

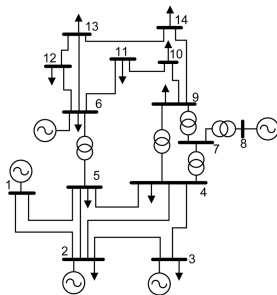
Optimal Decisions with Messy Data and Uncertainty

Bart P.G. Van Parys

CWI Amsterdam

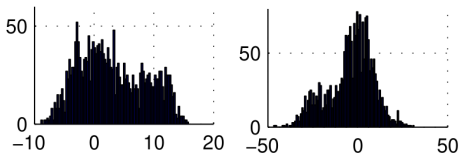
27 October 2025

Practical Decisions require Dealing with Uncertainty and Data



Optimal Power Flow (Roald et al., 2015)

- ▶ **minimize:** operation cost
- ▶ **decide:** generation commitment
- ▶ **uncertain:** forecast error

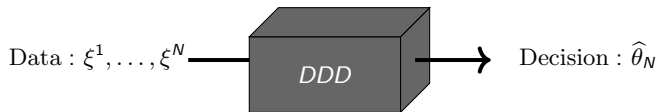


Classical approaches for decisions in the face of uncertainty

$$\min_{\theta \in \Theta} \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)]$$

use the **distribution** $\mathbb{P}$ of $\xi \in \Xi$ as a primitive.

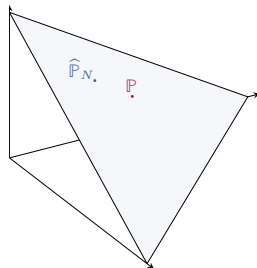
Modern approaches try to learn good decisions with **data**



Empirical Risk Minimization (ERM/SAA)

- ▶ Empirical distribution $\hat{\mathbb{P}}_N = \frac{1}{N} \sum_{i=1}^N \delta_{\xi^i}$.
- ▶ Empirical risk minimization (Birge and Louveaux, 2011)

$$\begin{aligned}\hat{\theta}_N &\in \arg \min_{\theta \in \Theta} \mathbf{E}_{\hat{\mathbb{P}}_N} [\ell(\theta, \xi)] \\ &\in \arg \min_{\theta \in \Theta} \frac{1}{N} \sum_{i=1}^N \ell(\theta, \xi^i).\end{aligned}$$



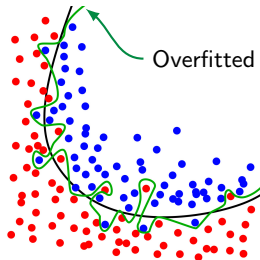
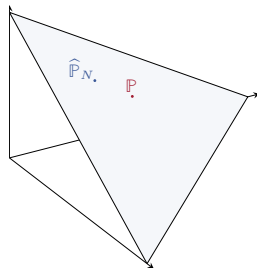
Empirical Risk Minimization (ERM/SAA)

- ▶ Empirical distribution $\hat{\mathbb{P}}_N = \frac{1}{N} \sum_{i=1}^N \delta_{\xi^i}$.
- ▶ Empirical risk minimization (Birge and Louveau, 2011)

$$\begin{aligned}\hat{\theta}_N &\in \arg \min_{\theta \in \Theta} \mathbf{E}_{\hat{\mathbb{P}}_N} [\ell(\theta, \xi)] \\ &\in \arg \min_{\theta \in \Theta} \frac{1}{N} \sum_{i=1}^N \ell(\theta, \xi^i).\end{aligned}$$

- ▶ “Overfitting”, “Error Maximization Effect”, “Optimizer’s Curse” (Michaud, 1989)

$$\underbrace{\mathbf{E}_{\mathbb{P}} [\ell(\hat{\theta}_N, \xi)]}_{\text{actual cost}} > \underbrace{\mathbf{E}_{\hat{\mathbb{P}}_N} [\ell(\hat{\theta}_N, \xi)]}_{\text{predicted cost}}$$



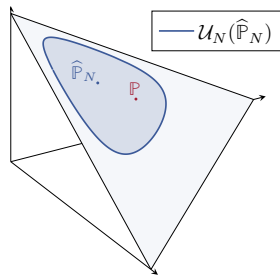
Worst-Case Formulations

Worst-case formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{Q \in \mathcal{U}_N(\hat{\mathbb{P}}_N)} \mathbf{E}_Q [\ell(\theta, \xi)]$$

and guard against overfitting, i.e.,

$$\mathbb{P} \in \mathcal{U}_N(\hat{\mathbb{P}}_N) \implies \underbrace{\max_{Q \in \mathcal{U}_N(\hat{\mathbb{P}}_N)} \mathbf{E}_Q [\ell(\hat{\theta}_N, \xi)]}_{\text{predicted cost}} \geq \underbrace{\mathbf{E}_{\mathbb{P}} [\ell(\hat{\theta}_N, \xi)]}_{\text{actual cost}}.$$



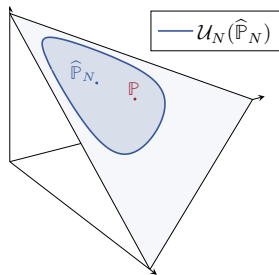
Worst-Case Formulations

Worst-case formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{Q \in \mathcal{U}_N(\hat{\mathbb{P}}_N)} \mathbf{E}_Q [\ell(\theta, \xi)]$$

and guard against overfitting, i.e.,

$$\mathbb{P} \in \mathcal{U}_N(\hat{\mathbb{P}}_N) \implies \underbrace{\max_{Q \in \mathcal{U}_N(\hat{\mathbb{P}}_N)} \mathbf{E}_Q [\ell(\hat{\theta}_N, \xi)]}_{\text{predicted cost}} \geq \underbrace{\mathbf{E}_{\mathbb{P}} [\ell(\hat{\theta}_N, \xi)]}_{\text{actual cost}}.$$



Tractability depends on whether the **distributional optimization** problem

$$\max_{Q \in \mathcal{U}_N(\hat{\mathbb{P}}_N)} \mathbf{E}_Q [\ell(\theta, \xi)]$$

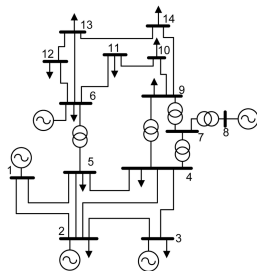
can be solved efficiently.

Table of Contents

1. The Dawn of DRO
2. Efficient Decisions from Data
3. Holistic Robustness
4. Robust Statistics
5. Conclusions

Power Flow Problem:

- Generators ($\odot$) generate power to satisfy demand ($\downarrow$)

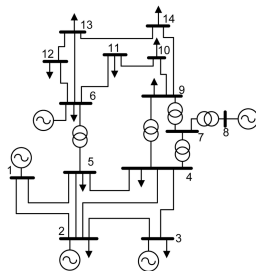


Power Flow Problem:

- Generators (⊗) generate power to satisfy demand (↓)
- Power through every line is limited

$$p_k = \underbrace{p_{nom,k}}_{nominal} < \underbrace{p_{max,k}}_{trip\ limit} \quad \forall k$$

if generators honor their power commitment

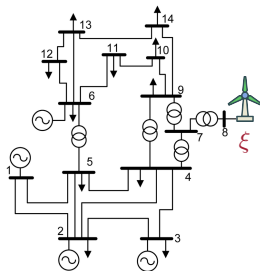


Power Flow Problem:

- Generators (⊙) generate power to satisfy demand (↓)
- Power through every line is limited

$$p_k = \underbrace{p_{nom,k}}_{\text{nominal}} + \underbrace{a_k^\top \xi}_{\text{correction}} < \underbrace{p_{max,k}}_{\text{trip limit}} \quad \forall k$$

- Renewable generators suffer **forecast errors** ξ



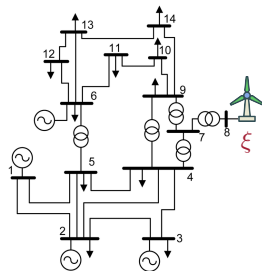
Example: Power Grid Safety (Roald et al., 2015)

Power Flow Problem:

- ▶ Generators ($\odot$) generate power to satisfy demand ($\downarrow$)
- ▶ Power through every line is limited

$$P = \{ \xi : \underbrace{p_{nom,k}}_{\text{nominal}} + \underbrace{a_k^\top \xi}_{\text{correction}} < \underbrace{p_{max,k}}_{\text{trip limit}} \quad \forall k \}$$

- ▶ Renewable generators suffer **forecast errors** ξ



Probability of a line tripping is $\mathbb{P}(\xi \notin P)$:

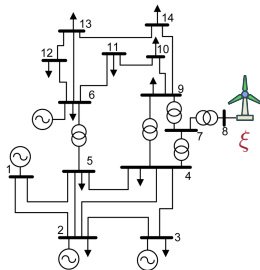
Example: Power Grid Safety (Roald et al., 2015)

Power Flow Problem:

- ▶ Generators (⊗) generate power to satisfy demand (↓)
- ▶ Power through every line is limited

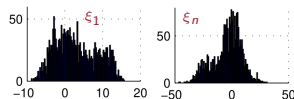
$$P = \{\xi : \underbrace{p_{nom,k}}_{nominal} + \underbrace{a_k^\top \xi}_{correction} < \underbrace{p_{max,k}}_{trip\ limit} \quad \forall k\}$$

- ▶ Renewable generators suffer **forecast errors** ξ



Probability of a line tripping is $P(\xi \notin P)$:

- ▶ Distribution $\xi \sim P$ is not known



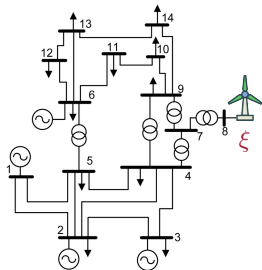
Example: Power Grid Safety (Roald et al., 2015)

Power Flow Problem:

- ▶ Generators (⊗) generate power to satisfy demand (↓)
- ▶ Power through every line is limited

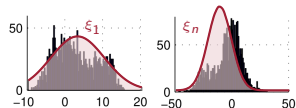
$$P = \{\xi : \underbrace{p_{nom,k}}_{nominal} + \underbrace{a_k^\top \xi}_{correction} < \underbrace{p_{max,k}}_{trip\ limit} \quad \forall k\}$$

- ▶ Renewable generators suffer **forecast errors** ξ



Probability of a line tripping is $\mathbb{P}(\xi \notin P)$:

- ▶ Distribution $\xi \sim \mathbb{P}$ is not known
- ▶ $\mathbb{P}$ is normal may **underestimate** trip probability



Uncertainty Quantification Problems with Moments

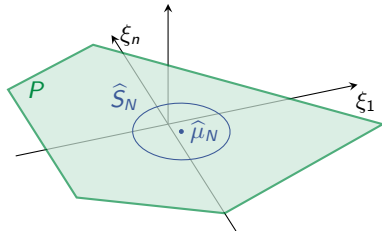
Worst-Case Probability:

$$\max \mathbb{Q}(\xi \notin P)$$

$$\text{s.t. } \mathbb{Q} \in \mathcal{P},$$

$$\int \xi d\mathbb{Q} = \hat{\mu}_N,$$

$$\int \xi \xi^\top d\mathbb{Q} = \hat{S}_N$$



Properties:

- Reduces to a tractable convex optimization problem (Bertsimas and Popescu, 2005)

Uncertainty Quantification Problems with Moments

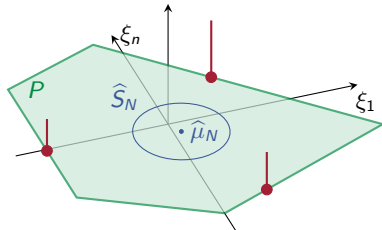
Worst-Case Probability:

$$\max \mathbb{Q}(\xi \notin P)$$

$$\text{s.t. } \mathbb{Q} \in \mathcal{P},$$

$$\int \xi d\mathbb{Q} = \hat{\mu}_N,$$

$$\int \xi \xi^\top d\mathbb{Q} = \hat{S}_N$$



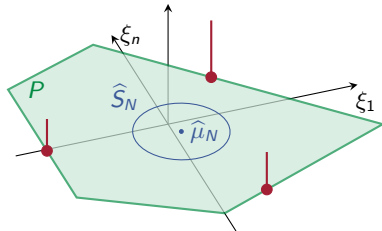
Properties:

- Reduces to a tractable convex optimization problem (Bertsimas and Popescu, 2005)
- **Pathological** worst-case distribution $\mathbb{Q}^*$

Uncertainty Quantification Problems with Moments

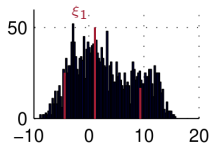
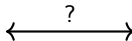
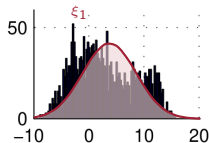
Worst-Case Probability:

$$\begin{aligned} \max \quad & \mathbb{Q}(\xi \notin P) \\ \text{s.t.} \quad & \mathbb{Q} \in \mathcal{P}, \\ & \int \xi d\mathbb{Q} = \hat{\mu}_N, \\ & \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N \end{aligned}$$



Properties:

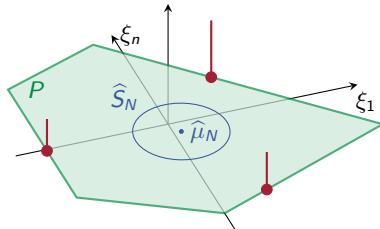
- ▶ Reduces to a tractable convex optimization problem (Bertsimas and Popescu, 2005)
- ▶ **Pathological** worst-case distribution $\mathbb{Q}^*$



Uncertainty Quantification Problems with Moments

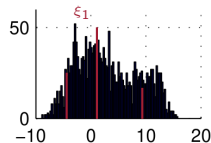
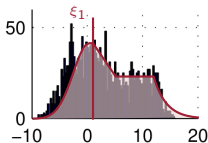
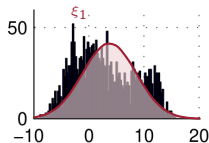
Worst-Case Probability:

$$\begin{aligned} \max \quad & \mathbb{Q}(\xi \notin P) \\ \text{s.t.} \quad & \mathbb{Q} \in \mathcal{P}, \\ & \int \xi d\mathbb{Q} = \hat{\mu}_N, \\ & \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N \end{aligned}$$



Properties:

- ▶ Reduces to a tractable convex optimization problem (Bertsimas and Popescu, 2005)
- ▶ **Pathological** worst-case distribution $\mathbb{Q}^*$
- ▶ Unimodal distribution $\mathbb{Q}^*$ (Van Parys, P. J. Goulart, et al., 2016)

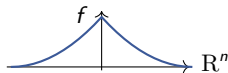


$$\max \mathbb{Q}(\xi \notin P)$$

$$\text{s.t. } \mathbb{Q} \in \mathcal{S}, \int \xi d\mathbb{Q} = \hat{\mu}_N, \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N$$

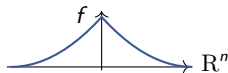
$$\begin{aligned} \max \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q} \in \mathcal{S}, \int \xi d\mathbb{Q} = \hat{\mu}_N, \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N \end{aligned}$$

- conditional value-at-risk quantification in financial risk management
- distributions with common tail properties (Van Parys, P. Goulart, and Morari, 2019)



$$\begin{aligned} \max \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q} \in \mathcal{S}, \int \xi d\mathbb{Q} = \hat{\mu}_N, \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N \end{aligned}$$

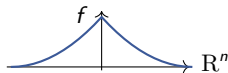
- ▶ conditional value-at-risk quantification in financial risk management
- ▶ distributions with common tail properties (Van Parys, P. Goulart, and Morari, 2019)



+ Transfer information to tail regions where little data is observed

$$\begin{aligned} \max \quad & \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] \\ \text{s.t.} \quad & \mathbb{Q} \in \mathcal{S}, \int \xi d\mathbb{Q} = \hat{\mu}_N, \int \xi \xi^\top d\mathbb{Q} = \hat{S}_N \end{aligned}$$

- ▶ conditional value-at-risk quantification in financial risk management
- ▶ distributions with common tail properties (Van Parys, P. Goulart, and Morari, 2019)



- + Transfer information to tail regions where little data is observed
 - Learns from data **indirectly** through its moments

Table of Contents

1. The Dawn of DRO
2. Efficient Decisions from Data
3. Holistic Robustness
4. Robust Statistics
5. Conclusions

Divergence Balls $\mathcal{U}_N(\hat{\mathbb{P}}_N) = \{\mathbb{Q} : D(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r\}$

Formulations

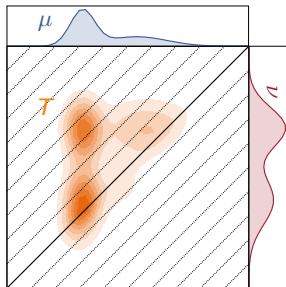
$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{D(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$$

based on tractable statistical distances can learn **directly** from data.

Examples:

- **Wasserstein distance** is defined as

$$W(\mu, \nu) := \inf_{T \in \mathcal{T}(\mu, \nu)} \int \|\xi - \xi'\| dT$$



Divergence Balls $\mathcal{U}_N(\hat{\mathbb{P}}_N) = \{\mathbb{Q} : D(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r\}$

Formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{D(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$$

based on tractable statistical distances can learn **directly** from data.

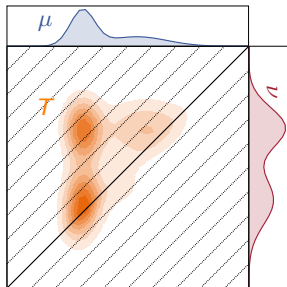
Examples:

- **Wasserstein distance** is defined as

$$W(\mu, \nu) := \inf_{T \in \mathcal{T}(\mu, \nu)} \int \|\xi - \xi'\| dT$$

- **KL-Divergence** is defined as

$$\text{KL}(\mu, \nu) = \int \log \left(\frac{d\mu}{d\nu} \right) d\mu$$



Example: Wasserstein DRO

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{w(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}}[\ell(\theta, \xi)] .$$

Example: Wasserstein DRO

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{w(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}}[\ell(\theta, \xi)].$$

- Assume data $\{\xi^i \in \mathbb{R}^n\}_{i=1}^N$ is drawn independently from $\mathbb{P}$.

Example: Wasserstein DRO

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{W(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}}[\ell(\theta, \xi)].$$

- Assume data $\{\xi^i \in \mathbb{R}^n\}_{i=1}^N$ is drawn independently from $\mathbb{P}$.
- For Wasserstein balls we can show the **concentration inequality** (Fournier and Guillin, 2015)

$$\text{Prob}_{\mathbb{P}} \left[W(\hat{\mathbb{P}}_N, \mathbb{P}) \leq r \right] \geq 1 - c_1 \exp(-c_2 N r^{\max(n, 2)})$$

with $c_1 > 0$ and $c_2 > 0$ positive constants.

Example: Wasserstein DRO

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \max_{W(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}}[\ell(\theta, \xi)].$$

- ▶ Assume data $\{\xi^i \in \mathbb{R}^n\}_{i=1}^N$ is drawn independently from $\mathbb{P}$.
- ▶ For Wasserstein balls we can show the **out-of-sample** (OOS) guarantee

$$\text{Prob}_{\mathbb{P}} \left[\underbrace{\max_{W(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}}[\ell(\theta, \xi)]}_{\text{predicted cost}} \geq \underbrace{\mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)]}_{\text{actual cost}} \quad \forall \theta \in \Theta \right] \geq 1 - c_1 \exp(-c_2 N r^{\max(n, 2)})$$

with $c_1 > 0$ and $c_2 > 0$ positive constants.

Why Robust Optimization?

There are many ways to protect decisions against overfitting:

- ▶ OR: Robustification,
- ▶ Statistics : Regularization,
- ▶ Deep Learning: Early stopping, drop out, ...

Why Robust Optimization?

There are many ways to protect decisions against overfitting:

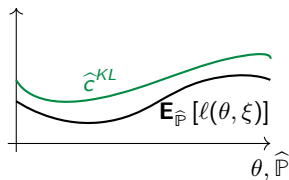
- ▶ OR: Robustification,
- ▶ Statistics : Regularization,
- ▶ Deep Learning: Early stopping, drop out, ...

Why should we use robust optimization?

Kullback-Leibler Robust Formulation

Kullback-Leibler (KL) formulation:

$$\begin{aligned}\hat{c}^{KL}(\theta, \hat{\mathbb{P}}) &= \max_{\mathbb{Q}} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q} &\in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}, \mathbb{Q}) &\leq r.\end{aligned}$$



Thm (Van Parys, Esfahani, et al., 2020): $\Xi \subset \mathbb{R}^n$ is compact, $\{\xi \mapsto \ell(\theta, \xi)\}_{\theta \in \Theta}$ is equicontinuous.

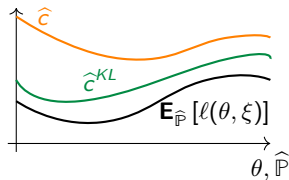
► (OOS Guarantee) we have

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}^{KL}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Kullback-Leibler Robust Formulation

Kullback-Leibler (KL) formulation:

$$\begin{aligned}\hat{c}^{KL}(\theta, \hat{\mathbb{P}}) &= \max \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q} &\in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}, \mathbb{Q}) &\leq r.\end{aligned}$$



Thm (Van Parys, Esfahani, et al., 2020): $\Xi \subset \mathbb{R}^n$ is compact, $\{\xi \mapsto \ell(\theta, \xi)\}_{\theta \in \Theta}$ is equicontinuous.

► (OOS Guarantee) we have

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}^{KL}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

► (Efficiency) Furthermore, let $\hat{c}$ satisfy

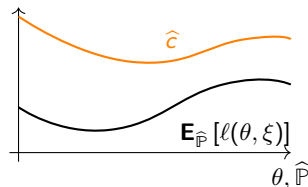
$$\text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Then, $\hat{c}(\theta, \hat{\mathbb{P}}) \geq \hat{c}^{KL}(\theta, \hat{\mathbb{P}})$ for all $\theta \in \Theta$ and $\hat{\mathbb{P}} \in \mathcal{P}$.

Meta Optimization Problem

Data driven formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \hat{c}(\theta, \hat{\mathbb{P}}_N)$$



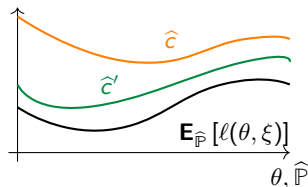
Feasibility: Consider formulations who enjoy **OOS guarantee**

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Meta Optimization Problem

Data driven formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \hat{c}(\theta, \hat{\mathbb{P}}_N)$$



Feasibility: Consider formulations who enjoy **OOS guarantee**

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

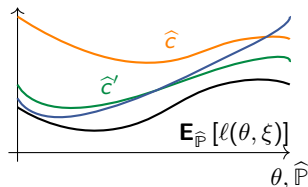
Conservatism: Formulation $\hat{c}'$ is **less conservative** than $\hat{c}$ if and only if

$$\hat{c}' \preceq \hat{c} \iff \hat{c}'(\theta, \hat{\mathbb{P}}) \leq \hat{c}(\theta, \hat{\mathbb{P}}) \quad \forall \theta, \hat{\mathbb{P}}.$$

Meta Optimization Problem

Data driven formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \hat{c}(\theta, \hat{\mathbb{P}}_N)$$



Feasibility: Consider formulations who enjoy **OOS guarantee**

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Conservatism: Formulation $\hat{c}'$ is **less conservative** than $\hat{c}$ if and only if

$$\hat{c}' \preceq \hat{c} \iff \hat{c}'(\theta, \hat{\mathbb{P}}) \leq \hat{c}(\theta, \hat{\mathbb{P}}) \quad \forall \theta, \hat{\mathbb{P}}.$$

Meta Optimization Problem

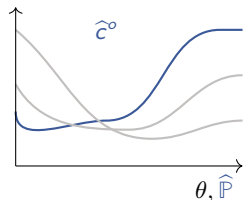
Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{cases}$$

Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{cases}$$

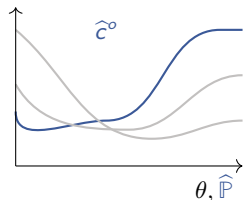


Pareto optimal solution

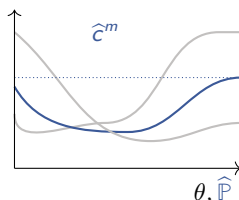
Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{cases}$$



Pareto optimal solution

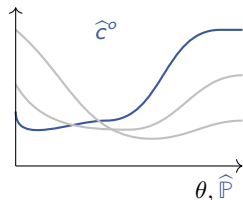


Minimax solution

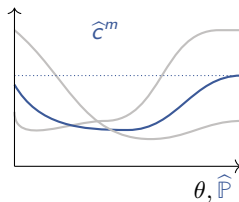
Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

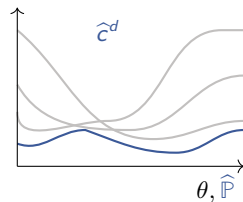
$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{cases}$$



Pareto optimal solution



Minimax solution



Pareto dominant solution

Pareto Dominant Solution

Thm (Van Parys, Esfahani, et al., 2020): The meta optimization problem

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{cases}$$

admits the **Pareto dominant** solution $\hat{c}^{KL}$.

Proof Crux: Indirect consequence of Sanov's theorem.

Pareto Dominant Solution

Thm (Van Parys, Esfahani, et al., 2020): The meta optimization problem

$$\left\{ \begin{array}{l} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P} \end{array} \right.$$

admits the **Pareto dominant** solution $\hat{c}^{KL}$.

Proof Crux: Indirect consequence of Sanov's theorem.

Generalizations:

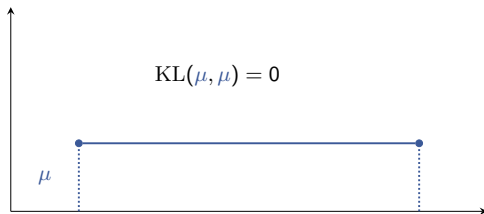
- ▶ IID data assumption can be relaxed to account for dependence (Sutter et al., 2020)
- ▶ Out-of-sample guarantee can be relaxed (Bennouna and Van Parys, 2021)

Table of Contents

1. The Dawn of DRO
2. Efficient Decisions from Data
- 3. Holistic Robustness**
4. Robust Statistics
5. Conclusions

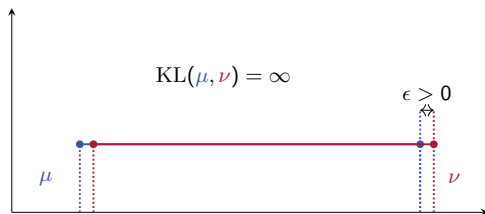
A Paradoxical Observation

Observation: KL-divergence lacks continuity



A Paradoxical Observation

Observation: KL-divergence lacks continuity

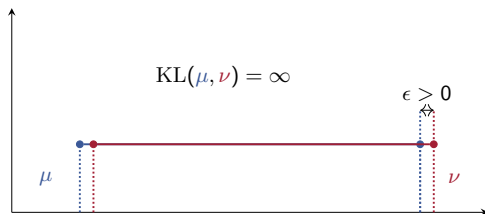


Remarks:

- KL robustness offers no protection against data corruption

A Paradoxical Observation

Observation: KL-divergence lacks continuity



Remarks:

- ▶ KL robustness offers no protection against data corruption
- ▶ Wasserstein distance does not suffer this shortcoming!

Statistical Error

Randomness of the sampled
data

$$\xi \sim \mathbb{P}$$



$$\{\xi^1, \dots, \xi^N\}$$

Statistical Error

Randomness of the sampled data

$$\xi \sim \mathbb{P}$$



$$\{\xi^1, \dots, \xi^N\}$$

Noise

Small ϵ -perturbation on each sample

$$\{\xi^1, \dots, \xi^N\}$$



$$\{\xi^1 + \delta^1, \dots, \xi^N + \delta^N\}$$
$$\|\delta^1\| \leq \epsilon, \dots, \|\delta^N\| \leq \epsilon$$

Contamination

(Huber, 1964)

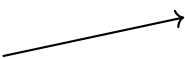
Small fraction α of all samples is wholly corrupted

$$\{\cancel{\xi^1}, \dots, \xi^N\}$$



$$\{\xi'^1, \dots, \xi^N\}$$

Existing Formulations Protect Against a Specific Overfitting Source

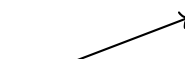
KL DRO $\max \left\{ \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] : \text{KL}(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r \right\}$  **Statistical Error**

Noise

Contamination

Existing Formulations Protect Against a Specific Overfitting Source

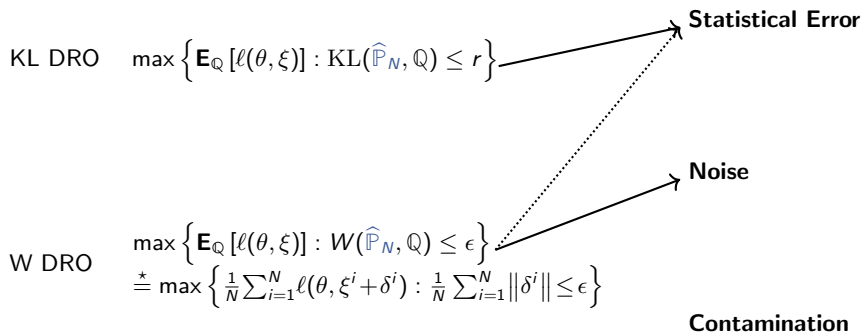
KL DRO $\max \left\{ \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] : \text{KL}(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r \right\}$  **Statistical Error**

W DRO $\max \left\{ \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)] : W(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq \epsilon \right\}$  **Noise**

$\stackrel{*}{=} \max \left\{ \frac{1}{N} \sum_{i=1}^N \ell(\theta, \xi^i + \delta^i) : \frac{1}{N} \sum_{i=1}^N \|\delta^i\| \leq \epsilon \right\}$ **Contamination**

[*] Assuming $\xi \mapsto \ell(\theta, \xi)$ is concave for all θ .

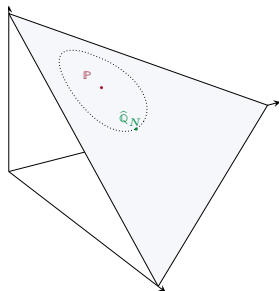
Existing Formulations Protect Against a Specific Overfitting Source



[*] Assuming $\xi \mapsto \ell(\theta, \xi)$ is concave for all θ .

Kullback-Leibler Distance

$$\begin{array}{ccc} \xi \sim \mathbb{P} & & \\ \downarrow & \text{Statistical Error} & \\ \{\xi^1, \xi^2, \dots, \xi^N\} & & \hat{\mathbb{Q}}_N \end{array}$$

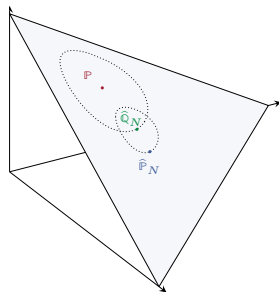


Thm: (Sanov, 1958) If $\hat{\mathbb{Q}}_N$ is an empirical distribution of independent samples from $\mathbb{P}$ then

$$\text{Prob}_{\mathbb{P}} \left[\text{KL}(\hat{\mathbb{Q}}_N, \mathbb{P}) \leq r \right] = 1 - \exp(-rN + o(N)).$$

Lévy-Prokhorov Distance

$$\begin{aligned} \xi &\sim \mathbb{P} \\ &\Downarrow \text{Statistical Error} \\ \{\xi^1, \xi^2, \dots, \xi^N\} &\quad \widehat{\mathbb{Q}}_N \\ &\Downarrow \epsilon\text{-Noise} \\ \{\xi^1 + \delta^1, \xi^2 + \delta^2, \dots, \xi^N + \delta^N\} \\ &\Downarrow \alpha\text{-Contamination} \\ \{\xi^1 + \delta^1, \xi'^2, \dots, \xi^N + \delta^N\} &\quad \widehat{\mathbb{P}}_N \end{aligned}$$

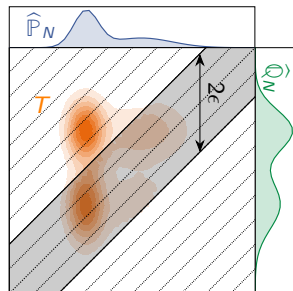


Prop: If $\widehat{\mathbb{P}}_N$ and $\widehat{\mathbb{Q}}_N$ differ by ϵ -norm bounded noise and α -contamination. Then

$$\text{LP}_\epsilon(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N) := \inf_{T \in \mathcal{T}(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N)} \int \mathbb{1} \{ \|\xi - \xi'\| \geq \epsilon \} dT \leq \alpha.$$

Lévy-Prokhorov Distance

$$\begin{aligned}
 &\xi \sim \mathbb{P} \\
 &\Downarrow \text{Statistical Error} \\
 &\{\xi^1, \xi^2, \dots, \xi^N\} \quad \widehat{\mathbb{Q}}_N \\
 &\Downarrow \epsilon\text{-Noise} \\
 &\{\xi^1 + \delta^1, \xi^2 + \delta^2, \dots, \xi^N + \delta^N\} \\
 &\Downarrow \alpha\text{-Contamination} \\
 &\{\xi^1 + \delta^1, \xi'^2, \dots, \xi^N + \delta^N\} \quad \widehat{\mathbb{P}}_N
 \end{aligned}$$



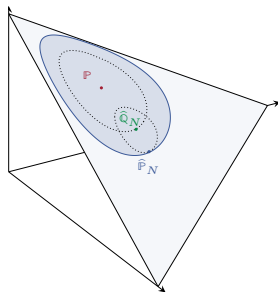
Prop: If $\widehat{\mathbb{P}}_N$ and $\widehat{\mathbb{Q}}_N$ differ by ϵ -norm bounded noise and α -contamination. Then

$$\text{LP}_\epsilon(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N) := \inf_{T \in \mathcal{T}(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N)} \int \mathbb{1} \{ \|\xi - \xi'\| \geq \epsilon \} dT \leq \alpha.$$

Distributions with marginals $\widehat{\mathbb{P}}_N$ and $\widehat{\mathbb{Q}}_N$

Lévy-Prokhorov Distance

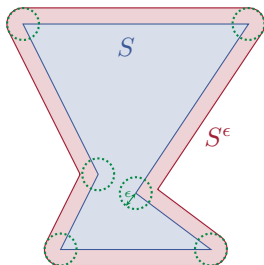
$$\begin{aligned} \xi &\sim \mathbb{P} \\ &\Downarrow \text{Statistical Error} \\ \{\xi^1, \xi^2, \dots, \xi^N\} &\quad \widehat{\mathbb{Q}}_N \\ &\Downarrow \epsilon\text{-Noise} \\ \{\xi^1 + \delta^1, \xi^2 + \delta^2, \dots, \xi^N + \delta^N\} \\ &\Downarrow \alpha\text{-Contamination} \\ \{\xi^1 + \delta^1, \xi'^2, \dots, \xi^N + \delta^N\} &\quad \widehat{\mathbb{P}}_N \end{aligned}$$



Prop: If $\widehat{\mathbb{P}}_N$ and $\widehat{\mathbb{Q}}_N$ differ by ϵ -norm bounded noise and α -contamination. Then

$$\text{LP}_\epsilon(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N) := \inf_{T \in \mathcal{T}(\widehat{\mathbb{P}}_N, \widehat{\mathbb{Q}}_N)} \int \mathbb{1} \{ \|\xi - \xi'\| \geq \epsilon \} dT \leq \alpha.$$

Intermezzo : Lévy-Prokhorov Distance

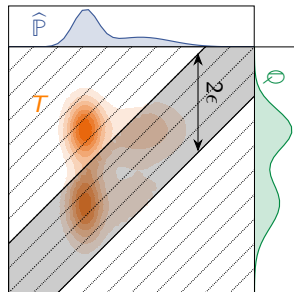
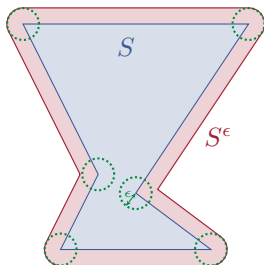


- **Prokhorov** in 1956: Metrize the topology of weak convergence on metric space using

$$\left\{ \mathbb{Q} : \text{LP}_\epsilon(\hat{\mathbb{P}}, \mathbb{Q}) < \alpha \right\} = \left\{ \mathbb{Q} : \mathbb{Q}(S) \leq \hat{\mathbb{P}}(S^\epsilon) + \alpha \quad \forall S \right\}$$

where $S^\epsilon := \{\xi' \in \Xi : \exists \xi \in S, \|\xi' - \xi\| \leq \epsilon\}$ and $\alpha = \epsilon$.

Intermezzo : Lévy-Prokhorov Distance



- **Prokhorov** in 1956: Metrize the topology of weak convergence on metric space using

$$\begin{aligned}\left\{Q : \text{LP}_\epsilon(\hat{\mathbb{P}}, Q) < \alpha\right\} &= \left\{Q : Q(S) \leq \hat{\mathbb{P}}(S^\epsilon) + \alpha \quad \forall S\right\} \\ &= \left\{Q : \exists T \in \mathcal{T}(\hat{\mathbb{P}}, Q), \int \mathbb{1}\{\|\xi - \xi'\| > \epsilon\} dT(\xi, \xi') < \alpha\right\}.\end{aligned}$$

where $S^\epsilon := \{\xi' \in \Xi : \exists \xi \in S, \|\xi' - \xi\| \leq \epsilon\}$ and $\alpha = \epsilon$.

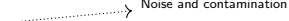
- **Strassen** in 1965: Established an optimal transport characterization

Holistic Robust Formulation

Holistic Robust (HR) Formulation:

$$\hat{c}^{HR}(\theta, \hat{\mathbb{P}}) = \max \mathbf{E}_{\mathbb{P}'} [\ell(\theta, \xi)]$$

$$\text{s.t. } \mathbb{Q}, \mathbb{P}' \in \mathcal{P},$$

$$\text{LP}_\epsilon(\hat{\mathbb{P}}, \mathbb{Q}) \leq \alpha,$$


Noise and contamination

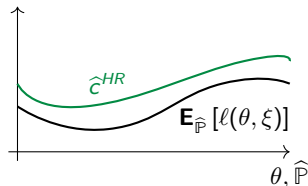
$$\text{KL}(\mathbb{Q}, \mathbb{P}') \leq r.$$


Statistical error

Holistic Robust Formulation

Holistic Robust (HR) Formulation:

$$\begin{aligned}\hat{c}^{HR}(\theta, \hat{\mathbb{P}}) &= \max \mathbf{E}_{\mathbb{P}'} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q}, \mathbb{P}' &\in \mathcal{P}, \\ \text{LP}_\epsilon(\hat{\mathbb{P}}, \mathbb{Q}) &\leq \alpha, \\ \text{KL}(\mathbb{Q}, \mathbb{P}') &\leq r.\end{aligned}$$



Thm (Bennouna and Van Parys, 2022): Consider ϵ -norm bounded noise and α -contamination. Then,

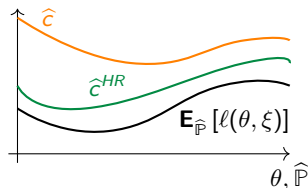
► (OOS Guarantee) we have

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}^{HR}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Holistic Robust Formulation

Holistic Robust (HR) Formulation:

$$\begin{aligned}\hat{c}^{HR}(\theta, \hat{\mathbb{P}}) &= \max \mathbf{E}_{\mathbb{P}'} [\ell(\theta, \xi)] \\ \text{s.t. } \mathbb{Q}, \mathbb{P}' &\in \mathcal{P}, \\ \text{LP}_\epsilon(\hat{\mathbb{P}}, \mathbb{Q}) &\leq \alpha, \\ \text{KL}(\mathbb{Q}, \mathbb{P}') &\leq r.\end{aligned}$$



Thm (Bennouna and Van Parys, 2022): Consider ϵ -norm bounded noise and α -contamination. Then,

- (OOS Guarantee) we have

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}^{HR}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

- (Efficiency) Furthermore, let $\hat{c}$ satisfy

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-rN + o(N)) \quad \forall \mathbb{P}.$$

Then, $\hat{c}(\theta, \hat{\mathbb{P}}) \geq \hat{c}^{HR}(\theta, \hat{\mathbb{P}})$ for all θ and $\hat{\mathbb{P}}$.

Even for $r = 0$, $\alpha = \epsilon > 0$ no exact algorithm was known (Erdoğ̃an and Iyengar, 2006).

Tractability

Even for $r = 0$, $\alpha = \epsilon > 0$ no exact algorithm was known (Erdoğlan and Iyengar, 2006).

HR Robust formulation is as tractable as Wasserstein robust formulation.

Even for $r = 0$, $\alpha = \epsilon > 0$ no exact algorithm was known (Erdoğlan and Iyengar, 2006).

HR Robust formulation is as tractable as Wasserstein robust formulation.

Thm: We have

$$\begin{aligned}\hat{c}^{HR}(\theta, \hat{\mathbb{P}}) &= \max \sum_{i=1}^N p'_i \ell^\epsilon(\theta, \xi^i) + p'_0 \ell^\infty(\theta) \\ \text{s.t. } &p' \in \mathbb{R}_+^{N+1}, q \in \mathbb{R}_+^{N+1}, s \in \mathbb{R}_+^N, \\ &\sum_{i=0}^N p'_i = \sum_{i=0}^N q_i = 1, q_i + s_i = \hat{\mathbb{P}}(\xi^i) \quad \forall i = 1, \dots, N, \\ &\sum_{i=0}^N q_i \log\left(\frac{q_i}{p'_i}\right) \leq r, \sum_{i=1}^N s_i \leq \alpha.\end{aligned}$$

with

- ▶ Inflated loss $\ell^\epsilon(\theta, \xi) = \max_{\|\xi' - \xi\| \leq \epsilon, \xi' \in \Xi} \ell(\theta, \xi')$.
- ▶ Worst-case loss $\ell^\infty(\theta) = \max_{\xi' \in \Xi} \ell(\theta, \xi')$

Table of Contents

1. The Dawn of DRO
2. Efficient Decisions from Data
3. Holistic Robustness
- 4. Robust Statistics**
5. Conclusions

Huber's TV Contamination

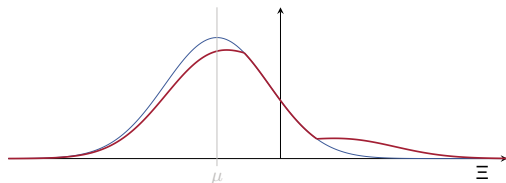
Contamination

Small fraction $\alpha \in [0, \frac{1}{2})$ of probability mass can be shifted

$$\xi \sim \mathbb{G}_\mu$$



$$\xi \sim \mathbb{P}^c$$



$$\text{TV}(\mathbb{G}_\mu, \mathbb{P}^c) = \frac{1}{2} \int |\mathbf{g}_\mu(\xi) - \mathbf{p}^c(\xi)| d\xi \leq \alpha$$

Statistical Error

Randomness of the sampled data

$$\xi \sim \mathbb{P}^c$$



$$\{\xi^1, \dots, \xi^N\}$$

Location Estimation (Huber, 1964)

Let $\{\xi^i \sim \mathbb{P}^c\}_{i=1}^N$.

- Classical location estimate is

$$\hat{\mu}^{\text{avg}}(\hat{\mathbb{P}}_N) = \sum_{i=1}^N \frac{1}{n} \cdot \xi^i$$

- Very sensitive to corruption, i.e., mapping $\hat{\mathbb{P}} \rightarrow \mu^{\text{avg}}(\hat{\mathbb{P}})$ is discontinuous.

Location Estimation (Huber, 1964)

Let $\{\xi^i \sim \mathbb{P}^c\}_{i=1}^N$.

- Classical location estimate is

$$\hat{\mu}^{\text{avg}}(\hat{\mathbb{P}}_N) = \sum_{i=1}^N \frac{1}{n} \cdot \xi^i \iff \sum_{i=1}^N \frac{1}{n} \cdot (\xi^i - \hat{\mu}^{\text{avg}}(\hat{\mathbb{P}}_N)) = 0.$$

- Very sensitive to corruption, i.e., mapping $\hat{\mathbb{P}} \rightarrow \mu^{\text{avg}}(\hat{\mathbb{P}})$ is discontinuous.

Location Estimation (Huber, 1964)

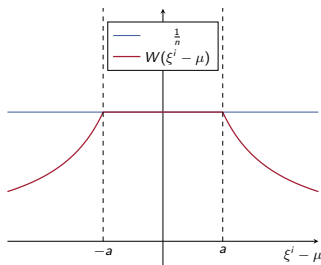
Let $\{\xi^i \sim \mathbb{P}^c\}_{i=1}^N$.

- Classical location estimate is

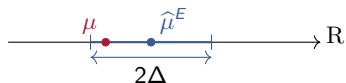
$$\hat{\mu}^{\text{avg}}(\hat{\mathbb{P}}_N) = \sum_{i=1}^N \frac{1}{n} \cdot \xi^i \iff \sum_{i=1}^N \frac{1}{n} \cdot (\xi^i - \hat{\mu}^{\text{avg}}(\hat{\mathbb{P}}_N)) = 0.$$

- Very sensitive to corruption, i.e., mapping $\hat{\mathbb{P}} \rightarrow \mu^{\text{avg}}(\hat{\mathbb{P}})$ is discontinuous.
- Huber's proposes instead

$$\sum_{i=1}^N W(\xi^i - \hat{\mu}^{\text{hub}}(\hat{\mathbb{P}}_N)) \cdot (\xi^i - \hat{\mu}^{\text{hub}}(\hat{\mathbb{P}}_N)) = 0.$$



Error Margin

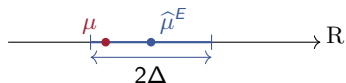


Take the **error margin** $\Delta(\hat{\mu}^E)$ of an estimator $\hat{\mu}^E$ to be the smallest $\Delta \geq 0$ for which

$$\text{Prob}_{\mathbb{P}^c} [|\hat{\mu}^E(\hat{\mathbb{P}}_N) - \mu| \leq \Delta] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ such that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$.

Error Margin



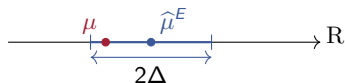
Take the **error margin** $\Delta(\hat{\mu}^E)$ of an estimator $\hat{\mu}^E$ to be the smallest $\Delta \geq 0$ for which

$$\text{Prob}_{\mathbb{P}^c} [|\hat{\mu}^E(\hat{\mathbb{P}}_N) - \mu| \leq \Delta] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ such that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$.

Remark: For $\alpha > 0$ we get $+\infty = \Delta(\hat{\mu}^{avg}) > \Delta(\hat{\mu}^{median})$.

Error Margin



Take the **error margin** $\Delta(\hat{\mu}^E)$ of an estimator $\hat{\mu}^E$ to be the smallest $\Delta \geq 0$ for which

$$\text{Prob}_{\mathbb{P}^c} [|\hat{\mu}^E(\hat{\mathbb{P}}_N) - \mu| \leq \Delta] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ such that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$.

Remark: For $\alpha > 0$ we get $+\infty = \Delta(\hat{\mu}^{avg}) > \Delta(\hat{\mu}^{median})$.

Thm: (Huber, 1964) We have that $\Delta(\hat{\mu}^E) \geq \Delta(\hat{\mu}^{hub})$ for all estimators $\hat{\mu}^E$.

► Huber's estimator yields “minimal” confidence **intervals**.

Our holistic robust perspective suggests the confidence **set**

$$\hat{S}^{hr}(\hat{\mathbb{P}}_N) = \left\{ \mu \in \mathbb{R} : \begin{array}{l} \exists \mathbb{P}^c \in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}_N, \mathbb{P}^c) \leq r, \\ \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha \end{array} \right\}.$$

Confidence Set Estimators



Our holistic robust perspective suggests the confidence **set**

$$\hat{S}^{hr}(\hat{\mathbb{P}}_N) = \left\{ \mu \in \mathbb{R} : \begin{array}{l} \exists \mathbb{P}^c \in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}_N, \mathbb{P}^c) \leq r, \\ \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha \end{array} \right\}.$$

Confidence Set Estimators



Our holistic robust perspective suggests the confidence **set**

$$\hat{S}^{hr}(\hat{\mathbb{P}}_N) = \left\{ \mu \in \mathbb{R} : \begin{array}{l} \exists \mathbb{P}^c \in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}_N, \mathbb{P}^c) \leq r, \\ \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha \end{array} \right\}.$$

Thm (Informal):

► (OOS Guarantee) We have

$$\text{Prob}_{\mathbb{P}^c}[\mu \in \hat{S}^{hr}(\hat{\mathbb{P}}_N)] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ so that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$.

Confidence Set Estimators



Our holistic robust perspective suggests the confidence **set**

$$\hat{S}^{hr}(\hat{\mathbb{P}}_N) = \left\{ \mu \in \mathbb{R} : \begin{array}{l} \exists \mathbb{P}^c \in \mathcal{P}, \\ \text{KL}(\hat{\mathbb{P}}_N, \mathbb{P}^c) \leq r, \\ \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha \end{array} \right\}.$$

Thm (Informal):

- (OOS Guarantee) We have

$$\text{Prob}_{\mathbb{P}^c}[\mu \in \hat{S}^{hr}(\hat{\mathbb{P}}_N)] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ so that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$.

- (Efficiency) Furthermore, let $\hat{S}$ satisfy

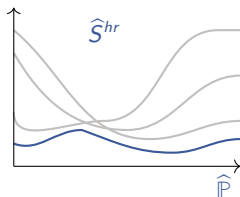
$$\text{Prob}_{\mathbb{P}^c}[\mu \in \hat{S}(\hat{\mathbb{P}}_N)] \geq 1 - \exp(-rN + o(N))$$

for all $\mu \in \mathbb{R}$, $\mathbb{P}^c \in \mathcal{P}$ so that $\text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha$. Then, $\hat{S}^{hr}(\hat{\mathbb{P}}) \subseteq \hat{S}(\hat{\mathbb{P}})$ for all $\hat{\mathbb{P}}$.

Meta Optimization Problem

Among all confidence sets with **out-of-sample guarantee** find the **smallest** one

$$\left\{ \begin{array}{l} \min_{\hat{S}} [\subseteq] \hat{S} \\ \text{s.t. } \text{Prob}_{\mathbb{P}^c}[\mu \in \hat{S}(\hat{\mathbb{P}}_n)] \geq 1 - \exp(-rN + o(N)) \\ \forall \mu \in \mathbb{R}, \mathbb{P}^c \in \mathcal{P} \text{ st } \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha. \end{array} \right.$$

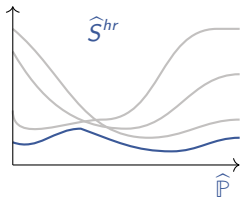


Pareto dominant solution

Meta Optimization Problem

Among all confidence sets with **out-of-sample guarantee** find the **smallest** one

$$\left\{ \begin{array}{l} \min_{\hat{S}} [\subseteq] \hat{S} \\ \text{s.t. } \text{Prob}_{\mathbb{P}^c} [\mu \in \hat{S}(\hat{\mathbb{P}}_n)] \geq 1 - \exp(-rN + o(N)) \\ \forall \mu \in \mathbb{R}, \mathbb{P}^c \in \mathcal{P} \text{ st } \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha. \end{array} \right.$$

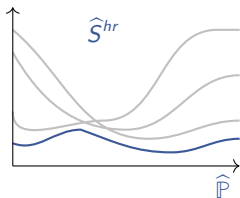


Pareto dominant solution

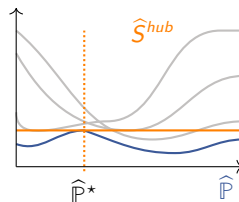
Meta Optimization Problem

Among all confidence sets with **out-of-sample guarantee** find the **smallest** one

$$\begin{cases} \min_{\hat{\mathcal{S}}} [\subseteq] \hat{\mathcal{S}} \\ \text{s.t. } \text{Prob}_{\mathbb{P}^c}[\mu \in \hat{\mathcal{S}}(\hat{\mathbb{P}}_n)] \geq 1 - \exp(-rN + o(N)) \\ \forall \mu \in \mathbb{R}, \mathbb{P}^c \in \mathcal{P} \text{ st } \text{TV}(\mathbb{P}^c, \mathbb{G}_\mu) \leq \alpha. \end{cases}$$



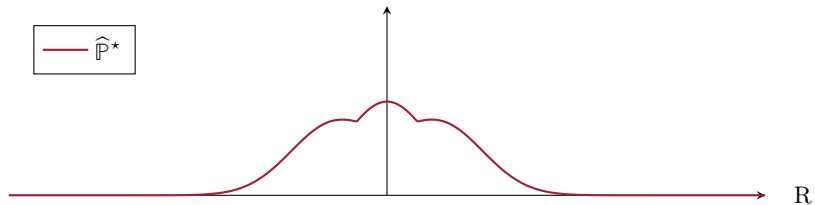
Pareto dominant solution



Minimax solution

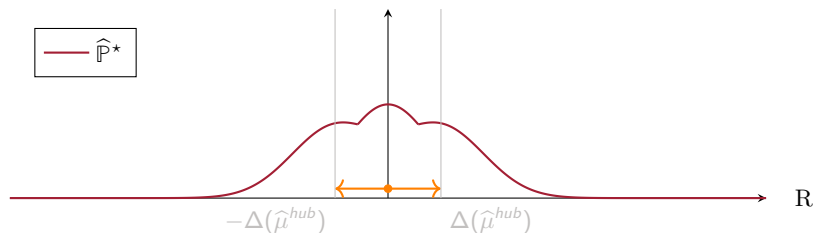
Thm: There exists $\hat{\mathbb{P}}^*$ so that we have $\text{diam}(\hat{\mathcal{S}}^{hr}(\hat{\mathbb{P}}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Minimax Solution



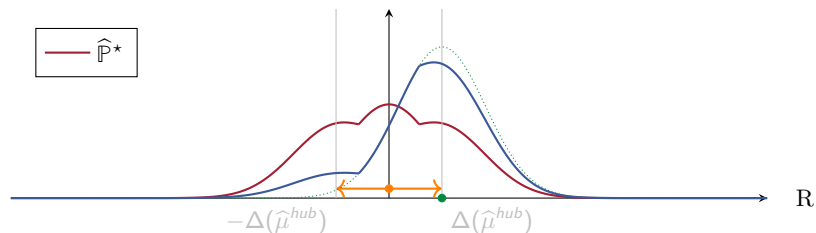
Thm: There exists $\hat{P}^*$ so that we have $\text{diam}(\hat{S}^{hr}(\hat{P}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Minimax Solution



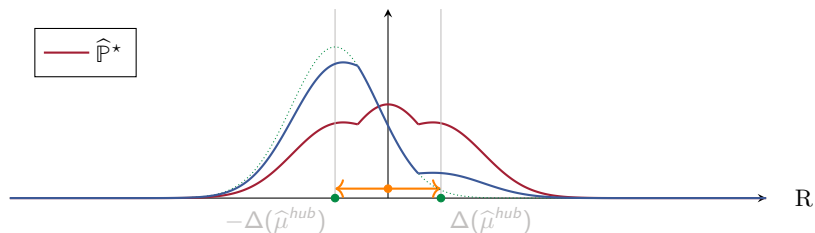
Thm: There exists $\hat{P}^*$ so that we have $\text{diam}(\hat{S}^{hr}(\hat{P}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Minimax Solution



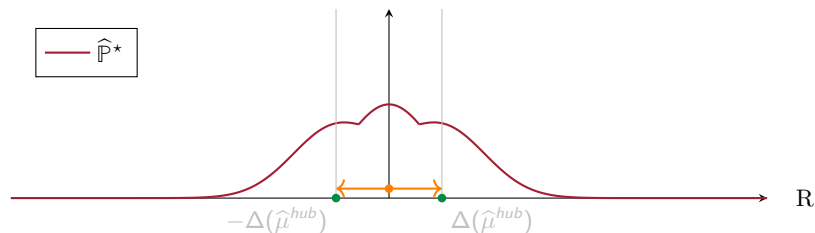
Thm: There exists $\hat{P}^*$ so that we have $\text{diam}(\hat{S}^{hr}(\hat{P}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Minimax Solution



Thm: There exists $\hat{P}^*$ so that we have $\text{diam}(\hat{S}^{hr}(\hat{P}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Minimax Solution



Thm: There exists $\hat{P}^*$ so that we have $\text{diam}(\hat{S}^{hr}(\hat{P}^*)) = 2\Delta(\hat{\mu}^{hub})$.

Remarks

- ▶ We may have $\hat{S}^{hr}(\hat{P}) \subset \mu^{hub}(\hat{P}) + [-\Delta(\hat{\mu}^{hub}), \Delta(\hat{\mu}^{hub})]$ for $\hat{P} \neq \hat{P}^*$.
- ▶ However, the estimator $\hat{S}^{hr}$ may be hard to evaluate.

Table of Contents

1. The Dawn of DRO
2. Efficient Decisions from Data
3. Holistic Robustness
4. Robust Statistics
- 5. Conclusions**

1. Robust optimization can make data-driven decisions efficiently
2. Holistic robust optimization additionally protects decisions against corruption
3. Parametric DRO is intimately related to classical robust statistics

-  Bennouna, Amine and Bart Van Parys (2021). "Learning and decision-making with data: Optimal formulations and phase transitions". In: **arXiv preprint arXiv:2109.06911**.
-  Bennouna, Amine and Bart Van Parys (2022). "Holistic Robust Data-Driven Decisions". In: **arXiv preprint arXiv:2207.09560**.
-  Bertsimas, Dimitris and Ioana Popescu (2005). "Optimal inequalities in probability theory: A convex optimization approach". In: **SIAM Journal on Optimization** 15.3, pp. 780–804.
-  Birge, John R and Francois Louveau (2011). **Introduction to stochastic programming**. Springer Science & Business Media.
-  Dharmadhikari, Sudhakar and Kumar Joag-Dev (1988). **Unimodality, convexity, and applications**. Elsevier.
-  Erdoğan, Emre and Garud Iyengar (2006). "Ambiguous chance constrained problems and robust optimization". In: **Mathematical Programming** 107, pp. 37–61.
-  Fournier, Nicolas and Arnaud Guillin (2015). "On the rate of convergence in Wasserstein distance of the empirical measure". In: **Probability theory and related fields** 162.3-4, pp. 707–738.

References II

-  Huber, Peter J. (1964). “Robust estimation of a location parameter”. In: **Ann. Math. Statist.** 35, pp. 73–101.
-  Michaud, Richard O. (1989). “The Markowitz optimization enigma: Is “optimized” optimal?” In: **Financial analysts journal** 45.1, pp. 31–42.
-  Roald, Line, Frauke Oldewurtel, Bart Van Parys, and Göran Andersson (2015). “Security constrained optimal power flow with distributionally robust chance constraints”. In: **arXiv preprint arXiv:1508.06061**.
-  Sanov, Ivan N (1958). **On the probability of large deviations of random variables**. United States Air Force, Office of Scientific Research.
-  Strassen, Volker (1965). “The existence of probability measures with given marginals”. In: **The Annals of Mathematical Statistics** 36.2, pp. 423–439.
-  Sutter, Tobias, Bart Van Parys, and Daniel Kuhn (2020). “A general framework for optimal data-driven optimization”. In: **arXiv preprint arXiv:2010.06606**.
-  Van Parys, Bart, P. Esfahani Mohajerin Esfahani, and D. Kuhn (2020). “From data to decisions: Distributionally robust optimization is optimal”. In: **Management Science**.
-  Van Parys, Bart, Paul Goulart, and Paul Embrechts (2016). “Fréchet inequalities via convex optimization”. In: **Available on Optimization Online**.

-  Van Parys, Bart, Paul Goulart, and Manfred Morari (2019). "Distributionally robust expectation inequalities for structured distributions". In: **Mathematical Programming** 173, pp. 251–280.
-  Van Parys, Bart, Paul J. Goulart, and Daniel Kuhn (2016). "Generalized Gauss inequalities via semidefinite programming". In: **Mathematical Programming** 156.1-2, pp. 271–302.
-  Van Parys, Bart, Daniel Kuhn, Paul J Goulart, and Manfred Morari (2015). "Distributionally robust control of constrained stochastic systems". In: **IEEE Transactions on Automatic Control** 61.2, pp. 430–442.
-  Van Parys, Bart, Bing Feng Ng, Paul Goulart, and Rafael Palacios (2014). "Optimal control for load alleviation in wind turbines". In: **32nd ASME Wind Energy Symposium**, p. 1222.

6. Extreme Unimodal Distributions

7. α -Unimodal Distributions

8. Meta-Optimization Problem

9. Prescriptors without Predictors

10. Location Estimation

Extreme Unimodal Distributions

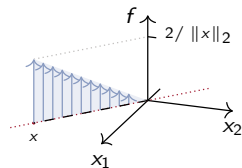
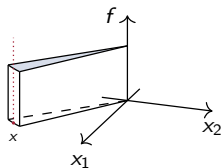
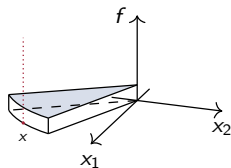


Table of Contents

6. Extreme Unimodal Distributions

7. α -Unimodal Distributions

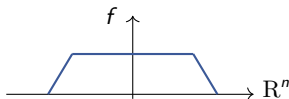
8. Meta-Optimization Problem

9. Prescriptors without Predictors

10. Location Estimation

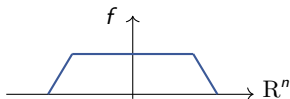
Structured Moment Sets $\mathcal{U}_\alpha(\mu, S) := \{Q \in \mathcal{S}_\alpha : \int \xi dQ = \mu, \int \xi \xi^\top dQ = S\}$

- Reduce pessimism by imposing additional structure



Structured Moment Sets $\mathcal{U}_\alpha(\mu, S) := \{Q \in \mathcal{S}_\alpha : \int \xi dQ = \mu, \int \xi \xi^\top dQ = S\}$

- Reduce pessimism by imposing additional structure



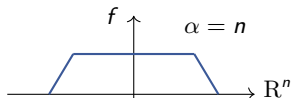
- **Definition:** If $\mathbb{P}$ has a continuous density function f , then $\mathbb{P}$ is **unimodal** if

$$t \mapsto f(tx)$$

is non-increasing for any $x \in \mathbb{R}^n$ (Dharmadhikari and Joag-Dev, 1988).

Structured Moment Sets $\mathcal{U}_\alpha(\mu, S) := \{Q \in \mathcal{S}_\alpha : \int \xi dQ = \mu, \int \xi \xi^\top dQ = S\}$

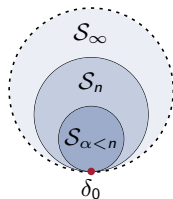
- Reduce pessimism by imposing additional structure



- **Definition:** If $\mathbb{P}$ has a continuous density function f , then $\mathbb{P}$ is α -**unimodal** if

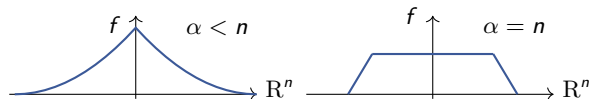
$$t \mapsto f(tx)/t^{\alpha-n}$$

is non-increasing for any $x \in \mathbb{R}^n$ (Dharmadhikari and Joag-Dev, 1988).



Structured Moment Sets $\mathcal{U}_\alpha(\mu, S) := \{Q \in \mathcal{S}_\alpha : \int \xi dQ = \mu, \int \xi \xi^\top dQ = S\}$

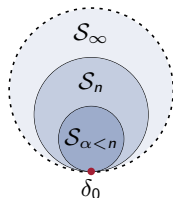
- Reduce pessimism by imposing additional structure



- **Definition:** If $\mathbb{P}$ has a continuous density function f , then $\mathbb{P}$ is α -**unimodal** if

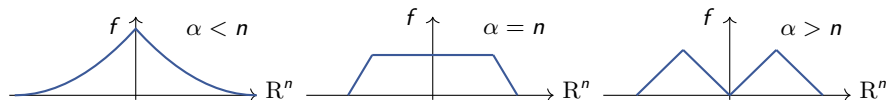
$$t \mapsto f(tx)/t^{\alpha-n}$$

is non-increasing for any $x \in \mathbb{R}^n$ (Dharmadhikari and Joag-Dev, 1988).



Structured Moment Sets $\mathcal{U}_\alpha(\mu, S) := \{Q \in \mathcal{S}_\alpha : \int \xi dQ = \mu, \int \xi \xi^\top dQ = S\}$

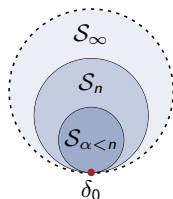
- Reduce pessimism by imposing additional structure



- **Definition:** If $\mathbb{P}$ has a continuous density function f , then $\mathbb{P}$ is α -**unimodal** if

$$t \mapsto f(tx)/t^{\alpha-n}$$

is non-increasing for any $x \in \mathbb{R}^n$ (Dharmadhikari and Joag-Dev, 1988).



Thm: Let $P := \{\xi \in \mathbb{R}^n : a_k^\top \xi - b_k < 0 \ \forall k = 1, \dots, K\}$ be a polytope. Then,

$$\max \mathbb{Q}(\xi \notin P)$$

$$\text{s.t. } \mathbb{Q} \in \mathcal{U}_\alpha(\hat{\mu}_N, \hat{S}_N)$$

is a convex **tractable optimization** problem.

Structured Moment Problems

Thm: Let $P := \{\xi \in \mathbb{R}^n : \mathbf{a}_k^\top \xi - b_k < 0 \ \forall k = 1, \dots, K\}$ be a polytope. Then,

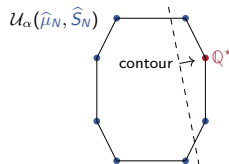
$$\begin{aligned} \max \mathbb{Q}(\xi \notin P) &= \max \sum_{k=1}^K \lambda_k \left(1 - (\mathbf{b}_k \lambda_k / \mathbf{a}_k^\top \mathbf{z}_k)^\alpha\right) \\ \text{s.t. } \mathbb{Q} \in \mathcal{U}_\alpha(\hat{\mu}_N, \hat{S}_N) &\quad \text{s.t. } \sum_{k=1}^K \begin{pmatrix} \mathbf{z}_k & \mathbf{z}_k \\ \mathbf{z}_k^\top & \lambda_k \end{pmatrix} \preceq \begin{pmatrix} (\alpha+2)\hat{S}_n/\alpha & (\alpha+1)\hat{\mu}_n/\alpha \\ (\alpha+1)\hat{\mu}_n^\top/\alpha & 1 \end{pmatrix} \\ &\quad \begin{pmatrix} \mathbf{z}_k & \mathbf{z}_k \\ \mathbf{z}_k^\top & \lambda_k \end{pmatrix} \succeq 0 \quad \forall k = 1, \dots, K \end{aligned}$$

is a convex **tractable optimization** problem.

Proof Sketch

- We only need to optimize over **extreme distributions**

$$\begin{aligned} \max \mathbb{Q}(\xi \in P) &= \max \mathbb{Q}(\xi \in P) \\ \text{s.t. } \mathbb{Q} \in \mathcal{U}_\alpha(\hat{\mu}_N, \hat{S}_N) &\quad \text{s.t. } \mathbb{Q} \in \text{ex}\mathcal{U}_\alpha(\hat{\mu}_N, \hat{S}_N) \end{aligned}$$



- The set $\mathcal{S}_\alpha$ is a **Choquet Simplex** extreme points $\{e_x : x \in \mathbb{R}^n\}$.
- Via Caratheodory's theorem we can establish that in fact

$$Q^* = \sum_{k=0}^K p_k e_{x_k} \text{ with } a_k^\top x_k \geq b_k \quad \forall k = 1, \dots, K.$$

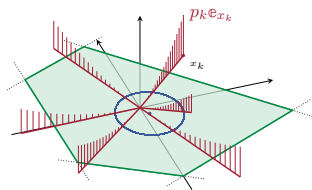


Table of Contents

6. Extreme Unimodal Distributions

7. α -Unimodal Distributions

8. Meta-Optimization Problem

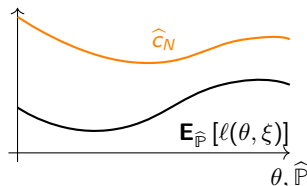
9. Prescriptors without Predictors

10. Location Estimation

Meta Optimization Problem

Data driven formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \hat{c}_N(\theta, \hat{\mathbb{P}}_N)$$



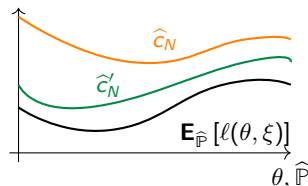
Feasibility: Consider formulations so that

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P}.$$

Meta Optimization Problem

Data driven formulations

$$\hat{\theta}_N \in \arg \min_{\theta \in \Theta} \hat{c}_N(\theta, \hat{\mathbb{P}}_N)$$



Feasibility: Consider formulations so that

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}}[\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P}.$$

Conservatism: Formulation $\tilde{c}$ is less conservative than $\hat{c}$ if and only if

$$\tilde{c} \preceq \hat{c} \iff \lim_{N \rightarrow \infty} \frac{|\tilde{c}_N(\theta, \hat{\mathbb{P}}) - \mathbf{E}_{\hat{\mathbb{P}}}[\ell(\theta, \xi)]|}{|\hat{c}_N(\theta, \hat{\mathbb{P}}) - \mathbf{E}_{\hat{\mathbb{P}}}[\ell(\theta, \xi)]|} \leq 1 \quad \forall \theta, \hat{\mathbb{P}}.$$

Meta Optimization Problem

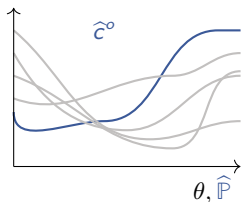
Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P} \end{cases}$$

Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P} \end{cases}$$

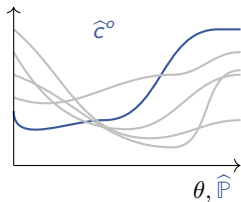


Pareto optimal solution

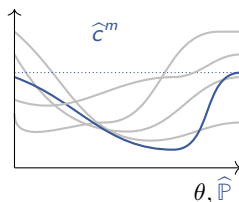
Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P} \end{cases}$$



Pareto optimal solution

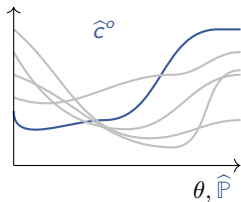


Minimax solution

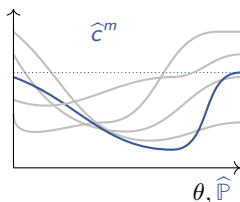
Meta Optimization Problem

Among all predictors with **out-of-sample guarantee** find the **least conservative** one

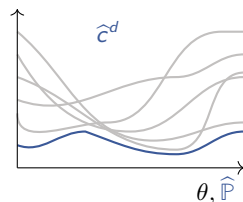
$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P} \end{cases}$$



Pareto optimal solution



Minimax solution



Pareto dominant solution

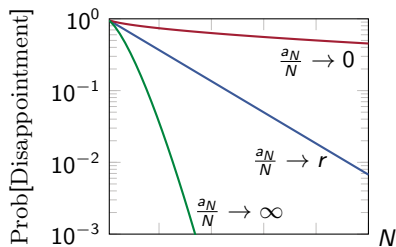
Pareto Dominant Solution

Thm: The meta optimization problem

$$\begin{cases} \min_{\hat{c}} [\preceq] \hat{c} \\ \text{s.t. } \text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P} \end{cases}$$

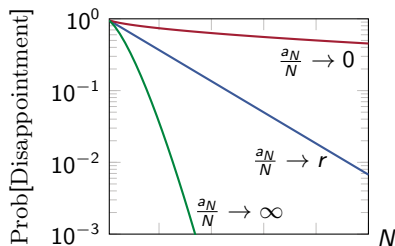
admits a **Pareto dominant** solution. We observe three regimes depending on **rate** a_N .

Phase Transitions



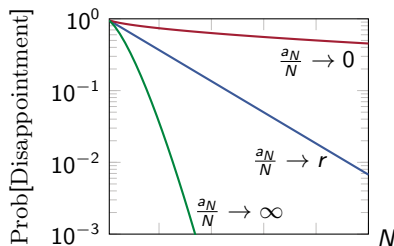
$a_N \gg N$	$a_N \sim rN$	$a_N \ll N$

Phase Transitions



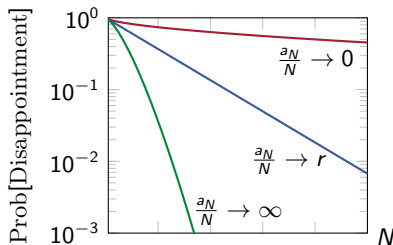
$a_N \gg N$	$a_N \sim rN$	$a_N \ll N$
	DRO: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$	
	Inconsistent	

Phase Transitions



$a_N \gg N$	$a_N \sim rN$	$a_N \ll N$
RO: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_n, \mathbb{Q}) \leq \infty} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$ $= \max_{\xi \in \Xi} \ell(\theta, \xi)$	DRO: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$	
Inconsistent	Inconsistent	

Phase Transitions



$a_N \gg N$	$a_N \sim rN$	$a_N \ll N$
RO: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq \infty} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$ $= \max_{\xi \in \Xi} \ell(\theta, \xi)$	DRO: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq r} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$	VR: $\hat{c}_N^d(\theta, \hat{\mathbb{P}}_N)$ $= \max_{KL(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq \frac{a_N}{N}} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]$ $\equiv \mathbf{E}_{\hat{\mathbb{P}}_N} [\ell(\theta, \xi)] + \sqrt{\frac{2a_N}{N}} \mathbf{V}_{\hat{\mathbb{P}}_N} [\ell(x, \xi)]$
Inconsistent	Inconsistent	Consistent

Table of Contents

6. Extreme Unimodal Distributions

7. α -Unimodal Distributions

8. Meta-Optimization Problem

9. Prescriptors without Predictors

10. Location Estimation

Not all data driven **prescriptors** are of the form

$$\hat{\theta}_N(\hat{\mathbb{P}}_N) \in \arg \min_{\theta \in \Theta} \hat{c}_N(\theta, \hat{\mathbb{P}}_N).$$

Prescriptors without Predictors

Not all data driven **prescriptors** are of the form

$$\hat{\theta}_N(\hat{\mathbb{P}}_N) \in \arg \min_{\theta \in \Theta} \hat{c}_N(\theta, \hat{\mathbb{P}}_N).$$

Assumption: $\theta \mapsto \ell(\theta, \xi)$ is L -Lipschitz. Consider a **prescriptor** $\hat{\theta}_N : \mathcal{P}(\Xi) \rightarrow \Theta$ which enjoys some **out-of-sample guarantee**

$$\text{Prob}_{\mathbb{P}} \left[\hat{g}_N(\hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} \left[\ell(\hat{\theta}_N(\hat{\mathbb{P}}_N), \xi) \right] \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P}.$$

for some function $\hat{g}_N$.

Reduction

Any prescriptor $\hat{\theta}_N(\hat{\mathbb{P}}_N)$ can be **represented** as minimizing the proxy cost function

$$\hat{c}_N(\theta, \hat{\mathbb{P}}_N) := \hat{g}_N(\hat{\mathbb{P}}_N) + L\|\theta - \hat{\theta}_N(\hat{\mathbb{P}}_N)\|$$

Furthermore, observe that

$$\hat{g}_N(\hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\hat{\theta}_N(\hat{\mathbb{P}}_N), \xi)] \implies \hat{g}_N(\hat{\mathbb{P}}_N) + L\|\theta - \hat{\theta}_N(\hat{\mathbb{P}}_N)\| \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta$$

and consequently

$$\text{Prob}_{\mathbb{P}} \left[\hat{c}_N(\theta, \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} [\ell(\theta, \xi)] \quad \forall \theta \in \Theta \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P}.$$

Dominance

Recall the **KL prescriptor**

$$\hat{\theta}_N^{\text{KL}} \in \arg \min_{\theta \in \Theta} \{c_N^{\text{KL}}(\theta, \hat{\mathbb{P}}_N) := \max_{D_{\text{KL}}(\hat{\mathbb{P}}_N, \mathbb{Q}) \leq a_N/N} \mathbf{E}_{\mathbb{Q}} [\ell(\theta, \xi)]\}.$$

From the **Pareto dominance** of $\hat{c}_N^d$ we have $\forall \theta \in \Theta$ and $\forall \hat{\mathbb{P}}$ that $\hat{c}_N(\theta, \hat{\mathbb{P}}) \geq \hat{c}_N^{\text{KL}}(\theta, \hat{\mathbb{P}})$ and in particular

$$\hat{g}_N(\hat{\mathbb{P}}_N) = \hat{c}_N(\hat{\theta}_N(\hat{\mathbb{P}}_N), \hat{\mathbb{P}}_N) \geq \hat{c}_N^{\text{KL}}(\hat{\theta}_N(\hat{\mathbb{P}}_N), \hat{\mathbb{P}}_N) \geq \hat{c}_N^{\text{KL}}(\hat{\theta}_{\text{KL}}(\hat{\mathbb{P}}_N), \hat{\mathbb{P}}_N)$$

The KL prescriptor $\hat{\theta}_N^{\text{KL}}$ also dominates $\hat{\theta}_N$. That is,

$$\text{Prob}_{\mathbb{P}} \left[\hat{g}_N(\hat{\mathbb{P}}_N) \geq \hat{c}_N^{\text{KL}}(\hat{\theta}_{\text{KL}}(\hat{\mathbb{P}}_N), \hat{\mathbb{P}}_N) \geq \mathbf{E}_{\mathbb{P}} \left[\ell(\hat{\theta}_{\text{KL}}(\hat{\mathbb{P}}_N), \xi) \right] \right] \geq 1 - \exp(-a_N + o(a_N)) \quad \forall \mathbb{P}.$$

Table of Contents

6. Extreme Unimodal Distributions

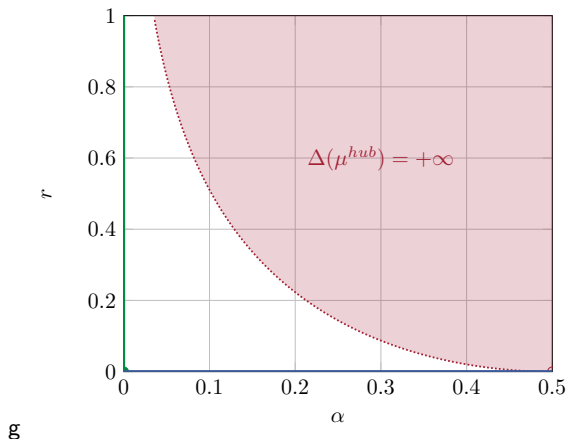
7. α -Unimodal Distributions

8. Meta-Optimization Problem

9. Prescriptors without Predictors

10. Location Estimation

Phase Diagram



- ▶ Green: Mean is optimal.
- ▶ Blue: Median is optimal.
- ▶ Red: Error margin is unbounded.