right)^k V^0$$

$$V^k = \|x^k - x^*\|^2 + 2n\gamma^2 \sigma_k^2$$

$$\gamma = \frac{1}{6L}$$

Key assumption:

$$\begin{aligned} \mathbb{E} [\|g^k\|^2 \mid x^k] &\leq 2A (f(x^k) - f(x^*)) + B\sigma_k^2 + D_1 \\ \mathbb{E} [\sigma_{k+1}^2 \mid \sigma_k^2] &\leq (1 - \rho)\sigma_k^2 + 2C (f(x^k) - f(x^*)) + D_2 \end{aligned}$$

The End