

Optimization over Trained Neural Networks: Difference-of-Convex Algorithm and Application to Data Center Scheduling

Xinwei Liu, Vladimir Dvorkin

Department of Electrical Engineering and Computer Science

University of Michigan – Ann Arbor

INFORMS 2025, Atlanta, GA

October 24, 2025

1. Motivation and Background
2. Mathematical Formulation
3. Application

$$\begin{array}{ll} \underset{\mathbf{x} \in \mathcal{X}}{\text{minimize}} & c(\mathbf{x}) \\ \text{subject to} & g(\mathbf{x}) \leq 0 \end{array}$$

- Usually, we know all elements $(c, g, \mathcal{X})$ of an optimization problem.
- **What should we do if the constraint function g is unknown?**
- In practice, we can approximate the function g from the data.

$$\begin{array}{ll} \underset{\mathbf{x} \in \mathcal{X}}{\text{minimize}} & c(\mathbf{x}) \\ \text{subject to} & g(\mathbf{x}) \leq 0 \end{array} \quad \Longrightarrow \quad \begin{array}{ll} \underset{\mathbf{x} \in \mathcal{X}}{\text{minimize}} & c(\mathbf{x}) \\ \text{subject to} & \text{NN}(\mathbf{x}) \leq 0 \end{array}$$

- Usually, we know all elements $(c, g, \mathcal{X})$ of an optimization problem.
- **What should we do if the constraint function g is unknown?**
- In practice, we can approximate the function g from the data.
- **Consider an approximation by a neural network $\text{NN} : \mathcal{X} \mapsto \mathbb{R}$**
- The problem becomes an optimization over trained neural network.

A neural network mapping $\mathbf{d}$ to λ is mathematically represented as

$$\lambda = \mathbf{W}_{N+1} \text{ReLU}(\dots \text{ReLU}(\mathbf{W}_1 \mathbf{d} + \mathbf{b}_1) \dots) + \mathbf{b}_{N+1}$$

with fixed weights and biases $\mathbf{W}$ and $\mathbf{b}$, and variable input $\mathbf{d}$ and output λ

[Turner et al., 2024]

A neural network mapping $\mathbf{d}$ to λ is mathematically represented as

$$\lambda = \mathbf{W}_{N+1} \text{ReLU}(\dots \text{ReLU}(\mathbf{W}_1 \mathbf{d} + \mathbf{b}_1) \dots) + \mathbf{b}_{N+1}$$

with fixed weights and biases $\mathbf{W}$ and $\mathbf{b}$, and variable input $\mathbf{d}$ and output λ

We can write the ReLU activation function using complementarity logic:

$$y = \text{ReLU}(a) \implies \begin{aligned} y &= a + v \\ 0 &\leq y \perp v \leq 0 \end{aligned}$$

a input variable
 y output variable
 v auxiliary variable
 $\perp$ complementarity

A neural network mapping $\mathbf{d}$ to λ is mathematically represented as

$$\lambda = \mathbf{W}_{N+1} \text{ReLU}(\dots \text{ReLU}(\mathbf{W}_1 \mathbf{d} + \mathbf{b}_1) \dots) + \mathbf{b}_{N+1}$$

with fixed weights and biases $\mathbf{W}$ and $\mathbf{b}$, and variable input $\mathbf{d}$ and output λ

We can write the ReLU activation function using complementarity logic:

$$y = \text{ReLU}(a) \implies \begin{aligned} y &= a + v \\ 0 &\leq y \perp v \leq 0 \end{aligned}$$

a input variable
 y output variable
 v auxiliary variable
 $\perp$ complementarity

Extremely difficult problem to solve in practice due to large computational burden

[Turner et al., 2024]

State of the art approaches

- Mix-integer constraint[Tjeng et al., 2017, Grimstad and Andersson, 2019]
- Big-M reformulation of ReLU—numerical issue
[Tjeng et al., 2017, Grimstad and Andersson, 2019]
- SOS1 reformulation with decision trees[Turner et al., 2024]
- Semidefinite relaxation of ReLU [Dathathri et al., 2020, Fazlyab et al., 2020]
- Bound propagation [Wang et al., 2021]

Our contribution

- Use difference of convex approach to optimize over trained neural network.
- Relax the ReLU constraints by penalizing them in the objective function.
- Avoid the computational complexity of standard MIP formulations.
- Formulate an algorithm to compute the hyperparameter to achieve fast convergence.
- Demonstrate in the data center load allocation example.

$$\begin{aligned} & \underset{\mathbf{d} \in \mathcal{D}, \mathbf{y}, \boldsymbol{\lambda}}{\text{minimize}} \quad c(\boldsymbol{\lambda}) \quad \text{convex function} \\ & \text{subject to} \quad \mathbf{y}_1 = \text{ReLU}(\mathbf{W}_1 \mathbf{d} + \mathbf{b}_1), \\ & \quad \mathbf{y}_i = \text{ReLU}(\mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i), \\ & \quad \forall i = 2, \dots, N \\ & \quad \boldsymbol{\lambda} = \mathbf{W}_{N+1} \mathbf{y}_N + \mathbf{b}_{N+1}, \end{aligned}$$

$$\begin{aligned} & \underset{\mathbf{d} \in \mathcal{D}, \mathbf{y}, \boldsymbol{\lambda}}{\text{minimize}} \quad c(\boldsymbol{\lambda}) \quad \text{convex function} \\ & \text{subject to} \quad \mathbf{y}_1 = \text{ReLU}(\mathbf{W}_1 \mathbf{d} + \mathbf{b}_1), \\ & \quad \mathbf{y}_i = \text{ReLU}(\mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i), \\ & \quad \forall i = 2, \dots, N \\ & \quad \boldsymbol{\lambda} = \mathbf{W}_{N+1} \mathbf{y}_N + \mathbf{b}_{N+1}, \end{aligned} \quad \Rightarrow$$

$$\begin{aligned} & \underset{\mathbf{d} \in \mathcal{D}, \mathbf{y}, \mathbf{v}, \boldsymbol{\lambda}}{\text{minimize}} \quad c(\boldsymbol{\lambda}) \\ & \text{subject to} \quad \mathbf{y}_1 = \mathbf{W}_1 \mathbf{d} + \mathbf{b}_1 + \mathbf{v}_1, \\ & \quad \mathbf{y}_i = \mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i + \mathbf{v}_i, \\ & \quad \forall i = 2, \dots, N, \\ & \quad \mathbf{0} \leq \mathbf{y}_i \perp \mathbf{v}_i \leq \mathbf{0}, \\ & \quad \forall i = 1, \dots, N, \\ & \quad \boldsymbol{\lambda} = \mathbf{W}_{N+1} \mathbf{y}_N + \mathbf{b}_{N+1}, \end{aligned}$$

- **Problem:** NP hard.
- **Solution:** Introduce the complementarity condition as a bilinear penalty term in the objective function.

$$\underset{(\mathbf{d}, \mathbf{y}, \mathbf{v}, \boldsymbol{\lambda}) \in \mathcal{O}}{\text{minimize}} \quad c(\boldsymbol{\lambda}) + \rho \sum_{i=1}^N \mathbf{y}_i^\top \mathbf{v}_i$$

$$\text{subject to } \mathbf{y}_1 = \mathbf{W}_1 \mathbf{d} + \mathbf{b}_1 + \mathbf{v}_1,$$

$$\mathbf{y}_i = \mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i + \mathbf{v}_i, \quad \forall i = 2, \dots, N$$

$$\boldsymbol{\lambda} = \mathbf{W}_{N+1} \mathbf{y}_N + \mathbf{b}_{N+1},$$

- Non-negative condition on $\mathbf{y}_i, \mathbf{v}_i \geq \mathbf{0}$, $\forall i = 1, \dots, N$ is included in set $\mathcal{O}$
- **Problem:** Non-convex problem.
- **Solution:** Reformulate the objective function into the difference of two convex functions and solve iteratively using the Difference of Convex Approach!

$$c(\lambda) + \rho \sum_{i=1}^N \mathbf{y}_i^\top \mathbf{v}_i \quad \Rightarrow \quad \underbrace{c(\lambda) + \frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i + \mathbf{v}_i\|_2^2}_{f_1(\mathbf{y}, \mathbf{v}), \text{convex function}} - \underbrace{\frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i - \mathbf{v}_i\|_2^2}_{f_2(\mathbf{y}, \mathbf{v}), \text{convex function}}$$

$$c(\boldsymbol{\lambda}) + \rho \sum_{i=1}^N \mathbf{y}_i^\top \mathbf{v}_i \quad \Rightarrow \quad \underbrace{c(\boldsymbol{\lambda}) + \frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i + \mathbf{v}_i\|_2^2}_{f_1(\mathbf{y}, \mathbf{v}), \text{convex function}} - \underbrace{\frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i - \mathbf{v}_i\|_2^2}_{f_2(\mathbf{y}, \mathbf{v}), \text{convex function}}$$

Algorithm 1 DCA for solving the bilinear problem

input: feasible guess $\mathbf{d}^0, \mathbf{y}^0, \mathbf{v}^0, \boldsymbol{\lambda}^0$, tolerance $\varepsilon_{\text{tol}} > 0$

output: optimized NN input $\mathbf{d}^*$

repeat

 set $k \leftarrow k + 1$.

 get $\mathbf{d}^{k+1}, \mathbf{y}^{k+1}, \mathbf{v}^{k+1}, \boldsymbol{\lambda}^{k+1}$ by solving DCA subproblem

until $f(\mathbf{y}^k, \mathbf{v}^k) - f(\mathbf{y}^{k+1}, \mathbf{v}^{k+1}) \leq \varepsilon_{\text{tol}}$

return $\mathbf{d}^* \leftarrow \mathbf{d}^{k+1}$

$f(\mathbf{y}^k, \mathbf{v}^k) = f_1(\mathbf{y}^k, \mathbf{v}^k) - f_2(\mathbf{y}^k, \mathbf{v}^k)$. Here, we write $\mathbf{d}, \mathbf{y}, \mathbf{v}, \boldsymbol{\lambda}$ to represent $\mathbf{d}_i, \mathbf{y}_i, \mathbf{v}_i, \boldsymbol{\lambda}_i \forall i = 1, \dots, N$

[Jara-Moroni et al., 2018]

$$\underbrace{-\frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i - \mathbf{v}_i\|_2^2}_{-f_2(\mathbf{y}, \mathbf{v}), \text{concave function}} \quad \Rightarrow \quad \underbrace{\frac{-\rho}{4} \sum_{i=1}^N (\mathbf{y}_i^k - \mathbf{v}_i^k)^\top (\mathbf{y}_i - \mathbf{v}_i)}_{-f_2(\mathbf{y}, \mathbf{v}) \text{ Linear overestimation of the concave part}}$$

$$\underbrace{-\frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i - \mathbf{v}_i\|_2^2}_{-f_2(\mathbf{y}, \mathbf{v}), \text{concave function}} \quad \Rightarrow \quad \underbrace{\frac{-\rho}{4} \sum_{i=1}^N (\mathbf{y}_i^k - \mathbf{v}_i^k)^\top (\mathbf{y}_i - \mathbf{v}_i)}_{-f_2(\mathbf{y}, \mathbf{v}) \text{ Linear overestimation of the concave part}}$$

DCA subproblem

$$\begin{aligned}
 & \underset{(\mathbf{d}, \mathbf{y}, \mathbf{v}, \boldsymbol{\lambda}) \in \mathcal{O}}{\text{minimize}} \quad c(\boldsymbol{\lambda}) + \frac{\rho}{4} \sum_{i=1}^N \|\mathbf{y}_i + \mathbf{v}_i\|_2^2 \\
 & \quad - \frac{\rho}{2} \sum_{i=1}^N (\mathbf{y}_i^k - \mathbf{v}_i^k)^\top (\mathbf{y}_i - \mathbf{v}_i) \quad \text{Linear overestimation of the concave part} \\
 & \text{subject to} \quad \mathbf{y}_1 = \mathbf{W}_1 \mathbf{d} + \mathbf{b}_1 + \mathbf{v}_1, \\
 & \quad \mathbf{y}_i = \mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i + \mathbf{v}_i, \quad \forall i = 2, \dots, N \\
 & \quad \boldsymbol{\lambda} = \mathbf{W}_{N+1} \mathbf{y}_N + \mathbf{b}_{N+1},
 \end{aligned}$$

The solution is extremely sensitive to ρ

- **Too Large:** Convergence slow
- **Too Small:** Violate complementarity condition

Need to find the smallest ρ that satisfies the complementarity condition.

General compact formulation

$$\begin{aligned} & \underset{\mathbf{d}, \mathbf{y}, \mathbf{v}}{\text{minimize}} && \mathbf{c}^\top \mathbf{y} \\ & \text{subject to} && \mathbf{A}\mathbf{d} \geq \mathbf{f}, \\ & && \mathbf{V}\mathbf{d} + \mathbf{W}\mathbf{y} + \mathbf{b} = \mathbf{v}, \\ & && \mathbf{0} \leq \mathbf{y}_i \perp \mathbf{v}_i \geq \mathbf{0}, \\ & && \forall i = 1, \dots, N, \end{aligned}$$

$\Rightarrow$

Bilinear compact formulation

$$\begin{aligned} & \underset{\mathbf{d}, \mathbf{y}, \mathbf{v}}{\text{minimize}} && \mathbf{c}^\top \mathbf{y} + \rho \mathbf{y}^\top \mathbf{v} \\ & \text{subject to} && \mathbf{A}\mathbf{d} \geq \mathbf{f} \quad : \lambda, \\ & && \mathbf{V}\mathbf{d} + \mathbf{W}\mathbf{y} + \mathbf{b} = \mathbf{v} \quad : \mu, \\ & && \mathbf{y}_i, \mathbf{v}_i \geq 0, \\ & && \forall i = 1, \dots, N \end{aligned}$$

Part of the derivation draw inspiration from Prop.16 in [Jara-Moroni et al., 2018]

Sensitivity to penalty parameter ρ



Strongly stationary point seeking relaxation

Let $(\tilde{\mathbf{d}}, \tilde{\mathbf{y}}, \tilde{\mathbf{v}})$ be some feasible point for the general compact formulation

$$\underset{\mathbf{d}, \mathbf{y}, \mathbf{v}}{\text{minimize}} \quad \mathbf{c}^\top \mathbf{y}$$

$$\text{subject to} \quad \mathbf{A}\mathbf{d} \geq \mathbf{f} \quad : \lambda,$$

$$\mathbf{V}\mathbf{d} + \mathbf{W}\mathbf{y} + \mathbf{b} = \mathbf{v} \quad : \mu,$$

$$y_i = 0 \quad : \mu_i^{y_0} \quad \forall i \in \mathcal{I}_y(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}),$$

$$y_i \geq 0 \quad : \mu_i^{y_+} \quad \forall i \in \mathcal{I}_v(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}),$$

$$v_i = 0 \quad : \mu_i^{v_0} \quad \forall i \in \mathcal{I}_v(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}),$$

$$v_i \geq 0 \quad : \mu_i^{v_+} \quad \forall i \in \mathcal{I}_y(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}),$$

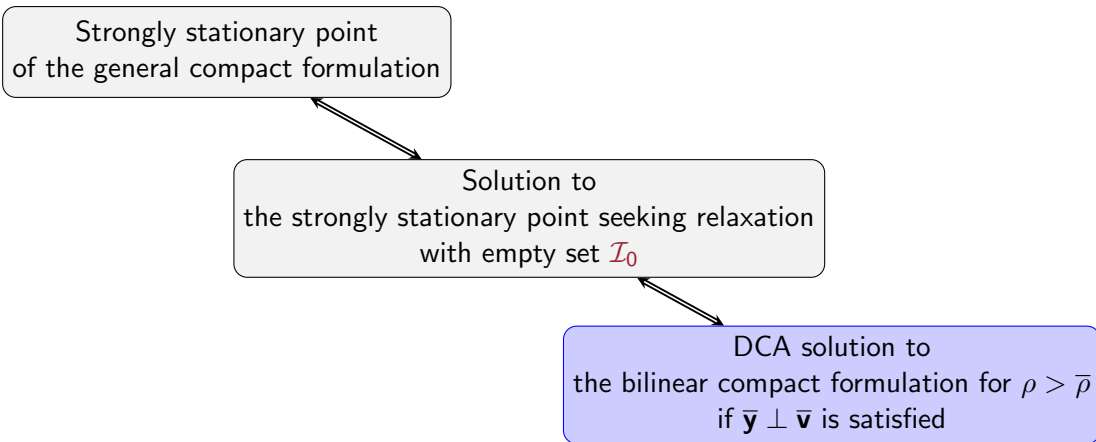
$$y_i, v_i \geq 0, \quad \forall i \in \mathcal{I}_0(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}),$$

$$\mathcal{I}_y(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}) \triangleq \{i \mid \tilde{y}_i = 0 < \tilde{v}_i\},$$

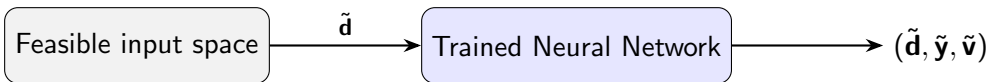
$$\mathcal{I}_v(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}) \triangleq \{i \mid \tilde{y}_i > 0 = \tilde{v}_i\},$$

$$\mathcal{I}_0(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}) \triangleq \{i \mid \tilde{y}_i = 0 = \tilde{v}_i\}.$$

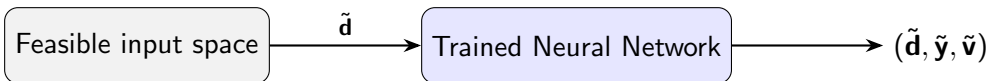
The solution $(\bar{\mathbf{d}}, \bar{\mathbf{y}}, \bar{\mathbf{v}})$ is the strongly stationary point of the general compact formulation if $\bar{\mathbf{y}}_i \perp \bar{\mathbf{v}}_i \quad \forall i = 1, \dots, N$, where the strongly stationary point is the point that satisfies the KKTs of the original non-convex optimization.



Obtain feasible point $(\tilde{\mathbf{d}}, \tilde{\mathbf{y}}, \tilde{\mathbf{v}})$



Obtain feasible point $(\tilde{\mathbf{d}}, \tilde{\mathbf{y}}, \tilde{\mathbf{v}})$



Feasibility insight of trained neural network

$$\mathcal{I}_0(\tilde{\mathbf{y}}, \tilde{\mathbf{v}}) = \emptyset$$



$$\tilde{y}_i \neq \tilde{v}_i, \forall i = 1, \dots, N$$



$$\mathbf{W}_i \mathbf{y}_{i-1} + \mathbf{b}_i \neq \mathbf{0}, \forall i = 1, \dots, N$$

- Often true
- If not, we can losslessly reduce NN size — no more zero neuron!

- $(\bar{\mathbf{d}}, \bar{\mathbf{y}}, \bar{\mathbf{v}})$: From strongly stationary point seeking relaxation
- **Obtain ρ** : From partial Karush–Kuhn–Tucker conditions of the strongly stationary point seeking relaxation formulation and the bilinear compact formulation
- Combine lagrangian multipliers $\boldsymbol{\mu}^v = [\boldsymbol{\mu}^{v_0^\top} \boldsymbol{\mu}^{v_+^\top}]^\top$, $\boldsymbol{\mu}^y = [\boldsymbol{\mu}^{y_0^\top} \boldsymbol{\mu}^{y_+^\top}]^\top$

Relaxation partial KKT

$$\mathbf{A}^\top \boldsymbol{\lambda} + \mathbf{V}^\top \boldsymbol{\mu}^v = \mathbf{0},$$

$$\mathbf{W}^\top \boldsymbol{\mu}^v + \boldsymbol{\mu}^y = \mathbf{c},$$

$$\mathbf{0} \leq \boldsymbol{\lambda} \perp \mathbf{A}\mathbf{d} - \mathbf{f} \geq \mathbf{0},$$

$$\mathbf{y}^\top \boldsymbol{\mu}^y = 0, \quad \mathbf{v}^\top \boldsymbol{\mu}^v = 0,$$

$$\mu_i^y \geq 0, \quad \forall i \in \mathcal{I}_v(\tilde{\mathbf{y}}, \tilde{\mathbf{v}})$$

$$\mu_i^v \geq 0, \quad \forall i \in \mathcal{I}_y(\tilde{\mathbf{y}}, \tilde{\mathbf{v}})$$

Bilinear partial KKT

$$\mathbf{A}^\top \boldsymbol{\lambda} + \mathbf{V}^\top \boldsymbol{\mu} = \mathbf{0},$$

$$\mathbf{0} \leq \mathbf{y} \perp -\mathbf{W}^\top \boldsymbol{\mu} + \mathbf{c} + \rho \mathbf{v} \geq \mathbf{0},$$

$$\mathbf{0} \leq \boldsymbol{\lambda} \perp \mathbf{A}\mathbf{d} - \mathbf{f} \geq \mathbf{0},$$

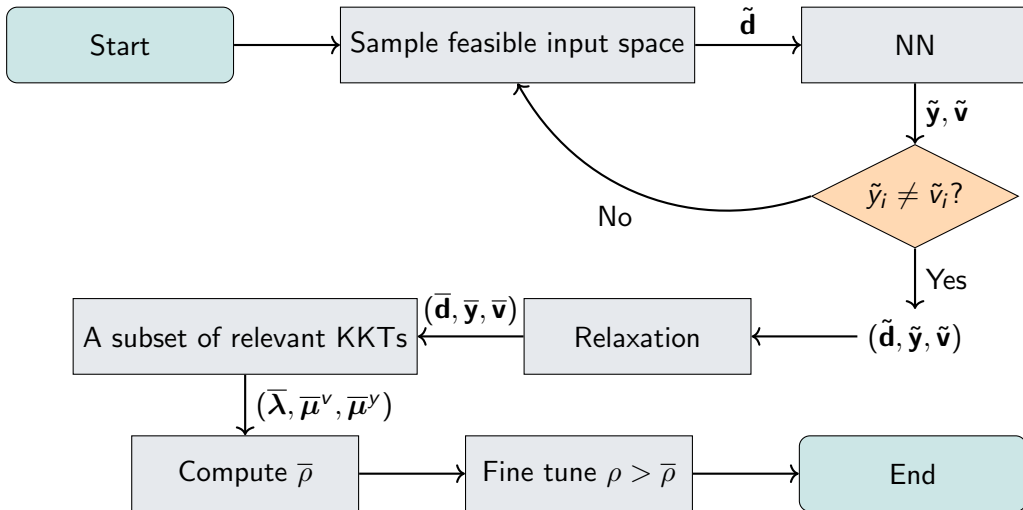
$$\mathbf{0} \leq \mathbf{v} \perp \rho \mathbf{y} + \boldsymbol{\mu} \geq \mathbf{0}.$$

We substitute μ^v from relation partial KKT into μ from bilinear partial KKT. Through some algebraic manipulation, we retrieve the following condition on $\bar{\rho}$

Lower bound for ρ

$$\bar{\rho} \triangleq \max \left\{ 0, \left\{ -\frac{\bar{\mu}_i^y}{\bar{v}_i} \mid \bar{v}_i > 0 \right\}, \left\{ -\frac{\bar{\mu}_i^v}{\bar{y}_i} \mid \bar{y}_i > 0 \right\} \right\}.$$

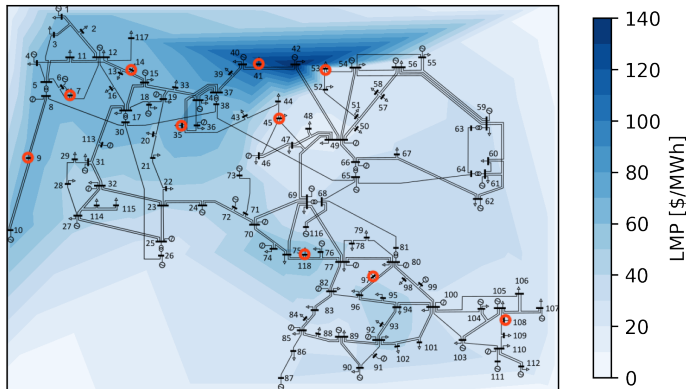
Computation for DCA penalty term ρ algorithm



Data center workload allocation in power grids



- Consider a network of spatially distributed data centers.
- Operators allocate workloads to minimize the cost of electricity consumption.



- In practice, we do not have network data or other parameters.
- Use a neural network to map the workload allocation to the cost of electricity.

- **Goal:** Allocate data center load to minimize the cost of electricity consumption.
- **Given:** Forecast of total electricity demand Δ .
- Data center allocation formulated as a bilevel optimization problem where $\mathbf{d}$ is the data center loads with upper and lower bounds $\bar{\mathbf{d}}$ and $\underline{\mathbf{d}}$. $\boldsymbol{\pi}$ is the price per demand.
- **Challenge:** Unknown underlying physical model

$$\begin{aligned} & \underset{\mathbf{d}}{\text{minimize}} && \boldsymbol{\pi}^\top \mathbf{d} \\ & \text{subject to} && \underline{\mathbf{d}} \leq \mathbf{d} \leq \bar{\mathbf{d}} \\ & && \mathbf{1}^\top \mathbf{d} = \Delta \\ & && \boldsymbol{\pi} \in \\ & && \text{Unknown model} \end{aligned}$$

- **Goal:** Allocate data center load to minimize the cost of electricity consumption.
- **Given:** Forecast of total electricity demand Δ .
- Data center allocation formulated as a bilevel optimization problem where $\mathbf{d}$ is the data center loads with upper and lower bounds $\bar{\mathbf{d}}$ and $\underline{\mathbf{d}}$. π is the price per demand.
- **Challenge:** Unknown underlying physical model
- **Solution:** Replace by a neural network trained with a labeled dataset on historical record $\{(\mathbf{d}_1, \lambda_1), \dots, (\mathbf{d}_n, \lambda_n)\}$.

$$\begin{aligned} & \underset{\mathbf{d}}{\text{minimize}} && \pi^\top \mathbf{d} \\ & \text{subject to} && \underline{\mathbf{d}} \leq \mathbf{d} \leq \bar{\mathbf{d}} \\ & && \mathbf{1}^\top \mathbf{d} = \Delta \end{aligned}$$

$$\pi \in$$

Unknown model

$\Rightarrow$

$$\begin{aligned} & \underset{\mathbf{d}}{\text{minimize}} && \text{NN}(\mathbf{d}) \\ & \text{subject to} && \underline{\mathbf{d}} \leq \mathbf{d} \leq \bar{\mathbf{d}} \\ & && \mathbf{1}^\top \mathbf{d} = \Delta \end{aligned}$$

- **Method:** Using locational marginal prices (LMPs) as π derived from the optimal power flow (OPF) model. Use this data to train neural network.

- **Method:** Using locational marginal prices (LMPs) as π derived from the optimal power flow (OPF) model. Use this data to train neural network.

$$\begin{array}{ll} \underset{\underline{\mathbf{p}} \leq \mathbf{p} \leq \bar{\mathbf{p}}}{\text{minimize}} & \mathbf{p}^\top \mathbf{C} \mathbf{p} + \mathbf{c}^\top \mathbf{p} \\ \text{subject to} & \mathbf{1}^\top (\mathbf{p} - \boldsymbol{\ell} - \mathbf{d}) = 0 \quad : \lambda \\ & |\mathbf{F}(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})| \leq \bar{\mathbf{f}} \quad : \underline{\boldsymbol{\mu}}, \bar{\boldsymbol{\mu}} \end{array}$$

- $\mathbf{p}$: generator dispatch
- $[\underline{\mathbf{p}}, \bar{\mathbf{p}}]$: feasible range of dispatch
- $\boldsymbol{\ell}$: conventional loads
- $\mathbf{d}$: data center loads
- $\mathbf{F}$: the matrix of power transfer distribution factors, which maps net power injections $(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})$ to power flows as $\mathbf{F}(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})$
- $\bar{\mathbf{f}}$: line capacity

- **Method:** Using locational marginal prices (LMPs) as π derived from the optimal power flow (OPF) model. Use this data to train neural network.

$$\begin{aligned} & \underset{\underline{\mathbf{p}} \leq \mathbf{p} \leq \bar{\mathbf{p}}}{\text{minimize}} && \mathbf{p}^\top \mathbf{C} \mathbf{p} + \mathbf{c}^\top \mathbf{p} \\ & \text{subject to} && \mathbf{1}^\top (\mathbf{p} - \boldsymbol{\ell} - \mathbf{d}) = 0 \quad : \lambda \\ & && |\mathbf{F}(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})| \leq \bar{\mathbf{f}} \quad : \underline{\boldsymbol{\mu}}, \bar{\boldsymbol{\mu}} \end{aligned}$$

- $\mathbf{p}$: generator dispatch
- $[\underline{\mathbf{p}}, \bar{\mathbf{p}}]$: feasible range of dispatch
- $\boldsymbol{\ell}$: conventional loads
- $\mathbf{d}$: data center loads
- $\mathbf{F}$: the matrix of power transfer distribution factors, which maps net power injections $(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})$ to power flows as $\mathbf{F}(\mathbf{p} - \boldsymbol{\ell} - \mathbf{d})$
- $\bar{\mathbf{f}}$: line capacity

Locational Marginal Price(LMP)

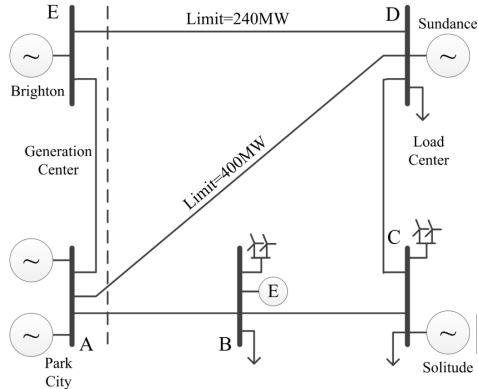
$$\pi = \mathbf{1}\lambda^* - \mathbf{F}^\top \bar{\boldsymbol{\mu}}^* + \mathbf{F}^\top \underline{\boldsymbol{\mu}}^*$$

[Chatzivasileiadis, 2018]

Experiment on PJM-5 bus system



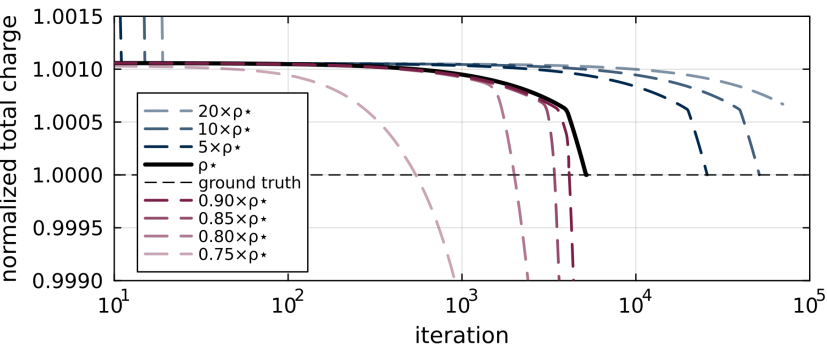
- 3 loads act as data centers with demand in the range of $[0.8, 1]$ of the nominal value.
- NN with 2 hidden layers, 50 neurons each. 10,000 training samples
- Global solution is computed using SOS1 constraints and Gurobi solver.



[Zhou, 2023]

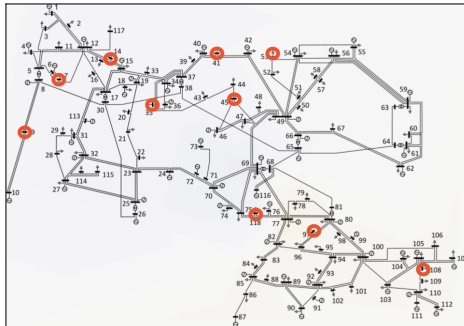
>>> X. Liu

Experiment on PJM-5 bus system

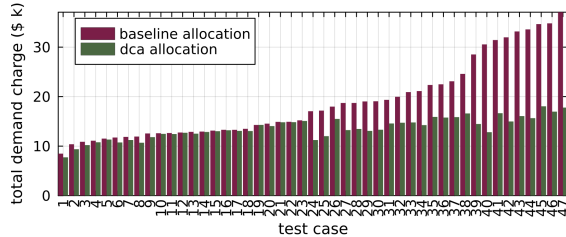


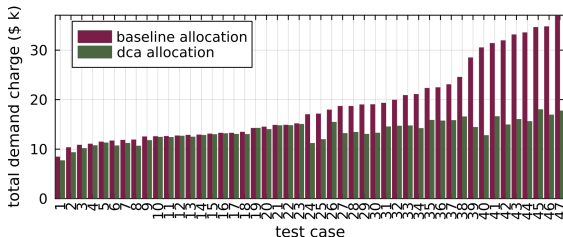
- ρ^* : **fast convergence**
- $\rho < \rho^*$: Does not satisfy complementarity condition
- $\rho > \rho^*$: Slow convergence

- **NN training:** NN with 5-hidden layer and 1,000 ReLUs, trained on 12,500 demand samples in the range of $[0, 100]$ MWh.
- **DCA solution:** Computes a new demand allocation to improve the baseline.
- **Samples:** Randomly picked 50 baseline cases, of which DCA could successfully converge and 47 of them roughly match the output of the OPF model.

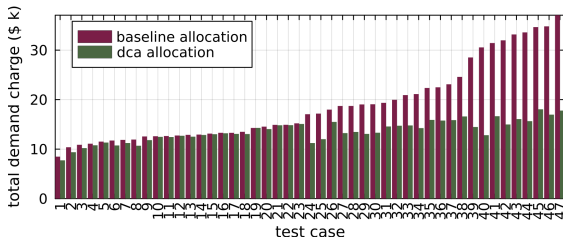


Experiment on IEEE 118 bus system

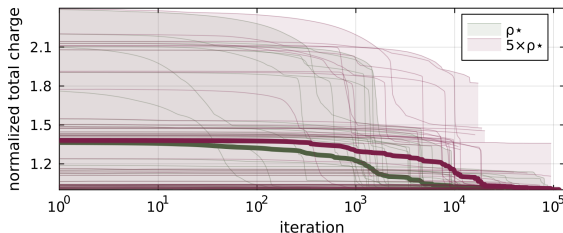


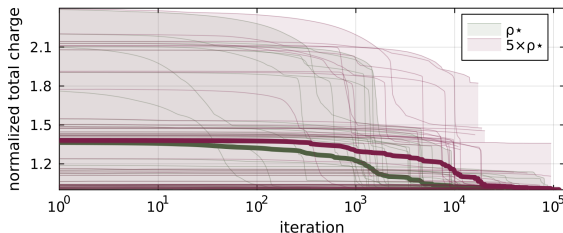
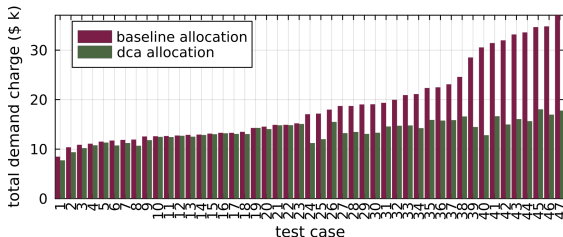


- **First 23 cases:** The savings are small. Lightly congested with little to no benefit from spatial load redistribution
- **Remaining cases:** Significant improvement



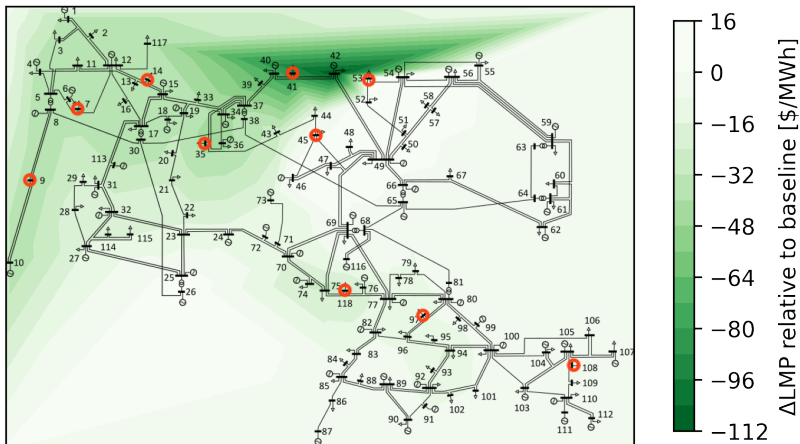
- **First 23 cases:** The savings are small. Lightly congested with little to no benefit from spatial load redistribution
- **Remaining cases:** Significant improvement





- **First 23 cases:** The savings are small. Lightly congested with little to no benefit from spatial load redistribution
- **Remaining cases:** Significant improvement
- The smaller ρ consistently leads to faster convergence.
- Each iteration takes ≈ 0.37 seconds, allowing for convergence in around 1 hour on average.

Experiment on IEEE 118 bus system



- The difference between the LMPs induced at baseline and the DCA solutions.
- LMPs were significantly reduced at the data centers.

- Developed a difference-of-convex algorithm to solve decision-making problems with objectives or constraints represented by trained neural networks.
- Avoided the computational bottleneck of ReLU logical constraints through informed penalization of such constraints in the objective function.
- Demonstrated on the data center load allocation application and showed significant reduction in cost compared to baseline.

Questions?





Chatzivasileiadis, S. (2018).

Lecture notes on optimal power flow (opf).

arXiv preprint arXiv:1811.00943.



Dathathri, S., , et al. (2020).

Enabling certification of verification-agnostic networks via memory-efficient semidefinite programming.

Advances in Neural Information Processing Systems, pages 5318–5331.



Dvorkin, V. (2025).

Lecture slides: Computational power systems.

Unpublished lecture slides, accessed on 2025-04-30.



Fazlyab, M., Morari, M., and Pappas, G. J. (2020).

Safety verification and robustness analysis of neural networks via quadratic constraints and semidefinite programming.

IEEE Trans. Autom. Control., 67(1):1–15.



Grimstad, B. and Andersson, H. (2019).
ReLU networks as surrogate models in mixed-integer linear programs.
Comput. Chem. Eng., 131:106580.



Jara-Moroni, F., Pang, J.-S., and Wächter, A. (2018).
A study of the difference-of-convex approach for solving linear programs with complementarity constraints.
Mathematical Programming, 169:221–254.



Tjeng, V., Xiao, K., and Tedrake, R. (2017).
Evaluating robustness of neural networks with mixed integer programming.
Preprint.



Turner, M. et al. (2024).
Pyscipopt-ml: Embedding trained machine learning models into mixed-integer programs.
Preprint.



Wang, S. et al. (2021).

Beta-crown: Efficient bound propagation with per-neuron split constraints for neural network robustness verification.

Advances in neural information processing systems, 34:29909–29921.



Zhou, A. (2023).

Pjm 5 bus data.

<https://dx.doi.org/10.21227/7ex2-1t83>.

Accessed on 2025-04-30.

The End

- **Partial Lagrangian function** (dualize the coupling constraints only):

$$\begin{aligned}\mathcal{L}(\mathbf{p}, \lambda, \bar{\underline{\mu}}, \underline{\mu}) = & c(\mathbf{p}) - \lambda \mathbf{1}^\top (\mathbf{p} - \mathbf{d}) \\ & + \bar{\underline{\mu}}^\top (\mathbf{F}(\mathbf{p} - \mathbf{d}) - \bar{\mathbf{f}}) + \underline{\mu}^\top (-\mathbf{F}(\mathbf{p} - \mathbf{d}) - \bar{\mathbf{f}})\end{aligned}$$

- Group terms corresponding to dispatch $\mathbf{p}$, demand $\mathbf{d}$ and line limits $\bar{\mathbf{f}}$:

$$\mathcal{L} = \mathcal{L}^{\mathbf{p}} + \mathcal{L}^{\mathbf{d}} + \mathcal{L}^{\mathbf{f}}, \quad \text{where}$$

$$\mathcal{L}^{\mathbf{p}}(\mathbf{p}, \lambda, \bar{\underline{\mu}}, \underline{\mu}) = c(\mathbf{p}) - (\mathbf{1}\lambda - \mathbf{F}^\top \bar{\underline{\mu}} + \mathbf{F}^\top \underline{\mu})^\top \mathbf{p}$$

$$\mathcal{L}^{\mathbf{d}}(\lambda, \bar{\underline{\mu}}, \underline{\mu}) = (\mathbf{1}\lambda - \mathbf{F}^\top \bar{\underline{\mu}} + \mathbf{F}^\top \underline{\mu})^\top \mathbf{d}$$

$$\mathcal{L}^{\mathbf{f}}(\bar{\underline{\mu}}, \underline{\mu}) = -(\underline{\mu} + \bar{\underline{\mu}})^\top \bar{\mathbf{f}}$$

Power dispatch $\mathbf{p}$ and demand $\mathbf{d}$ share the same multiplier but with opposite signs

[Dvorkin, 2025]

$$\pi^*(\lambda^*, \bar{\underline{\mu}}^*, \underline{\mu}^*) = \underbrace{\mathbf{1}\lambda^*}_{\text{uniform}} - \underbrace{\mathbf{F}^\top(\bar{\underline{\mu}}^* - \underline{\mu}^*)}_{\text{congestion}} \in \mathbb{R}^N$$

- π_n^* is the cost of supplying the next unit of demand at node n
- in case of congestion ($\bar{\underline{\mu}}^* > 0$ or $\underline{\mu}^* > 0$), electricity price varies across the grid
- price at the reference bus is λ^* (the ref. column of $\mathbf{F}$ is zero)

- Use matrix $\mathbf{F} \in \mathbb{R}^{E \times N}$ of power transfer distribution factors (PTDF)
 - how the power flow in line e changes w.r.t. to the change of power injection at node n ?
 - obtained by manipulating the DC bus admittance matrix $\mathbf{B}$
- Power flows $f = \mathbf{F}(\mathbf{p} - \mathbf{d})$ (distribution of net injections across power lines)
 - minimize $c(\mathbf{p})$ *generation cost*
 - subject to $\mathbf{1}^\top (\mathbf{p} - \mathbf{d}) = 0$ *active power balance*
 - $|\mathbf{F}(\mathbf{p} - \mathbf{d})| \leq \bar{f}$ *power flow limits*
 - $\underline{\mathbf{p}} \leq \mathbf{p} \leq \bar{\mathbf{p}}$ *min/max gen p-limits*

Constraints:

$$z \geq 0$$

$$\mu \geq 0$$

$$z\mu = 0$$

Big-M reformulation:

$$\begin{cases} \mu \leq Mu \\ z \leq M(1 - u) \end{cases}$$

where $u \in \{0, 1\}$ is an auxiliary binary variable, and M is a large enough constant

[Dvorkin, 2025]

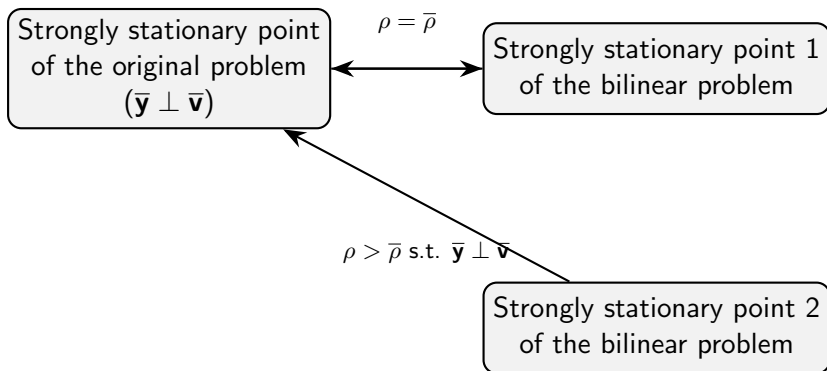
$$x_{i+1} = \text{ReLU}(\mathbb{L}_i(x_i)).$$

$$x_{i+1} \geq 0 [\lambda_i^a], x_{i+1} \geq \mathbb{L}_i(x_i) [\lambda_i^b]$$

$$x_{i+1} \odot (x_{i+1} - \mathbb{L}_i(x_i)) \leq 0 [\lambda_i^c], x_i \odot x_i - (\ell_i + u_i) \odot x_i + \ell_i \odot u_i \leq 0 [\lambda_i^d].$$

$$\begin{aligned} \mathcal{L}(x_i, x_{i+1}, \lambda_i) &= (-x_{i+1})^\top \lambda_i^a + (\mathbb{L}_i(x_i) - x_{i+1})^\top \lambda_i^b \\ &\quad + (x_{i+1} \odot (x_{i+1} - \mathbb{L}_i(x_i)))^\top \lambda_i^c + (x_i \odot x_i - (\ell_i + u_i) \odot x_i + \ell_i \odot u_i)^\top \lambda_i^d \\ &= \underbrace{(\ell_i \odot u_i)^\top \lambda_i^d}_{\text{independent of } x_i, x_{i+1}} - \underbrace{x_{i+1}^\top \lambda_i^a + (\mathbb{L}_i(x_i))^\top \lambda_i^b - x_{i+1}^\top \lambda_i^b - x_i^\top ((\ell_i + u_i) \odot \lambda_i^d)}_{\text{linear in } x_i, x_{i+1}} \\ &\quad + \underbrace{x_{i+1}^\top \text{diag}(\lambda_i^c) x_{i+1} - x_{i+1}^\top \text{diag}(\lambda_i^c) \mathbb{L}_i(x_i) + x_i^\top \text{diag}(\lambda_i^d) x_i}_{\text{Quadratic in } x_i, x_{i+1}}. \end{aligned}$$

- ReLU relaxation: $w^\top \text{ReLU}(v) \geq w^\top \mathbf{D}v + b'$ where $\mathbf{D}$ is a diagonal matrix containing free variables $0 \leq \alpha_j \leq 1$ only when $\mathbf{u}_j > 0 > \mathbf{l}_j$ and $w_j \geq 0$, while its rest values as well as constant b' are determined by $\mathbf{l}, \mathbf{u}, w$.
- $\min_{x \in \mathcal{C}} f(x) \geq \min_{x \in \mathcal{C}} \mathbf{a}_{\text{CROWN}}^\top x + c_{\text{CROWN}}$ where $\mathbf{a}_{\text{CROWN}}$ and c_{CROWN} can be computed using $\mathbf{W}^{(i)}, \mathbf{b}^{(i)}, \mathbf{l}^{(i)}, \mathbf{u}^{(i)}$ in polynomial time.



- The fine tuning to make sure complementarity condition is satisfied is easy because intuitively, if $\bar{\rho}$ works as a penalty parameter in one case, the penalty is penalize the bilinear part enough to satisfy the complementarity condition.
- This result is also proven experimentally.