

Non-Euclidean Monotone Operator Theory for Robustness of Implicit Neural Networks

Saber Jafarpour



University of Colorado **Boulder**

February 14, 2025

An increase in deployment of learning-based algorithms

Extremely fragile wrt input perturbations

Adversarial Perturbations¹

Small changes in the input

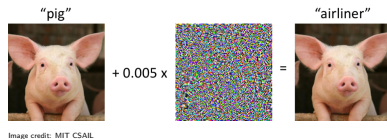


Large changes in the output

¹C. Szegedy, et al. Intriguing properties of neural networks, ICLR, 2014

An increase in deployment of learning-based algorithms

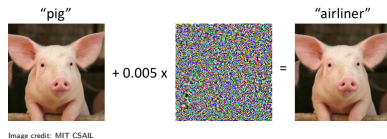
Extremely fragile wrt input perturbations



¹C. Szegedy, et al. Intriguing properties of neural networks, ICLR, 2014

An increase in deployment of learning-based algorithms

Extremely fragile wrt input perturbations

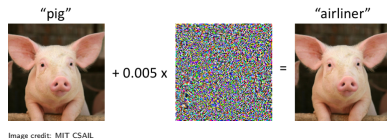


- Guaranteeing robustness of learning algorithms is critical in their real-world applications

¹C. Szegedy, et al. Intriguing properties of neural networks, ICLR, 2014

An increase in deployment of learning-based algorithms

Extremely fragile wrt input perturbations

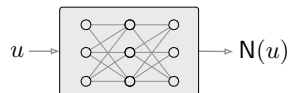


- Guaranteeing robustness of learning algorithms is critical in their real-world applications

Robustness of learning algorithm

Input perturbation set $\mathcal{U}$ and unsafe output domain $\mathcal{S}_{\text{unsafe}}$:

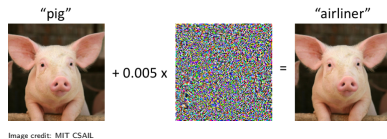
$$\mathcal{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset.$$



¹C. Szegedy, et al. Intriguing properties of neural networks, ICLR, 2014

An increase in deployment of learning-based algorithms

Extremely fragile wrt input perturbations

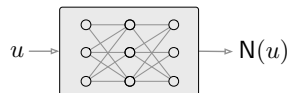


- Guaranteeing robustness of learning algorithms is critical in their real-world applications

Robustness of learning algorithm

Input perturbation set $\mathcal{U}$ and unsafe output domain $\mathcal{S}_{\text{unsafe}}$:

$$\mathcal{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset.$$



- 1 **Verification:** For a given learning algorithm can we check its robustness?
- 2 **Training:** how to design robust learning algorithms?

¹C. Szegedy, et al. Intriguing properties of neural networks, ICLR, 2014

Input-output Lipschitz Bounds

A framework for robustness analysis

Goal: over-approximate $N(\mathcal{U})$ with $\bar{N}(\mathcal{U})$ and check if $\bar{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset$.

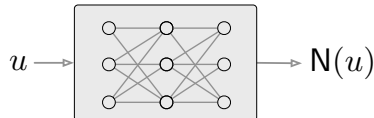
Input-output Lipschitz Bounds

A framework for robustness analysis

Goal: over-approximate $N(\mathcal{U})$ with $\bar{N}(\mathcal{U})$ and check if $\bar{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset$.

Lipschitz bound

$$\|N(u) - N(v)\| \leq \text{Lip}N \|u - v\|, \text{ for every } u, v \in \mathcal{U}$$



Input-output Lipschitz Bounds

A framework for robustness analysis

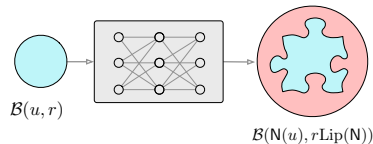
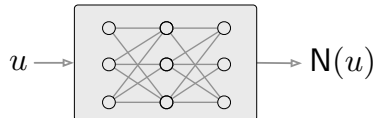
Goal: over-approximate $N(\mathcal{U})$ with $\bar{N}(\mathcal{U})$ and check if $\bar{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset$.

Lipschitz bound

$$\|N(u) - N(v)\| \leq \text{Lip}N \|u - v\|, \text{ for every } u, v \in \mathcal{U}$$

Robustness via Lipschitz bounds

$$N(\mathcal{B}(u, r)) \subseteq \mathcal{B}(N(u), r \text{Lip}N) := \bar{N}(\mathcal{U})$$



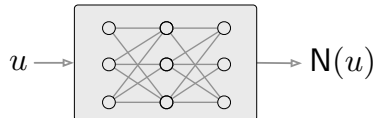
Input-output Lipschitz Bounds

A framework for robustness analysis

Goal: over-approximate $N(\mathcal{U})$ with $\bar{N}(\mathcal{U})$ and check if $\bar{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset$.

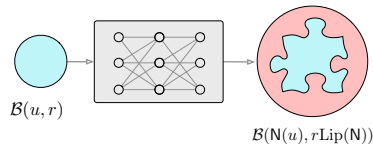
Lipschitz bound

$$\|N(u) - N(v)\| \leq \text{Lip}N \|u - v\|, \text{ for every } u, v \in \mathcal{U}$$



Robustness via Lipschitz bounds

$$N(\mathcal{B}(u, r)) \subseteq \mathcal{B}(N(u), r\text{Lip}N) := \bar{N}(\mathcal{U})$$



- A. Virmaux and K. Scaman. [Lipschitz regularity of deep neural networks: analysis and efficient estimation.](#) *NeurIPS*, 2018
- Mahyar Fazlyab, et al, [Efficient and Accurate Estimation of Lipschitz Constants for Deep Neural Networks.](#) *NeurIPS*, 2019
- Alexandre Araujo, et al, [A Unified Algebraic Perspective on Lipschitz Neural Networks,](#) *ICLR*, 2023

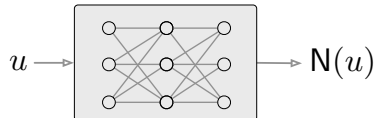
Input-output Lipschitz Bounds

A framework for robustness analysis

Goal: over-approximate $N(\mathcal{U})$ with $\bar{N}(\mathcal{U})$ and check if $\bar{N}(\mathcal{U}) \cap \mathcal{S}_{\text{unsafe}} = \emptyset$.

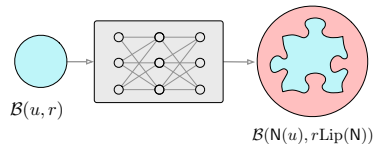
Lipschitz bound

$\|N(u) - N(v)\| \leq \text{Lip}N \|u - v\|$, for every $u, v \in \mathcal{U}$



Robustness via Lipschitz bounds

$N(\mathcal{B}(u, r)) \subseteq \mathcal{B}(N(u), r\text{Lip}N) := \bar{N}(\mathcal{U})$



In this talk: use Monotone Operator Theory to estimate Lipschitz bound of learning algorithms

Monotone Operator Theory

A classical framework in functional analysis

Let $\mathcal{H}$ be a Hilbert space with inner product $\langle \cdot, \cdot \rangle_{\mathcal{H}}$. Then $F : \mathcal{H} \rightarrow \mathcal{H}$ is a monotone operator if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq 0, \quad \text{for all } x, y$$

and is strongly monotone with parameter $m \geq 0$ if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq m \|x - y\|_{\mathcal{H}}^2, \quad \text{for all } x, y$$

Monotone Operator Theory

A classical framework in functional analysis

Let $\mathcal{H}$ be a Hilbert space with inner product $\langle \cdot, \cdot \rangle_{\mathcal{H}}$. Then $F : \mathcal{H} \rightarrow \mathcal{H}$ is a monotone operator if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq 0, \quad \text{for all } x, y$$

and is strongly monotone with parameter $m \geq 0$ if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq m \|x - y\|_{\mathcal{H}}^2, \quad \text{for all } x, y$$

- Minty (1962), Browder (1967), Rockafellar (1966)
- Bauschke and Combettes (2017), Ryu and Boyd (2016)

Monotone Operator Theory

A classical framework in functional analysis

Let $\mathcal{H}$ be a Hilbert space with inner product $\langle \cdot, \cdot \rangle_{\mathcal{H}}$. Then $F : \mathcal{H} \rightarrow \mathcal{H}$ is a monotone operator if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq 0, \quad \text{for all } x, y$$

and is strongly monotone with parameter $m \geq 0$ if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq m \|x - y\|_{\mathcal{H}}^2, \quad \text{for all } x, y$$

- Minty (1962), Browder (1967), Rockafellar (1966)
- Bauschke and Combettes (2017), Ryu and Boyd (2016)

Optimization

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

f is convex iff ∇f is monotone

Dynamical systems

$$\dot{x} = f(x)$$

f is contracting wrt ℓ_2 -norm iff $-f$ is strongly monotone

Monotone Operator Theory

A classical framework in functional analysis

Let $\mathcal{H}$ be a Hilbert space with inner product $\langle \cdot, \cdot \rangle_{\mathcal{H}}$. Then $F : \mathcal{H} \rightarrow \mathcal{H}$ is a monotone operator if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq 0, \quad \text{for all } x, y$$

and is strongly monotone with parameter $m \geq 0$ if

$$\langle F(x) - F(y), x - y \rangle_{\mathcal{H}} \geq m \|x - y\|_{\mathcal{H}}^2, \quad \text{for all } x, y$$

- Minty (1962), Browder (1967), Rockafellar (1966)
- Bauschke and Combettes (2017), Ryu and Boyd (2016)

Theorem (classical)

Let $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a strongly monotone operator wrt to $\langle \cdot, \cdot \rangle_{\mathbb{R}^n}$, then

- ① $F(x) = 0$ has a unique solution x^* , and
- ② x^* can be computed using the average iteration $x_{k+1} = (1 - \theta)x_k + \theta F(x_k)$.

Monotone Operator Theory

Logarithmic norm

Definition: Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

- directional derivative of matrix norm $\|\cdot\|$ in direction of A at point I_n .

³ A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

Monotone Operator Theory

Logarithmic norm

Definition: Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

- directional derivative of matrix norm $\|\cdot\|$ in direction of A at point I_n .
- **In the literature:** one-sided Lipschitz constant, matrix measure

³ A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

Monotone Operator Theory

Logarithmic norm

Definition: Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

- directional derivative of matrix norm $\|\cdot\|$ in direction of A at point I_n .
- **In the literature**: one-sided Lipschitz constant, matrix measure

Theorem²(Logarithmic norms and monotone operators)

$F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is a monotone operator if and only if $\mu_2(-D_x F(x)) \leq 0$, for every $x \in \mathbb{R}^n$

³ A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

Monotone Operator Theory

Logarithmic norm

Definition: Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

- directional derivative of matrix norm $\|\cdot\|$ in direction of A at point I_n .
- **In the literature**: one-sided Lipschitz constant, matrix measure

Theorem²(Logarithmic norms and monotone operators)

$F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is a monotone operator if and only if $\mu_2(-D_x F(x)) \leq 0$, for every $x \in \mathbb{R}^n$

Extend monotone operator theory to non-Euclidean norm spaces

³ A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

Non-Euclidean Monotone Operator Theory

Definition and characterizations

Let $\|\cdot\|$ be a norm on $\mathbb{R}^n$. Then $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is a monotone wrt to $\|\cdot\|$ if

$$\mu_{\|\cdot\|}(-DF(x)) \leq 0, \quad \text{for all } x$$

and is strongly monotone with parameter $m \geq 0$ if

$$\mu_{\|\cdot\|}(-DF(x)) \leq -m, \quad \text{for all } x$$

³A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

Non-Euclidean Monotone Operator Theory

Definition and characterizations

Let $\|\cdot\|$ be a norm on $\mathbb{R}^n$. Then $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ is a monotone wrt to $\|\cdot\|$ if

$$\mu_{\|\cdot\|}(-DF(x)) \leq 0, \quad \text{for all } x$$

and is strongly monotone with parameter $m \geq 0$ if

$$\mu_{\|\cdot\|}(-DF(x)) \leq -m, \quad \text{for all } x$$

Theorem (Non-Euclidean version)³

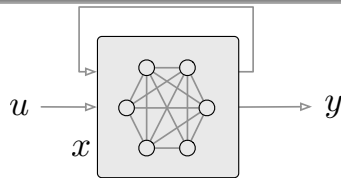
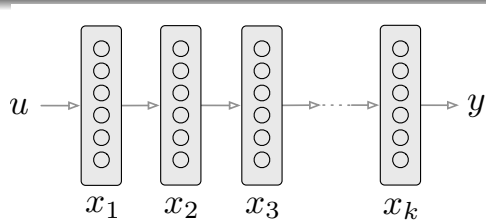
Let $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a strongly monotone operator wrt to a norm $\|\cdot\|$, then

- ① $F(x) = 0$ has a unique solution x^* , and
- ② x^* can be computed using the average iteration $x_{k+1} = (1 - \theta)x_k + \theta F(x_k)$.

³ A. Davydov, SJ, A. V. Proskurnikov, and F. Bullo. Non-Euclidean monotone operator theory and applications. JMLR, 2024

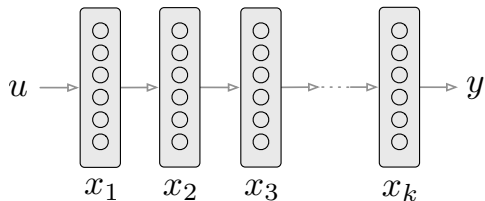
Implicit Neural Networks

Definition via fixed-point equations



Implicit Neural Networks

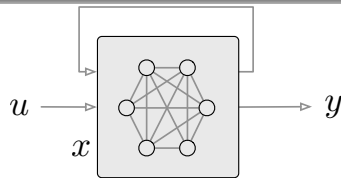
Definition via fixed-point equations



- Feedforward neural networks:

$$x^{i+1} = \Phi(A_i x^i + b_i), \quad x^0 = u$$

$$y = A_k x^k + b_k$$



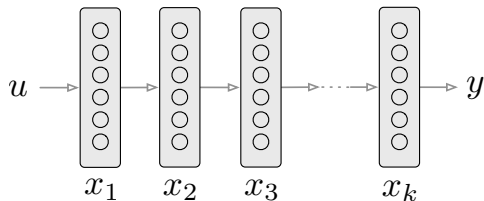
- Implicit neural networks:

$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

Implicit Neural Networks

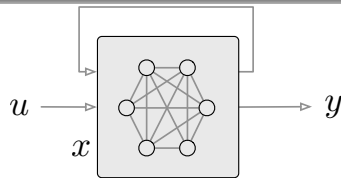
Definition via fixed-point equations



- Feedforward neural networks:

$$x^{i+1} = \Phi(A_i x^i + b_i), \quad x^0 = u$$

$$y = A_k x^k + b_k$$



- Implicit neural networks:

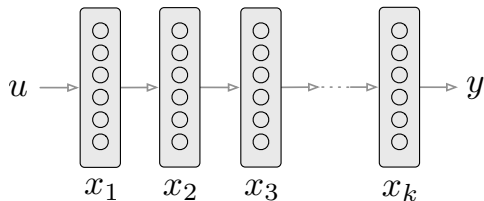
$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

- $\Phi(y_1, \dots, y_n) = (\phi_1(y_1), \dots, \phi_n(y_n))^T$ is a diagonal activation function
- activation functions are slope-restricted in $[0, 1]$, i.e., $0 \leq \frac{\phi_i(x) - \phi_i(y)}{x - y} \leq 1$ for all $x, y \in \mathbb{R}$

Implicit Neural Networks

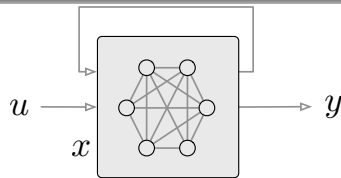
Definition via fixed-point equations



- Feedforward neural networks:

$$x^{i+1} = \Phi(A_i x^i + b_i), \quad x^0 = u$$

$$y = A_k x^k + b_k$$



- Implicit neural networks:

$$x = \Phi(Ax + Bu + b)$$

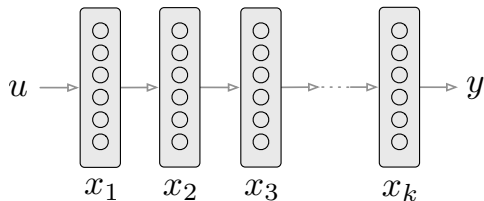
$$y = Cx + c$$

- $\Phi(y_1, \dots, y_n) = (\phi_1(y_1), \dots, \phi_n(y_n))^T$ is a diagonal activation function
- activation functions are slope-restricted in $[0, 1]$, i.e., $0 \leq \frac{\phi_i(x) - \phi_i(y)}{x - y} \leq 1$ for all $x, y \in \mathbb{R}$

Notion of Layer: output is defined **implicitly** as a function of input
e.g., fixed-point equation, differential equations, optimization problem

Implicit Neural Networks

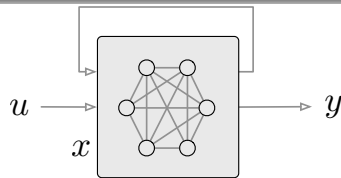
Definition via fixed-point equations



- Feedforward neural networks:

$$x^{i+1} = \Phi(A_i x^i + b_i), \quad x^0 = u$$

$$y = A_k x^k + b_k$$



- Implicit neural networks:

$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

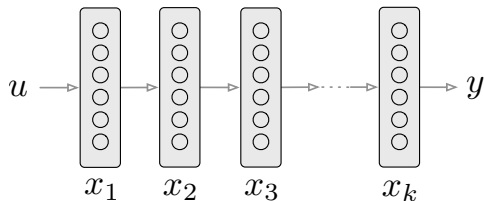
- $\Phi(y_1, \dots, y_n) = (\phi_1(y_1), \dots, \phi_n(y_n))^T$ is a diagonal activation function
- activation functions are slope-restricted in $[0, 1]$, i.e., $0 \leq \frac{\phi_i(x) - \phi_i(y)}{x - y} \leq 1$ for all $x, y \in \mathbb{R}$

① S. Bai, J. Z. Kolter, and V. Koltun. [Deep equilibrium models](#), NeurIPS, 2019

② L. El Ghaoui, F. Gu, B. Travacca, A. Askari, and A. Y. Tsai. [Implicit deep learning](#). SIMODS, 2019

Implicit Neural Networks

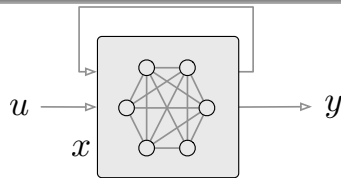
Definition via fixed-point equations



- Feedforward neural networks:

$$x^{i+1} = \Phi(A_i x^i + b_i), \quad x^0 = u$$

$$y = A_k x^k + b_k$$



- Implicit neural networks:

$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

- $\Phi(y_1, \dots, y_n) = (\phi_1(y_1), \dots, \phi_n(y_n))^T$ is a diagonal activation function
- activation functions are slope-restricted in $[0, 1]$, i.e., $0 \leq \frac{\phi_i(x) - \phi_i(y)}{x - y} \leq 1$ for all $x, y \in \mathbb{R}$

Advantages: Representation, Performance, Memory

Implicit Neural Networks

An operator theoretic perspective

Main Questions

$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

- 1 Existence and computation of solutions?
- 2 How to estimate the input-output $x \mapsto u$ robustness?

Implicit Neural Networks

An operator theoretic perspective

Main Questions

$$x = \Phi(Ax + Bu + b)$$

$$y = Cx + c$$

- 1 Existence and computation of solutions?
- 2 How to estimate the input-output $x \mapsto u$ robustness?

Key insight

Fixed-point equation
 $x = \Phi(Ax + Bu + b)$



Operator theory

$$N_u(x) = x - \Phi(Ax + Bu + b)$$

fixed-points



zeros of $N_u(x)$

robustness



sensitivity wrt u

- We can use tools from monotone operator theory to study implicit neural networks

Implicit Neural Networks

Well-posedness and robustness

Main observation

If $\mu_\infty(A) < 1$ then N_u is a monotone operator wrt $\|\cdot\|_\infty$.

⁴SJ, A. Davydov, A. V. Proskurnikov, and F. Bullo. Robust implicit networks via non-Euclidean contractions. In NeurIPs 2021

Implicit Neural Networks

Well-posedness and robustness

Main observation

If $\mu_\infty(A) < 1$ then N_u is a monotone operator wrt $\|\cdot\|_\infty$.

Theorem⁴

If $\mu_\infty(A) < 1$ then

- ❶ $x = \Phi(Ax + Bu + b)$ has a unique solution x_u^*
- ❷ x_u^* can be computed using average iterations for $x = \Phi(Ax + Bu + b)$
- ❸ $\|x_u^* - x_v^*\|_\infty \leq \frac{\|B\|_\infty}{1 - [\mu_\infty(A)]_+} \|u - v\|_\infty$

⁴ SJ, A. Davydov, A. V. Proskurnikov, and F. Bullo. Robust implicit networks via non-Euclidean contractions. In NeurIPs 2021

Main observation

If $\mu_\infty(A) < 1$ then N_u is a monotone operator wrt $\|\cdot\|_\infty$.

Theorem⁴

If $\mu_\infty(A) < 1$ then

- ❶ $x = \Phi(Ax + Bu + b)$ has a unique solution x_u^*
- ❷ x_u^* can be computed using average iterations for $x = \Phi(Ax + Bu + b)$
- ❸ $\|x_u^* - x_v^*\|_\infty \leq \frac{\|B\|_\infty}{1 - [\mu_\infty(A)]_+} \|u - v\|_\infty$

$$u \underbrace{\mapsto}_{\text{Lip}_{u \rightarrow x_u^*}} x^* \underbrace{\mapsto}_{\text{Lip}_{x_u^* \rightarrow y}} y \implies \text{Lip}_{u \rightarrow y} = \text{Lip}_{u \rightarrow x_u^*} \text{Lip}_{x_u^* \rightarrow y}$$

$$\implies \text{Lip}_{u \rightarrow y} = \frac{\|C\|_\infty \|B\|_\infty}{1 - [\mu_\infty(A)]_+}$$

⁴ SJ, A. Davydov, A. V. Proskurnikov, and F. Bullo. Robust implicit networks via non-Euclidean contractions. In NeurIPs 2021

Training Implicit Neural Networks

Promoting robustness

- 1 loss function $\mathcal{L}$
- 2 training data $(\hat{u}_i, \hat{y}_i)_{i=1}^N$

$$\min_{A,B,C,b,c} \sum_{i=1}^N \mathcal{L}(\hat{y}_i, Cx_i + c) \quad + \lambda \quad \text{Lip}_{u \rightarrow y}$$
$$x_i = \Phi(Ax_i + B\hat{u}_i + b)$$
$$\mu_\infty(A) \leq \gamma,$$

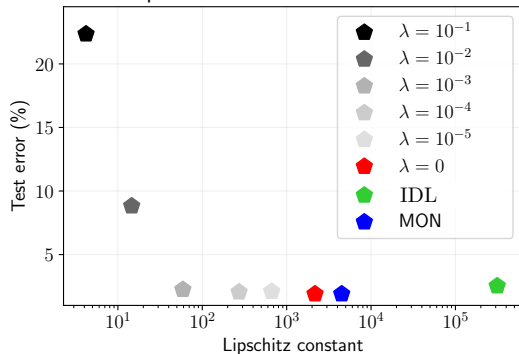
- $\gamma < 1$ is a hyperparameter and $\lambda \geq 0$ is a regularization parameter
- training optimization problem is solved via SGD
- at each step of SGD, $x_i = \Phi(Ax_i + B\hat{u}_i + b)$ is solved using the average-iterations

Numerical Experiments

Lipschitz bound for Implicit Neural Networks

- MNIST dataset: 28×28 pixel handwritten digits between 0 – 9, 60,000 training images and 10,000 test images.
- implicit neural network order: $n = 100$ and $\gamma = 0.95$
- loss function: cross entropy

Test error vs Lipschitz constant on MNIST handwritten digits



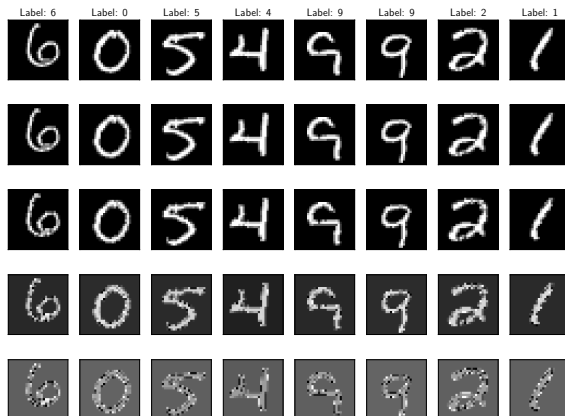
Improvements:

- ($\lambda = 0$): two orders of magnitude wrt. IDL and wrt. MON
- ($\lambda = 10^{-3}$): three orders of magnitude wrt. IDL and one order of magnitude wrt. MON
- ($\lambda = 10^{-2}$): four orders of magnitude wrt. IDL and two orders of magnitude wrt. MON

Numerical Experiments

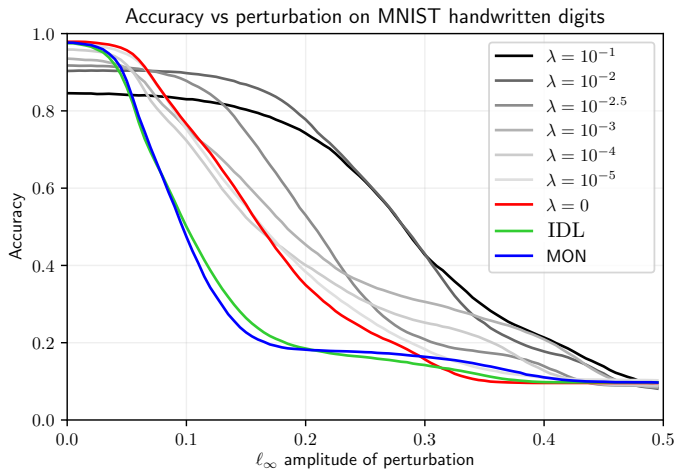
Empirical robustness of INNs

- perturbation: inversion attack $u_{\text{adv}} = u + \epsilon \text{sign}(\frac{1}{2}\mathbb{1}_{784} - u)$



Numerical Experiments

Empirical robustness of INNs

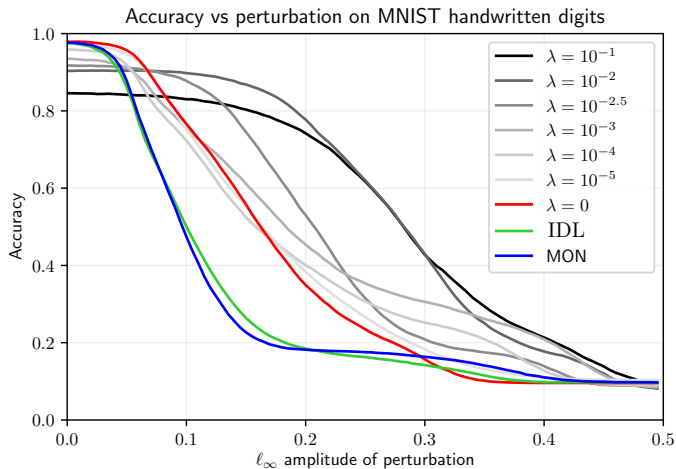


- ($\lambda = 0$): improved robustness than IDL and MON
- ($\lambda > 0$): improved robustness at sizable perturbations but losing some percentage accuracy in clean performance

Tradeoff between **clean performance** and **robustness**

Numerical Experiments

Empirical robustness of INNs

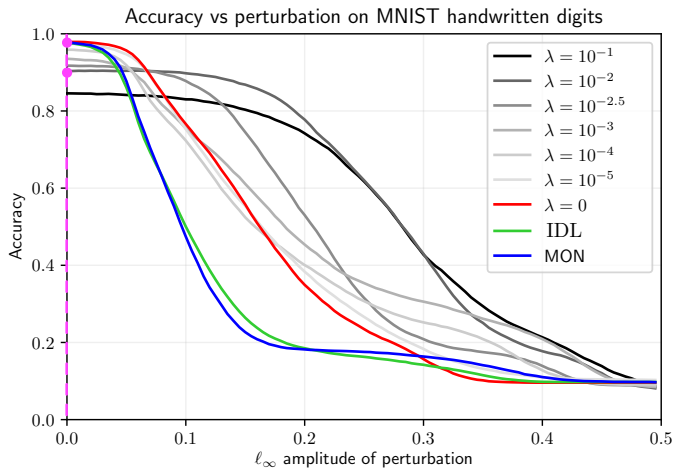


- ($\lambda = 0$): improved robustness than IDL and MON
- ($\lambda > 0$): improved robustness at sizable perturbations but losing some percentage accuracy in clean performance

Tradeoff between **clean performance** and **robustness**

Numerical Experiments

Empirical robustness of INNs

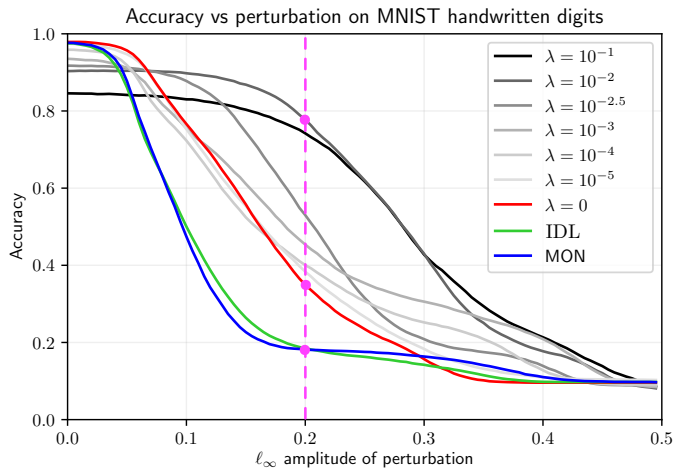


- ($\lambda = 0$): improved robustness than IDL and MON
- ($\lambda > 0$): improved robustness at sizable perturbations but losing some percentage accuracy in clean performance

Tradeoff between **clean performance** and **robustness**

Numerical Experiments

Empirical robustness of INNs



- ($\lambda = 0$): improved robustness than IDL and MON
- ($\lambda > 0$): improved robustness at sizable perturbations but losing some percentage accuracy in clean performance

Tradeoff between **clean performance** and **robustness**

- Extension of monotone operator theory to normed-spaces using Logarithmic norm
- Non-Euclidean contraction theory for well-posedness of INNs
- Lipschitz bounds of INNs using non-Euclidean monotone operator theory

Thank you for your attention!

Back up Slides

Contraction Theory

Logarithmic norm and weak pairings

Differential condition

Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

- Directional derivative of norm $\|\cdot\|$ in direction of A ,

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

¹A. Davydov, **SJ**, F. Bullo, Non-Euclidean contraction theory for robust nonlinear stability, 2022

Contraction Theory

Logarithmic norm and weak pairings

Differential condition

Logarithmic norm

Given a matrix $A \in \mathbb{R}^{n \times n}$ and a norm $\|\cdot\|$:

$$\mu_{\|\cdot\|}(A) := \lim_{h \rightarrow 0^+} \frac{\|I_n + hA\| - 1}{h}$$

- Directional derivative of norm $\|\cdot\|$ in direction of A ,

$$\mu_2(A) = \frac{1}{2} \lambda_{\max}(A + A^\top)$$

$$\mu_1(A) = \max_j (a_{jj} + \sum_{i \neq j} |a_{ij}|)$$

$$\mu_\infty(A) = \max_i (a_{ii} + \sum_{j \neq i} |a_{ij}|)$$

Integral condition

Weak pairing⁵

Given a norm $\|\cdot\|$, the associated weak pairing is $\llbracket \cdot, \cdot \rrbracket : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$:

- Subadditive and weakly homogeneity
- Positive definite
- Cauchy-Schwarz inequality
- $\llbracket x, x \rrbracket = \|x\|^2$

$$\llbracket x, y \rrbracket_2 = y^\top x$$

$$\llbracket x, y \rrbracket_1 = \text{sign}(y)^\top x$$

$$\llbracket x, y \rrbracket_\infty = \max_{i \in I_\infty(x)} x_i y_i$$

$$I_\infty(x) = \{i \mid |x_i| = \|x\|_\infty\}$$

¹A. Davydov, **SJ**, F. Bullo, Non-Euclidean contraction theory for robust nonlinear stability, 2022

Contraction theory

Characterization for non-Euclidean norms

Theorem⁶

$\dot{x} = f(x, u)$ is contracting wrt $\|\cdot\|$ with rate c iff

Differential: $\mu_{\|\cdot\|}(D_x f(x, u)) \leq -c, \quad \text{for all } x, u$

Integral: $\llbracket f(x, u) - f(y, u), x - y \rrbracket \leq -c\|x - y\|^2, \quad \text{for all } x, y, u$

² A. Davydov, S. Jafarpour, F. Bullo, TAC 2022

Contraction theory

Characterization for non-Euclidean norms

Theorem

$\dot{x} = f(x, u)$ is contracting wrt $\|\cdot\|$ with rate c iff

Differential: $\mu_{\|\cdot\|}(D_x f(x, u)) \leq -c, \quad \text{for all } x, u$

Integral: $\llbracket f(x, u) - f(y, u), x - y \rrbracket \leq -c\|x - y\|^2, \quad \text{for all } x, y, u$

- Connection between **contraction theory** and **monotone operator theory**

f is a contracting vector field wrt to $\|\cdot\|_2$
iff

$-f$ is a strongly monotone operator wrt to the inner product $\langle \cdot, \cdot \rangle$.

Contraction theory

Characterization for non-Euclidean norms

Theorem

$\dot{x} = f(x, u)$ is contracting wrt $\|\cdot\|$ with rate c iff

Differential: $\mu_{\|\cdot\|}(D_x f(x, u)) \leq -c, \quad \text{for all } x, u$

Integral: $\llbracket f(x, u) - f(y, u), x - y \rrbracket \leq -c\|x - y\|^2, \quad \text{for all } x, y, u$

- Connection between **contraction theory** and **monotone operator theory**

f is a contracting vector field wrt to $\|\cdot\|$
iff

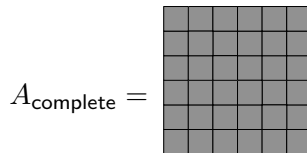
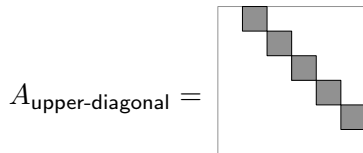
$-f$ is a strongly monotone operator wrt to the weak pairing $\llbracket \cdot, \cdot \rrbracket$.

Implicit neural networks

Origin and motivations

- Origins:
- Generalizing feedforward neural networks to fully-connected synaptic matrices

Intuition: $z^{i+1} = \phi_i(A_i z^i + b_i) \iff z = \Phi(Ax + Bu + b)$, where A has upper diagonal structure.

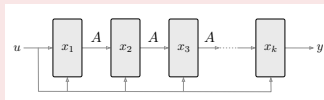


Implicit neural networks

Origin and motivations

- comparable accuracy to traditional neural networks with significant memory reduction

Intuition: implicit neural network = weight-tied infinite-layer network



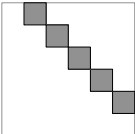
$z^{i+1} = \phi_i(Az^i + B_i x + b_i) \implies \lim_{i \rightarrow \infty} z^i = x^*$ solution to the implicit neural network

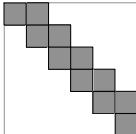
- suitable for learning constrained optimization problems

Intuition: casting KKT condition as an implicit layer

- vanishing and exploding gradient

Intuition: the notion of “autapse” (time-delayed self-feedback) from neuroscience

$$A_{\text{upper-diagonal}} =$$


$$A_{\text{Autapse}} =$$


- suitable for learning stiff problems or problems with discontinuity

Generalized Structure

Comparison with feedforward neural networks

- Feedforward neural networks:

$$z^{(\ell+1)} = \Phi(A_\ell z^{(\ell)} + b_\ell), \quad z^{(0)} = x$$

$$u = A_k z^{(k)} + b_k$$

$$z = \Phi \left(\begin{array}{|c|} \hline \text{Sparse Matrix} \\ \hline \end{array} z + \begin{array}{|c|} \hline \text{Bias Vector} \\ \hline \end{array} x + b \right)$$

$$u = \begin{array}{|c|} \hline \text{Sparse Vector} \\ \hline \end{array} z + b_k$$

- Implicit neural networks:

$$z = \Phi(Az + Bx + b)$$

$$u = Cz + c$$

$$z = \Phi \left(\begin{array}{|c|} \hline \text{Dense Matrix} \\ \hline \end{array} z + \begin{array}{|c|} \hline \text{Dense Vector} \\ \hline \end{array} x + b \right)$$

$$u = \begin{array}{|c|} \hline \text{Dense Vector} \\ \hline \end{array} z + c$$