

Asymptotic Normality and Optimality of Stochastic Linearized Augmented Lagrangian Method

Jiajin Li

Sauder School of Business

University of British Columbia



SIAM Conference on Optimization, 2026

Joint work with Feng Ruan (Northwestern).

What this talk is about.

Understand the statistical behavior (asymptotic normality) of stochastic primal–dual algorithms for constrained optimization.

Warm-up: the classical SGD

- **Setup.** Consider a strongly convex and smooth stochastic optimization problem:

$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x}) := \mathbb{E}[F(\mathbf{x}, \xi)]. \quad (\text{P0})$$

- **Algorithm.** The simplest method—stochastic gradient descent (SGD):

$$\mathbf{x}^{k+1} = \mathbf{x}^k - \alpha_k \nabla F(\mathbf{x}^k, \xi^k).$$

Here α_k is the stepsize and ξ^k denotes an i.i.d. sample.

Warm-up: the classical SGD

- **Setup.** Consider a strongly convex and smooth stochastic optimization problem:

$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x}) := \mathbb{E}[F(\mathbf{x}, \xi)]. \quad (\text{P0})$$

- **Algorithm.** The simplest method—stochastic gradient descent (SGD):

$$\boxed{\mathbf{x}^{k+1} = \mathbf{x}^k - \alpha_k \nabla F(\mathbf{x}^k, \xi^k)}.$$

Here α_k is the stepsize and ξ^k denotes an i.i.d. sample.

Theorem (Ruppert (1988); Polyak–Juditsky (1992))

If $\alpha_k \propto k^{-\theta}$ with $\theta \in (\frac{1}{2}, 1)$ and standard noise conditions hold, then

$$\sqrt{k}(\bar{\mathbf{x}}_k - \mathbf{x}^*) \xrightarrow{d} \mathcal{N}\left(0, \nabla^2 f(\mathbf{x}^*)^{-1} \Sigma \nabla^2 f(\mathbf{x}^*)^{-1}\right),$$

where $\bar{\mathbf{x}}_k := \frac{1}{k} \sum_{i=1}^k \mathbf{x}^i$, and $\Sigma := \text{Cov}(\nabla F(\mathbf{x}, \xi))$.

- Matches the covariance of SAA (sample-average approximation) $\Rightarrow$ same statistical efficiency.
- *Asymptotically optimal* in the Hájek–Le Cam sense (Hájek '70; Le Cam '71).
- Enables online covariance estimation and construction of confidence intervals for $\mathbf{x}^*$.

Theorem (Ruppert (1988); Polyak–Juditsky (1992))

If $\alpha_k \propto k^{-\theta}$ with $\theta \in (\frac{1}{2}, 1)$ and standard noise conditions hold, then

$$\sqrt{k}(\bar{\mathbf{x}}_k - \mathbf{x}^*) \xrightarrow{d} \mathcal{N}\left(0, \nabla^2 f(\mathbf{x}^*)^{-1} \Sigma \nabla^2 f(\mathbf{x}^*)^{-1}\right),$$

where $\bar{\mathbf{x}}_k := \frac{1}{k} \sum_{i=1}^k \mathbf{x}^i$, and $\Sigma := \text{Cov}(\nabla F(\mathbf{x}, \xi))$.

- Matches the covariance of SAA (sample-average approximation) $\Rightarrow$ same statistical efficiency.
- *Asymptotically optimal* in the Hájek–Le Cam sense (Hájek '70; Le Cam '71).
- Enables online covariance estimation and construction of confidence intervals for $\mathbf{x}^*$.

From unconstrained to constrained

We consider

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) := \mathbb{E}[F(\mathbf{x}, \xi)], \quad \mathcal{X} := \{\mathbf{x} \in \mathbb{R}^n : h_i(\mathbf{x}) \leq 0, i \in [m]\}. \quad (\text{P})$$

Assume $F(\cdot, \xi)$ is C^2 in $\mathbf{x}$ (a.s.) and each h_i is C^2 .

- **SAA baseline.** Sample-average approximation admits asymptotic normality and is *asymptotically optimal* (Hájek–Le Cam sense) (Dupacová–Wets '88; Shapiro '89; King–Rockafellar '93).
- Are practical online first-order methods also optimal?
 - **Dual Averaging** (Duchi–Ruan '18): optimal *only when $\mathcal{X}$ is polyhedral*.
 - **Projected SGD** (Davis–Drusvyatskiy–Jiang '23): optimal for general $\mathcal{X}$, but *lacks finite-time active-set identification* and requires stronger noise assumptions.
- Both are **primal-only** and require explicit projection onto $\mathcal{X}$.

From unconstrained to constrained

We consider

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) := \mathbb{E}[F(\mathbf{x}, \xi)], \quad \mathcal{X} := \{\mathbf{x} \in \mathbb{R}^n : h_i(\mathbf{x}) \leq 0, i \in [m]\}. \quad (\text{P})$$

Assume $F(\cdot, \xi)$ is C^2 in $\mathbf{x}$ (a.s.) and each h_i is C^2 .

- **SAA baseline.** Sample-average approximation admits asymptotic normality and is *asymptotically optimal* (Hájek–Le Cam sense) (Dupacová–Wets '88; Shapiro '89; King–Rockafellar '93).
- Are practical online first-order methods also optimal?
 - Dual Averaging (Duchi–Ruan '18): optimal *only when* $\mathcal{X}$ is polyhedral.
 - Projected SGD (Davis–Drusvyatskiy–Jiang '23): optimal for general $\mathcal{X}$, but lacks finite-time active-set identification and requires stronger noise assumptions.
- Both are primal-only and require explicit projection onto $\mathcal{X}$.

From unconstrained to constrained

We consider

$$\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) := \mathbb{E}[F(\mathbf{x}, \xi)], \quad \mathcal{X} := \{\mathbf{x} \in \mathbb{R}^n : h_i(\mathbf{x}) \leq 0, i \in [m]\}. \quad (\text{P})$$

Assume $F(\cdot, \xi)$ is C^2 in $\mathbf{x}$ (a.s.) and each h_i is C^2 .

- **SAA baseline.** Sample-average approximation admits asymptotic normality and is *asymptotically optimal* (Hájek–Le Cam sense) (Dupacová–Wets '88; Shapiro '89; King–Rockafellar '93).
- **Are practical online first-order methods also optimal?**
 - **Dual Averaging** (Duchi–Ruan '18): optimal *only when* $\mathcal{X}$ **is polyhedral**.
 - **Projected SGD** (Davis–Drusvyatskiy–Jiang '23): optimal for general $\mathcal{X}$, but **lacks finite-time active-set identification** and requires stronger noise assumptions.
- Both are **primal-only** and require explicit projection onto $\mathcal{X}$.

Any practical *and* optimal primal–dual method?

- **Asymptotically optimal** for general constrained problems
(Hájek–Le Cam efficiency; matches information-theoretic lower bounds)
- **Projection-free** onto $\mathcal{X}$ (only easy projection onto a redundant compact X_0 if needed)
- **Joint inference**: construct confidence intervals for both $\mathbf{x}^*$ *and* λ^*

Any practical *and* optimal primal–dual method?

- **Asymptotically optimal** for general constrained problems
(Hájek–Le Cam efficiency; matches information-theoretic lower bounds)
- **Projection-free** onto $\mathcal{X}$ (only easy projection onto a redundant compact X_0 if needed)
- **Joint inference**: construct confidence intervals for both $\mathbf{x}^*$ *and* λ^*

Stochastic Linearized Augmented Lagrangian Method (S-LALM)

Define the augmented Lagrangian function (Ben-Tal & Zibulevsky, 1997; Xu, 2021):

$$\mathcal{L}_\rho(\mathbf{x}, \boldsymbol{\lambda}) := f(\mathbf{x}) + \frac{\rho}{2} \sum_{i=1}^m \left[\max \left\{ h_i(\mathbf{x}) + \frac{\lambda_i}{\rho}, 0 \right\}^2 - \frac{\lambda_i^2}{\rho^2} \right].$$

Update rules:

$$\begin{aligned} \mathbf{x}^{k+1} &\leftarrow \Pi_{\mathcal{X}_0}(\mathbf{x}^k - \alpha_k (\nabla_{\mathbf{x}} \mathcal{L}_\rho(\mathbf{x}^k, \boldsymbol{\lambda}^k) + \xi^k)), \\ \boldsymbol{\lambda}^{k+1} &\leftarrow \boldsymbol{\lambda}^k + \beta_k \cdot \max \left\{ -\frac{\boldsymbol{\lambda}^k}{\rho}, h(\mathbf{x}^k) \right\}. \end{aligned}$$

- $\mathcal{X}_0$: a **compact redundant constraint set** with $\mathbf{x}^* \in \text{int}(\mathcal{X}_0)$ (e.g., a simple norm ball, easy to project onto).
- The stochastic noise ξ^k enters **only through the primal update**.
- Uses the **simplest gradient update** — no extrapolation or momentum.

Stochastic Linearized Augmented Lagrangian Method (S-LALM)

Define the augmented Lagrangian function (Ben-Tal & Zibulevsky, 1997; Xu, 2021):

$$\mathcal{L}_\rho(\mathbf{x}, \boldsymbol{\lambda}) := f(\mathbf{x}) + \frac{\rho}{2} \sum_{i=1}^m \left[\max \left\{ h_i(\mathbf{x}) + \frac{\lambda_i}{\rho}, 0 \right\}^2 - \frac{\lambda_i^2}{\rho^2} \right].$$

Update rules:

$$\begin{aligned} \mathbf{x}^{k+1} &\leftarrow \Pi_{\mathcal{X}_0}(\mathbf{x}^k - \alpha_k(\nabla_{\mathbf{x}} \mathcal{L}_\rho(\mathbf{x}^k, \boldsymbol{\lambda}^k) + \xi^k)), \\ \boldsymbol{\lambda}^{k+1} &\leftarrow \boldsymbol{\lambda}^k + \beta_k \cdot \max \left\{ -\frac{\boldsymbol{\lambda}^k}{\rho}, h(\mathbf{x}^k) \right\}. \end{aligned}$$

- $\mathcal{X}_0$: a compact redundant constraint set with $\mathbf{x}^* \in \text{int}(\mathcal{X}_0)$ (e.g., a simple norm ball, easy to project onto).
- The stochastic noise ξ^k enters **only through the primal update**.
- Uses the **simplest gradient update** — no extrapolation or momentum.

Stochastic Linearized Augmented Lagrangian Method (S-LALM)

Define the augmented Lagrangian function (Ben-Tal & Zibulevsky, 1997; Xu, 2021):

$$\mathcal{L}_\rho(\mathbf{x}, \boldsymbol{\lambda}) := f(\mathbf{x}) + \frac{\rho}{2} \sum_{i=1}^m \left[\max \left\{ h_i(\mathbf{x}) + \frac{\lambda_i}{\rho}, 0 \right\}^2 - \frac{\lambda_i^2}{\rho^2} \right].$$

Update rules:

$$\begin{aligned} \mathbf{x}^{k+1} &\leftarrow \Pi_{\mathcal{X}_0}(\mathbf{x}^k - \alpha_k(\nabla_{\mathbf{x}} \mathcal{L}_\rho(\mathbf{x}^k, \boldsymbol{\lambda}^k) + \xi^k)), \\ \boldsymbol{\lambda}^{k+1} &\leftarrow \boldsymbol{\lambda}^k + \beta_k \cdot \max \left\{ -\frac{\boldsymbol{\lambda}^k}{\rho}, h(\mathbf{x}^k) \right\}. \end{aligned}$$

- $\mathcal{X}_0$: a **compact redundant constraint set** with $\mathbf{x}^* \in \text{int}(\mathcal{X}_0)$ (e.g., a simple norm ball, easy to project onto).
- The stochastic noise ξ^k enters **only through the primal update**.
- Uses the **simplest gradient update** — no extrapolation or momentum.

Standard regularity conditions

Intuitively, we impose conditions ensuring that $(\mathbf{x}^*, \boldsymbol{\lambda}^*)$ is **identifiable**.

- **Active set:** $\mathcal{I} := \{ i \in [m] : h_i(\mathbf{x}^*) = 0 \}$.
- **(LICQ)** The gradients $\{\nabla h_i(\mathbf{x}^*)\}_{i \in \mathcal{I}}$ are linearly independent; equivalently, $D := [\nabla h_i(\mathbf{x}^*)]_{i \in \mathcal{I}}$ has full column rank.
- **(SC)** Strict complementarity: $\lambda_i^* > 0$ for all $i \in \mathcal{I}$.
- **(SOSC)** For all nonzero $d \in \ker(D^\top)$,
$$d^\top \nabla_{\mathbf{x}\mathbf{x}}^2 \mathcal{L}(\mathbf{x}^*, \boldsymbol{\lambda}^*) d > 0.$$

Standard regularity conditions

Intuitively, we impose conditions ensuring that $(\mathbf{x}^*, \boldsymbol{\lambda}^*)$ is **identifiable**.

- **Active set:** $\mathcal{I} := \{ i \in [m] : h_i(\mathbf{x}^*) = 0 \}$.
- **(LICQ)** The gradients $\{\nabla h_i(\mathbf{x}^*)\}_{i \in \mathcal{I}}$ are linearly independent; equivalently, $D := [\nabla h_i(\mathbf{x}^*)]_{i \in \mathcal{I}}$ has full column rank.
- **(SC)** Strict complementarity: $\lambda_i^* > 0$ for all $i \in \mathcal{I}$.
- **(SOSC)** For all nonzero $d \in \ker(D^\top)$,
$$d^\top \nabla_{\mathbf{x}\mathbf{x}}^2 \mathcal{L}(\mathbf{x}^*, \boldsymbol{\lambda}^*) d > 0.$$

Main theorem

Theorem (Li–Ruan, 2026)

Let $\alpha_k \propto k^{-\theta}$, $\beta_k \propto k^{-\theta}$ with $\theta \in (\frac{1}{2}, 1)$. Assume (LICQ), (SC), (SOSC) and standard noise conditions. Then

$$\sqrt{k} \begin{bmatrix} \bar{\mathbf{x}}_k - \mathbf{x}^* \\ \bar{\boldsymbol{\lambda}}_k - \boldsymbol{\lambda}^* \end{bmatrix} \xrightarrow{d} \mathcal{N}\left(0, \begin{bmatrix} J_p \Sigma J_p^\top & J_p \Sigma J_d^\top \\ J_d \Sigma J_p^\top & J_d \Sigma J_d^\top \end{bmatrix}\right),$$

where

$$J_p := (P_T H_0 P_T)^\dagger, \quad J_d := (D^\top D)^{-1} D^\top (I - H_0 (P_T H_0 P_T)^\dagger),$$

$T := \ker(D^\top)$ is the tangent space of the active manifold, $P_T = I - D(D^\top D)^{-1} D^\top$ is the orthogonal projector, and H_0 is the Hessian of the Lagrangian w.r.t. $\mathbf{x}$ at $(\mathbf{x}^*, \boldsymbol{\lambda}^*)$.

- The primal covariance $J_p \Sigma J_p^\top$ matches the information-theoretic lower bound (Duchi–Ruan '18; Davis–Drusvyatskiy–Jiang '23) — asymptotically optimal.

Theorem (Li–Ruan, 2026)

Let $\alpha_k \propto k^{-\theta}$, $\beta_k \propto k^{-\theta}$ with $\theta \in (\frac{1}{2}, 1)$. Assume (LICQ), (SC), (SOSC) and standard noise conditions. Then

$$\sqrt{k} \begin{bmatrix} \bar{\mathbf{x}}_k - \mathbf{x}^* \\ \bar{\boldsymbol{\lambda}}_k - \boldsymbol{\lambda}^* \end{bmatrix} \xrightarrow{d} \mathcal{N}\left(0, \begin{bmatrix} J_p \Sigma J_p^\top & J_p \Sigma J_d^\top \\ J_d \Sigma J_p^\top & J_d \Sigma J_d^\top \end{bmatrix}\right),$$

where

$$J_p := (P_T H_0 P_T)^\dagger, \quad J_d := (D^\top D)^{-1} D^\top (I - H_0 (P_T H_0 P_T)^\dagger),$$

$T := \ker(D^\top)$ is the tangent space of the active manifold, $P_T = I - D(D^\top D)^{-1} D^\top$ is the orthogonal projector, and H_0 is the Hessian of the Lagrangian w.r.t. $\mathbf{x}$ at $(\mathbf{x}^*, \boldsymbol{\lambda}^*)$.

- The primal covariance $J_p \Sigma J_p^\top$ **matches the information-theoretic lower bound** (Duchi–Ruan '18; Davis–Drusvyatskiy–Jiang '23) — *asymptotically optimal*.

Proof sketch

- (★★★ Almost sure convergence)

$$\mathbf{x}^k \xrightarrow{\text{a.s.}} \mathbf{x}^*, \quad \boldsymbol{\lambda}^k \xrightarrow{\text{a.s.}} \boldsymbol{\lambda}^*.$$

- (★ Finite-time active-set identification) With probability one, there exists a finite (random) index $K < \infty$ such that, for all $k \geq K$,

$$\{i \in [m] : \rho h_i(\mathbf{x}^k) + \lambda_i^k > 0\} = \mathcal{I}.$$

- (★★ Local convergence rate) As $k \rightarrow \infty$,

$$\frac{1}{\sqrt{k}} \sum_{i=1}^k \|\mathbf{x}^i - \mathbf{x}^*\|^2 \rightarrow 0, \quad \frac{1}{\sqrt{k}} \sum_{i=1}^k \|\boldsymbol{\lambda}^i - \boldsymbol{\lambda}^*\|^2 \rightarrow 0 \quad \text{a.s.}$$

Key challenge

Primal-dual asymmetry: there is *no explicit curvature* on the dual side.

Proof sketch

- (★★★ Almost sure convergence)

$$\mathbf{x}^k \xrightarrow{\text{a.s.}} \mathbf{x}^*, \quad \boldsymbol{\lambda}^k \xrightarrow{\text{a.s.}} \boldsymbol{\lambda}^*.$$

- (★ Finite-time active-set identification) With probability one, there exists a finite (random) index $K < \infty$ such that, for all $k \geq K$,

$$\{i \in [m] : \rho h_i(\mathbf{x}^k) + \lambda_i^k > 0\} = \mathcal{I}.$$

- (★★ Local convergence rate) As $k \rightarrow \infty$,

$$\frac{1}{\sqrt{k}} \sum_{i=1}^k \|\mathbf{x}^i - \mathbf{x}^*\|^2 \rightarrow 0, \quad \frac{1}{\sqrt{k}} \sum_{i=1}^k \|\boldsymbol{\lambda}^i - \boldsymbol{\lambda}^*\|^2 \rightarrow 0 \quad \text{a.s.}$$

Key challenge

Primal–dual asymmetry: there is *no explicit curvature* on the dual side.

Proof sketch

- (★★★ Almost sure convergence)

$$\mathbf{x}^k \xrightarrow{\text{a.s.}} \mathbf{x}^*, \quad \boldsymbol{\lambda}^k \xrightarrow{\text{a.s.}} \boldsymbol{\lambda}^*.$$

- (★ Finite-time active-set identification) With probability one, there exists a finite (random) index $K < \infty$ such that, for all $k \geq K$,

$$\{i \in [m] : \rho h_i(\mathbf{x}^k) + \lambda_i^k > 0\} = \mathcal{I}.$$

- (★★ Local convergence rate) As $k \rightarrow \infty$,

$$\frac{1}{\sqrt{k}} \sum_{i=1}^k \|\mathbf{x}^i - \mathbf{x}^*\|^2 \rightarrow 0, \quad \frac{1}{\sqrt{k}} \sum_{i=1}^k \|\boldsymbol{\lambda}^i - \boldsymbol{\lambda}^*\|^2 \rightarrow 0 \quad \text{a.s.}$$

Key challenge

Primal-dual asymmetry: there is *no explicit curvature* on the dual side.

Proof sketch

- (★★★ Almost sure convergence)

$$\mathbf{x}^k \xrightarrow{\text{a.s.}} \mathbf{x}^*, \quad \boldsymbol{\lambda}^k \xrightarrow{\text{a.s.}} \boldsymbol{\lambda}^*.$$

- (★ Finite-time active-set identification) With probability one, there exists a finite (random) index $K < \infty$ such that, for all $k \geq K$,

$$\{i \in [m] : \rho h_i(\mathbf{x}^k) + \lambda_i^k > 0\} = \mathcal{I}.$$

- (★★ Local convergence rate) As $k \rightarrow \infty$,

$$\frac{1}{\sqrt{k}} \sum_{i=1}^k \|\mathbf{x}^i - \mathbf{x}^*\|^2 \rightarrow 0, \quad \frac{1}{\sqrt{k}} \sum_{i=1}^k \|\boldsymbol{\lambda}^i - \boldsymbol{\lambda}^*\|^2 \rightarrow 0 \quad \text{a.s.}$$

Key challenge

Primal-dual asymmetry: there is *no explicit curvature* on the dual side.

Proof sketch

- (★★★ Almost sure convergence)

$$\mathbf{x}^k \xrightarrow{\text{a.s.}} \mathbf{x}^*, \quad \boldsymbol{\lambda}^k \xrightarrow{\text{a.s.}} \boldsymbol{\lambda}^*.$$

- (★ Finite-time active-set identification) With probability one, there exists a finite (random) index $K < \infty$ such that, for all $k \geq K$,

$$\{i \in [m] : \rho h_i(\mathbf{x}^k) + \lambda_i^k > 0\} = \mathcal{I}.$$

- (★★ Local convergence rate) As $k \rightarrow \infty$,

$$\frac{1}{\sqrt{k}} \sum_{i=1}^k \|\mathbf{x}^i - \mathbf{x}^*\|^2 \rightarrow 0, \quad \frac{1}{\sqrt{k}} \sum_{i=1}^k \|\boldsymbol{\lambda}^i - \boldsymbol{\lambda}^*\|^2 \rightarrow 0 \quad \text{a.s.}$$

Key challenge

Primal–dual asymmetry: there is *no explicit curvature* on the dual side.

Almost sure convergence of S-LALM

Theorem (Li–Ruan, 2026)

Let $\{(\mathbf{x}^k, \boldsymbol{\lambda}^k)\}_{k \geq 0}$ be generated by S-LALM under the (SOSC) condition and standard noise assumptions. Then, almost surely, every cluster point of $\{(\mathbf{x}^k, \boldsymbol{\lambda}^k)\}$ is a KKT point of Problem (P).

- The only regularity assumption required is (SOSC).
- (SOSC) implies the **quadratic growth property** of the nonsmooth augmented Lagrangian:
$$\mathcal{L}_\rho(\mathbf{x}, \boldsymbol{\lambda}^*) - \mathcal{L}_\rho(\mathbf{x}^*, \boldsymbol{\lambda}^*) \geq \varepsilon \|\mathbf{x} - \mathbf{x}^*\|^2, \quad \forall \mathbf{x} \in \mathcal{X}_0.$$
- This result is of **independent interest**—it establishes almost sure convergence under minimal assumptions.

Almost sure convergence of S-LALM

Theorem (Li–Ruan, 2026)

Let $\{(\mathbf{x}^k, \boldsymbol{\lambda}^k)\}_{k \geq 0}$ be generated by S-LALM under the (SOSC) condition and standard noise assumptions. Then, almost surely, every cluster point of $\{(\mathbf{x}^k, \boldsymbol{\lambda}^k)\}$ is a KKT point of Problem (P).

- The only regularity assumption required is **(SOSC)**.
- (SOSC) implies the **quadratic growth property** of the nonsmooth augmented Lagrangian:
$$\mathcal{L}_\rho(\mathbf{x}, \boldsymbol{\lambda}^*) - \mathcal{L}_\rho(\mathbf{x}^*, \boldsymbol{\lambda}^*) \geq \varepsilon \|\mathbf{x} - \mathbf{x}^*\|^2, \quad \forall \mathbf{x} \in \mathcal{X}_0.$$
- This result is of **independent interest**—it establishes almost sure convergence under minimal assumptions.

Example

Example

Suppose that $\xi \sim \mathcal{N}(0, I_3)$. We consider

$$\begin{aligned} \min_{\mathbf{x} \in \mathcal{X}_0} \quad & \mathbb{E}[(-\mathbf{e}_3 + \xi)^\top \mathbf{x}] \\ \text{s.t.} \quad & 2\sqrt{3}(\|\mathbf{x} - \mathbf{e}_1\|^2 - 4) \leq 0, \\ & 2\sqrt{3}(\|\mathbf{x} + \mathbf{e}_1\|^2 - 4) \leq 0, \end{aligned}$$

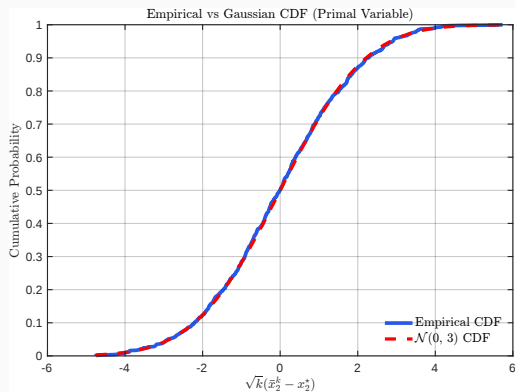
where $\mathcal{X}_0 = \{\mathbf{x} \in \mathbb{R}^3 : \|\mathbf{x}\|_\infty \leq 3\}$.

$$\sqrt{k} \begin{bmatrix} \bar{\mathbf{x}}^k - \mathbf{x}^* \\ \bar{\boldsymbol{\lambda}}^k - \boldsymbol{\lambda}^* \end{bmatrix} \xrightarrow{d} \mathcal{N}\left(0, \begin{bmatrix} \text{diag}(0, 3, 0) & 0 \\ 0 & \begin{bmatrix} 1 & -\frac{1}{2} \\ -\frac{1}{2} & 1 \end{bmatrix} \end{bmatrix}\right).$$

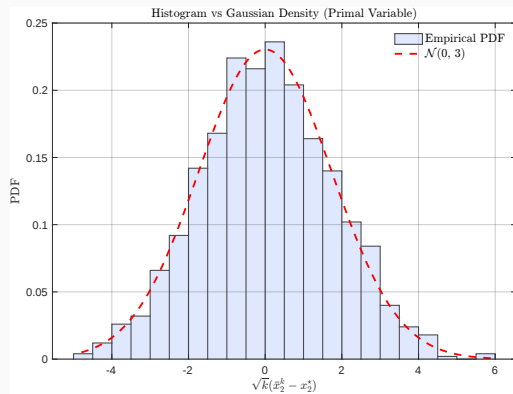
- Hyperparameters set up:

$$\rho = 0.2, \quad \alpha_k = \frac{0.1}{k^{0.51}}, \quad \beta_k = \frac{0.05}{k^{0.51}}.$$

Primal Variable

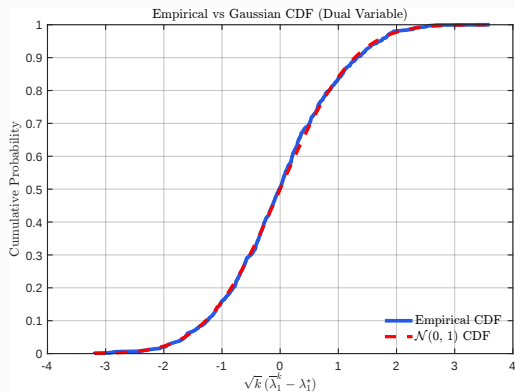


(a) Empirical CDF of the primal variable

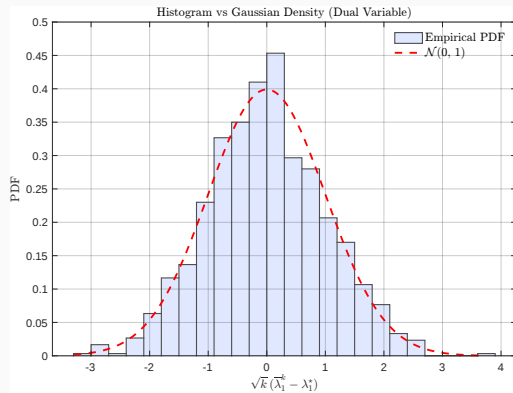


(b) Empirical PDF of the primal variable

Dual Variable

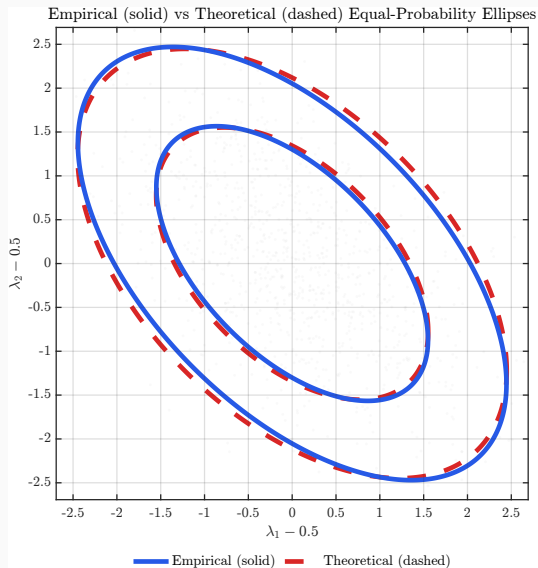


(a) Empirical CDF of the dual variable



(b) Empirical PDF of the dual variable

Covariance of Dual Variables



- S-LALM gives a projection-light online alternative to SAA.
- Averaging yields asymptotic normality with optimal primal covariance.
- The same limit theory gives joint inference for (x^*, λ^*) .

Thank you for your listening!
Any questions?