

Advancing Trustworthy & Efficient AI for Science, Engineering, and Medicine

Ju Sun
Oct 16, 2025

ECE Communications and Signal Processing (CSP) Seminar
University of Michigan



UNIVERSITY OF MINNESOTA
Driven to DiscoverSM

Success of deep learning (DL) not news anymore

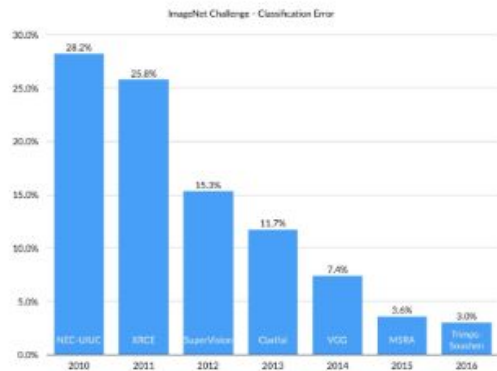


image classification

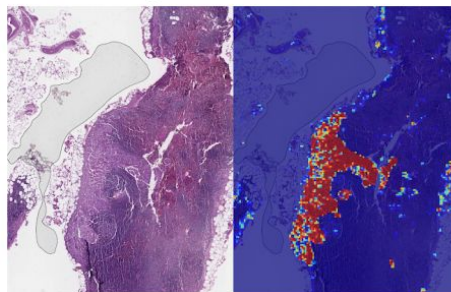


Go game (2017)

Commercial breakthroughs ...



self-driving vehicles credit: wired.com



healthcare credit: Google AI



smart-home devices credit: Amazon



robotics credit: Cornell U.

Robustness issues of DL not news anymore



"panda"

x

δ

FOOLING THE AI

Deep neural networks (DNNs) are brilliant at image recognition — but they can be easily hacked.

These stickers made an artificial-intelligence system read this stop sign as 'speed limit 45'.

Stop



Speed limit 45

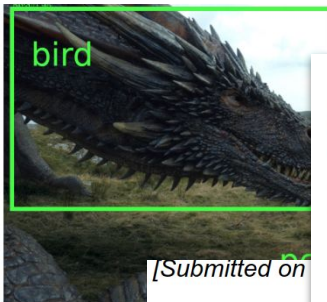


Robustness issues across domains/tasks



"panda"

x



[Submitted on 06/01/2018]

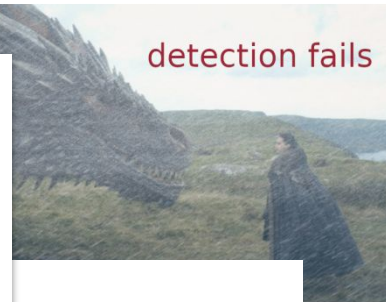
A Multilingual Inputs

Akshay Srinivasan

Adversarial
multilingual
input. Our

man and Hindi. While exact results differ depending on language/datasets, our key findings from these experiments can be summarized as follows:

1. NER models for all three languages are sensitive to adversarial input.
2. Adversarial fine-tuning and re-training could improve the performance of NER models both on original and adversarial test sets, without requiring additional manual labeled data.



Adversarial

we performed a
small perturbations in the
(German and Hindi) are

Name entry Recognition

are not very robust to such changes, as indicated by the fluctuations in the overall F1 score as well as in a more fine-grained evaluation. With that knowledge, we further explored whether it is possible to improve the existing NER

Robustness issues across models

Tutorial

Foundational Robustness of Foundation Models

NeurIPS 2022

Abstract

Foundation models adopting the methodology of deep learning with pre-training on large-scale unlabeled data and finetuning with task-specific supervision are becoming a mainstream technique in machine learning. Although foundation models hold many promises in learning general representations and few-shot/zero-shot generalization across domains and data modalities, at the same time they raise unprecedented challenges and considerable risks in robustness and privacy due to the use of the excessive volume of data and complex neural network architectures. This tutorial aims to deliver a Coursera-like online tutorial containing comprehensive lectures, a hands-on and interactive Jupyter/Colab live coding demo, and a panel discussion on different aspects of trustworthiness in foundation models. More information can be found at <https://sites.google.com/view/neurips2022-frfm-tutorial>

<https://research.ibm.com/publications/foundational-robustness-of-foundation-models>

Trustworthiness issues not news anymore

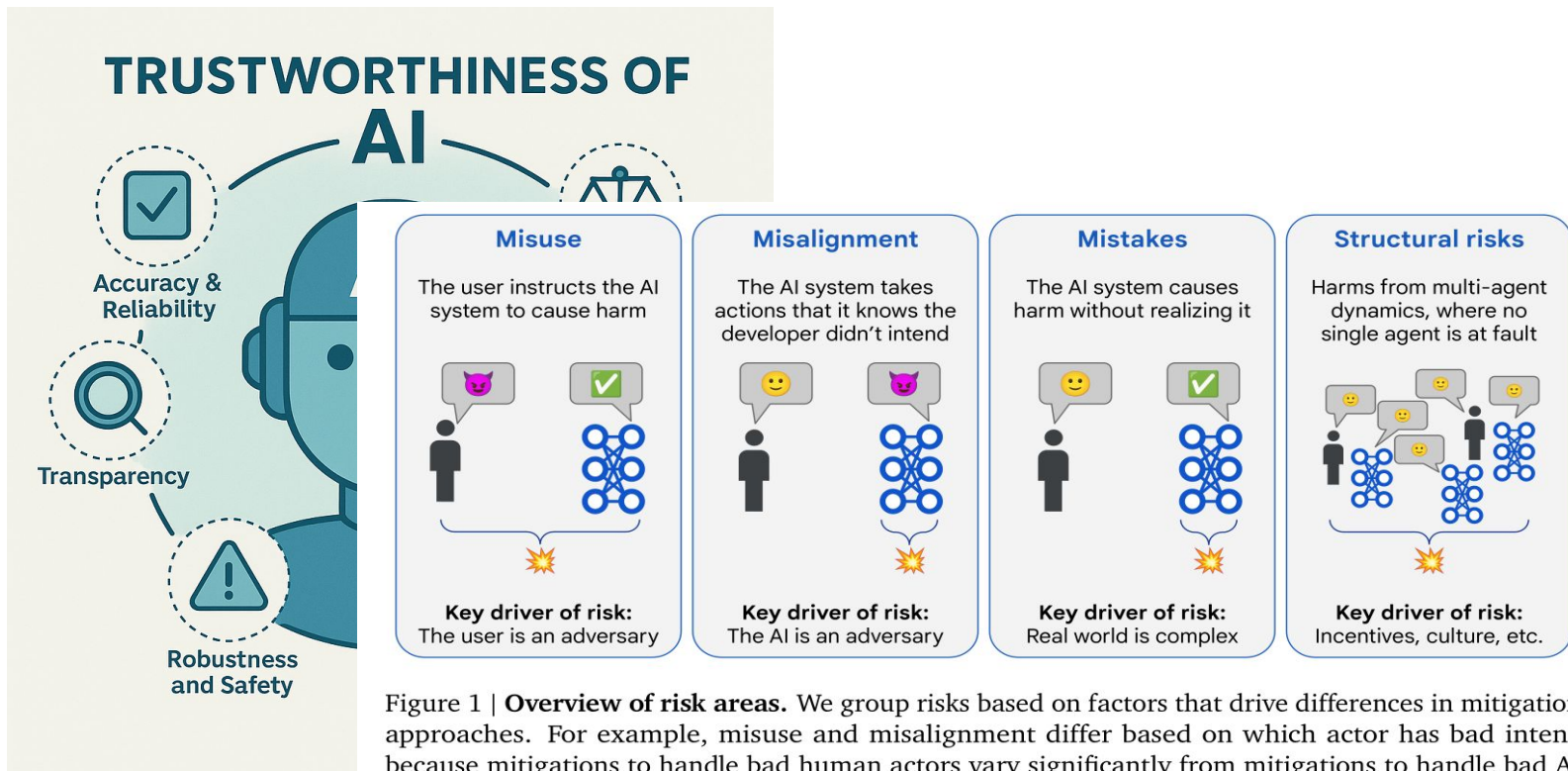


Figure 1 | **Overview of risk areas.** We group risks based on factors that drive differences in mitigation approaches. For example, misuse and misalignment differ based on which actor has bad intent, because mitigations to handle bad human actors vary significantly from mitigations to handle bad AI actors.

source :

<https://deepmind.google/discover/blog/taking-a-responsible-path-t>

International concerns and priorities



Administration Priorities The Record

OCTOBER 30, 2023

FACT SHEET: President Biden Issues Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence

 BRIEFING ROOM STATEMENTS AND RELEASES

Today, President Biden is issuing a landmark Executive Order to ensure that America leads the way in seizing the promise and managing the risks of artificial intelligence (AI). The Executive Order establishes new standards for AI safety and security, protects Americans' privacy, advances equity and civil rights, stands up for consumers and workers, promotes innovation and competition, advances American leadership around the world, and more.

- New Standards for AI Safety and Security
- Protecting Americans' Privacy
- Advancing Equity and Civil Rights
- Standing Up for Consumers, Patients, and Students
- Supporting Workers
- Promoting Innovation and Competition
- Advancing American Leadership Abroad
- Ensuring Responsible and Effective Government Use of AI

<https://www.whitehouse.gov/briefing-room/statements-releases/2023/10/30/fact-sheet-president-biden-issues-executive-order-on-safe-secure-and-trustworthy-artificial-intelligence/>

Three pillars of DL



GPU/TPU/FPGA/...

DATA



Specialized hardware



Specialized software

Key ingredients of DL have been in place for 25-30 years:

Landmark	Emblem	Epoch
Neocognitron	Fukushima	1980
CNN	Le Cun	mid 1980s'
Backprop	Hinton	mid 1980's
SGD	Le Cun, Bengio etc	mid 1990's
Various	Schmidhuber	mid 1980's
<i>CTF</i>	<i>DARPA etc</i>	<i>mid 1980's</i>

CV/NLP domains are lucky

LAION 

Large-scale Artificial Intelligence Open Network

TRULY OPEN AI. 100% NON-PROFIT. 100%

LAION, as a non-profit organization, provides datasets and models to liberate machine learning research. By doing so, we encourage open public education and a more environmentally friendly use of resources by reusing existing datasets and models.

Re-LAION 5B release (30.08.2024)

TABLE 2: Statistics of commonly-used data sources.

Corpora	Size	Source	Latest Update Time
BookCorpus [158]	5GB	Books	Dec-2015
Gutenberg [159]	-	Books	Dec-2021
C4 [82]	800GB	CommonCrawl	Apr-2019
CC-Stories-R [160]	31GB	CommonCrawl	Sep-2019
CC-NEWS [27]	78GB	CommonCrawl	Feb-2019
REALNEWS [161]	120GB	CommonCrawl	Apr-2019
OpenWebText [162]	38GB	Reddit links	Mar-2023
Pushift.io [163]	2TB	Reddit links	Mar-2023
Wikipedia [164]	21GB	Wikipedia	Mar-2023
BigQuery [165]	-	Codes	Mar-2023
the Pile [166]	800GB	Other	Dec-2020
ROOTS [167]	1.6TB	Other	Jun-2022

source: <https://arxiv.org/abs/2303.18223>

Not all fields are as lucky

Thrust B: How Should Domain Knowledge Be Incorporated into Supervised Machine Learning?

The central question for this thrust is “which knowledge should be leveraged in SciML, and how should this knowledge be included?” Any answers will naturally depend on the SciML task and computational budgets, thus mirroring standard considerations in traditional scientific computing.

Hard Constraints. One research avenue involves incorporation of domain knowledge through imposition of constraints that cannot be violated. These hard constraints could be enforced during training, replacing what typically is an unconstrained optimization problem with a constrained one. In general, such constraints could involve simulations or highly nonlinear functions of the training parameters. Therefore, there is a need to identify particular cases when constraint qualification conditions can be ensured as these conditions are necessary regularity conditions for constrained optimization [57–59]. Although incorporating constraints during training generally makes maximal use of training data, there may be additional opportunities to employ constraints at the time of prediction (e.g., by projecting predictions onto the region induced by the constraints).

Soft Constraints. A similar avenue for incorporating domain knowledge involves modifying the objective function (soft constraints) used in training. It is understood that ML loss function selection should be guided by the task and data. Therefore, opportunities exist for developing loss functions that incorporate domain knowledge and analyzing the resulting impact on solvability

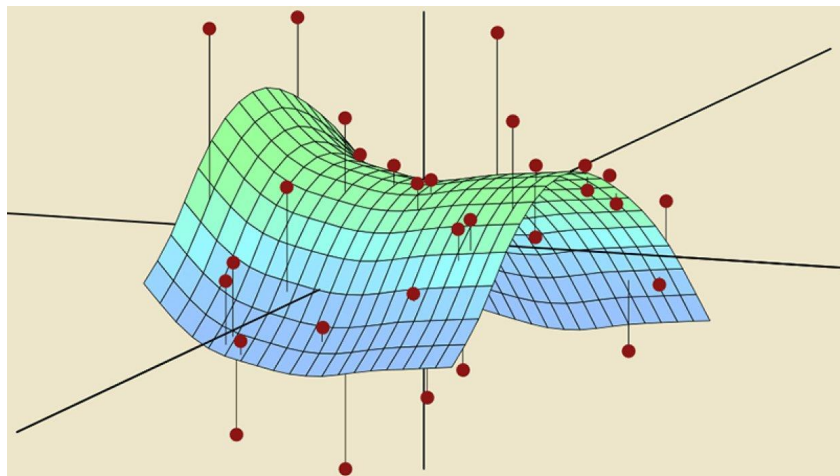


Ref <https://www.osti.gov/servlets/purl/1478744>

Domain-Aware Scientific Machine Learning

There's no free lunch!

Supervised learning as data fitting



Typically, #data points we need grow **exponentially** with respect to dimension (i.e., **curse of dimensionality**)

Knowledge

Small-data AI

Building in prior knowledge is **crucial** for reducing the data complexity
e.g., “convolutional” layers



Large-data AI

Data

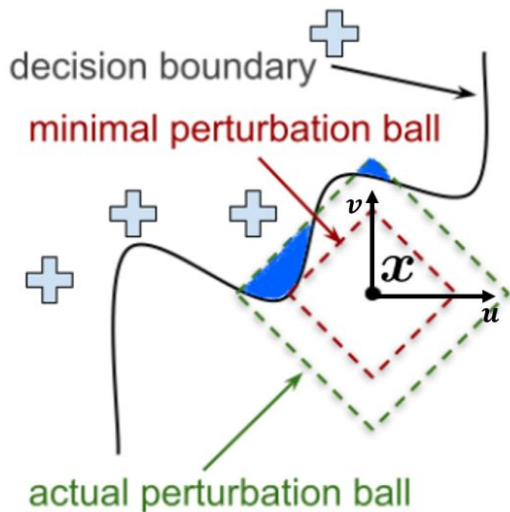
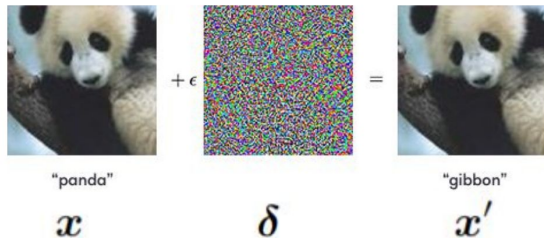
Today's talk: **toward trustworthy and efficient AI**

- **Robustness evaluation and selective classification**
- **Inverse problems with pretrained diffusion models**
- **AI for healthcare: handling data imbalance**

Today's talk: **toward trustworthy and efficient AI**

- **Robustness evaluation and selective classification**
- **Inverse problems with pretrained diffusion models**
- **AI for healthcare: handling data imbalance**

Robustness evaluation (RE)



$$\begin{aligned} & \max_{x'} \ell(y, f_{\theta}(x')) && \text{Maximize loss/error function} \\ \text{s. t. } & d(x, x') \leq \epsilon, && \text{Allowable perturbation} \\ & x' \in [0, 1]^n && \text{Valid image} \end{aligned}$$

$$\begin{aligned} & \min_{x'} d(x, x') && \text{Find robustness radius} \\ \text{s. t. } & \max_{i \neq y} f_{\theta}^i(x') \geq f_{\theta}^y(x'), && \text{On the decision boundary} \\ & x' \in [0, 1]^n && \text{Valid image} \end{aligned}$$

Report robust accuracy over an evaluation set

Constrained optimization problems

$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\boldsymbol{\theta}}(\mathbf{x}')) \\ \text{s. t. } & d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

$$\begin{aligned} & \min_{\mathbf{x}'} d(\mathbf{x}, \mathbf{x}') \\ \text{s. t. } & \max_{i \neq y} f_{\boldsymbol{\theta}}^i(\mathbf{x}') \geq f_{\boldsymbol{\theta}}^y(\mathbf{x}'), \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

Both objective and constraint functions are **nonconvex** in general, e.g., when containing DL models

Projected gradient descent (PGD) for RE

$$\max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}'))$$

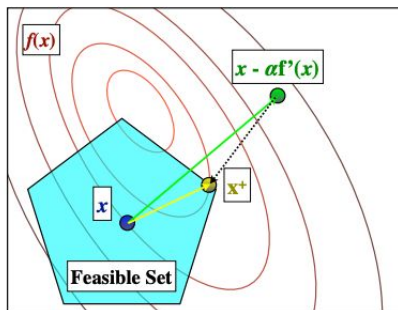
$$\text{s. t. } d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n$$

$$\min_{\mathbf{x} \in \mathcal{Q}} f(\mathbf{x})$$

$$\mathbf{x}_{k+1} = P_{\mathcal{Q}}\left(\mathbf{x}_k - \alpha_k \nabla f(\mathbf{x}_k)\right)$$

Step size

$$P_{\mathcal{Q}}(\mathbf{x}_0) = \arg \min_{\mathbf{x} \in \mathcal{Q}} \frac{1}{2} \|\mathbf{x} - \mathbf{x}_0\|_2^2 \quad \text{Projection operator}$$



Key hyperparameters:

- (1) step size
- (2) iteration number

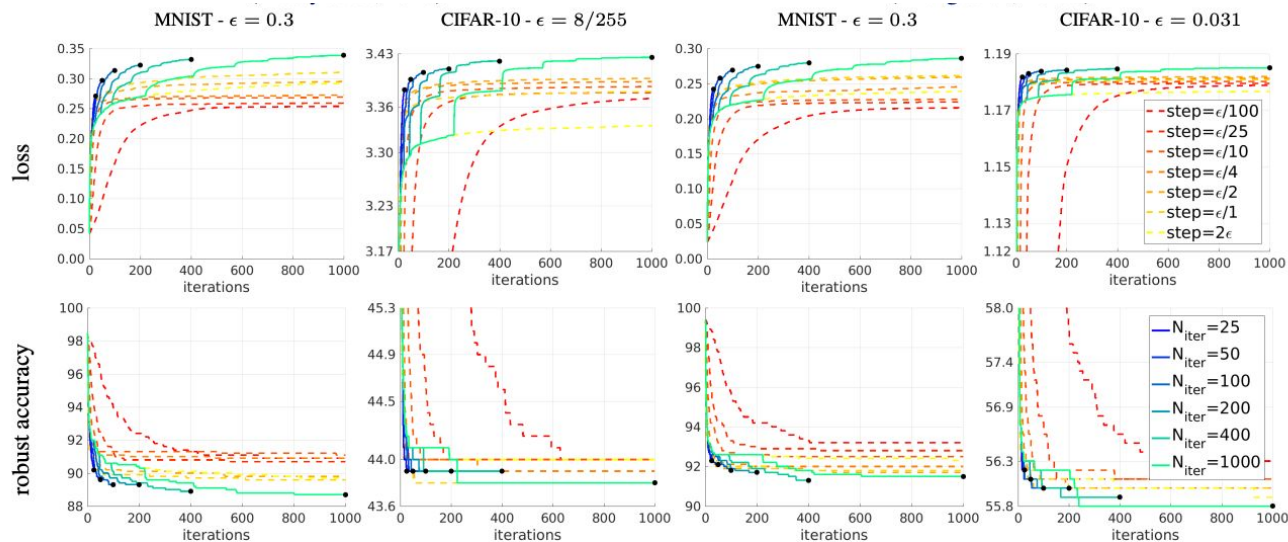
Algorithm 1 APGD

```

1: Input:  $f, S, x^{(0)}, \eta, N_{\text{iter}}, W = \{w_0, \dots, w_n\}$ 
2: Output:  $x_{\max}, f_{\max}$ 
3:  $x^{(1)} \leftarrow P_S(x^{(0)} + \eta \nabla f(x^{(0)}))$ 
4:  $f_{\max} \leftarrow \max\{f(x^{(0)}), f(x^{(1)})\}$ 
5:  $x_{\max} \leftarrow x^{(0)}$  if  $f_{\max} \equiv f(x^{(0)})$  else  $x_{\max} \leftarrow x^{(1)}$ 
6: for  $k = 1$  to  $N_{\text{iter}} - 1$  do
7:    $z^{(k+1)} \leftarrow P_S(x^{(k)} + \eta \nabla f(x^{(k)}))$ 
8:    $x^{(k+1)} \leftarrow P_S(x^{(k)} + \alpha(z^{(k+1)} - x^{(k)}))$ 
       $+ (1 - \alpha)(x^{(k)} - x^{(k-1)})$ 
9:   if  $f(x^{(k+1)}) > f_{\max}$  then
10:      $x_{\max} \leftarrow x^{(k+1)}$  and  $f_{\max} \leftarrow f(x^{(k+1)})$ 
11:   end if
12:   if  $k \in W$  then
13:     if Condition 1 or Condition 2 then
14:        $\eta \leftarrow \eta/2$  and  $x^{(k+1)} \leftarrow x_{\max}$ 
15:     end if
16:   end if
17: end for

```

Problem with projected gradient descent



$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}')) \\ & \text{s.t. } d(\mathbf{x}, \mathbf{x}') \leq \epsilon, \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

Tricky to set:
iteration number & step size
i.e., **tricky to decide where to stop**

Penalty methods for complicated d

$$\max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}'))$$

$$\text{s. t. } d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n$$

$$d(\mathbf{x}, \mathbf{x}') \doteq \|\phi(\mathbf{x}) - \phi(\mathbf{x}')\|_2$$

**perceptual
distance**

$$\text{where } \phi(\mathbf{x}) \doteq [\hat{g}_1(\mathbf{x}), \dots, \hat{g}_L(\mathbf{x})]$$

Projection onto the constraint is complicated

Penalty methods

$$\max_{\tilde{\mathbf{x}}} \mathcal{L}(f(\tilde{\mathbf{x}}), y) - \lambda \max\left(0, \|\phi(\tilde{\mathbf{x}}) - \phi(\mathbf{x})\|_2 - \epsilon\right)$$

Solve it for each fixed λ and then increase λ

Algorithm 2 Lagrangian Perceptual Attack (LPA)

```
1: procedure LPA(classifier network  $f(\cdot)$ , LPIPS distance  $d(\cdot, \cdot)$ , input  $\mathbf{x}$ , label  $y$ , bound  $\epsilon$ )
2:    $\lambda \leftarrow 0.01$ 
3:    $\tilde{\mathbf{x}} \leftarrow \mathbf{x} + 0.01 * \mathcal{N}(0, 1)$   $\triangleright$  initialize perturbations with random Gaussian noise
4:   for  $i$  in  $1, \dots, S$  do  $\triangleright$  we use  $S = 5$  iterations to search for the best value of  $\lambda$ 
5:     for  $t$  in  $1, \dots, T$  do  $\triangleright T$  is the number of steps
6:        $\Delta \leftarrow \nabla_{\tilde{\mathbf{x}}} [\mathcal{L}(f(\tilde{\mathbf{x}}), y) - \lambda \max(0, d(\tilde{\mathbf{x}}, \mathbf{x}) - \epsilon)]$   $\triangleright$  take the gradient of (5)
7:        $\hat{\Delta} = \Delta / \|\Delta\|_2$   $\triangleright$  normalize the gradient
8:        $\eta = \epsilon * (0.1)^{t/T}$   $\triangleright$  the step size  $\eta$  decays exponentially
9:        $m \leftarrow d(\tilde{\mathbf{x}}, \tilde{\mathbf{x}} + h\hat{\Delta})/h$   $\triangleright m \approx$  derivative of  $d(\tilde{\mathbf{x}}, \cdot)$  in the direction of  $\hat{\Delta}$ ;  $h = 0.1$ 
10:       $\tilde{\mathbf{x}} \leftarrow \tilde{\mathbf{x}} + (\eta/m)\hat{\Delta}$   $\triangleright$  take a step of size  $\eta$  in LPIPS distance
11:    end for
12:    if  $d(\tilde{\mathbf{x}}, \mathbf{x}) > \epsilon$  then
13:       $\lambda \leftarrow 10\lambda$   $\triangleright$  increase  $\lambda$  if the attack goes outside the bound
14:    end if
15:  end for
16:   $\tilde{\mathbf{x}} \leftarrow \text{PROJECT}(d, \tilde{\mathbf{x}}, \mathbf{x}, \epsilon)$ 
17:  return  $\tilde{\mathbf{x}}$ 
18: end procedure
```

Problem with penalty methods

Method	cross-entropy loss		margin loss	
	Viol. (%) ↓	Att. Succ. (%) ↑	Viol. (%) ↓	Att. Succ. (%) ↑
Fast-LPA	73.8	3.54	41.6	56.8
LPA	0.00	80.5	0.00	97.0
PPGD	5.44	25.5	0.00	38.5
PWCF (ours)	0.62	93.6	0.00	100

$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}')) \\ \text{s. t. } & d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n \\ & d(\mathbf{x}, \mathbf{x}') \doteq \|\phi(\mathbf{x}) - \phi(\mathbf{x}')\|_2 \\ \text{where } & \phi(\mathbf{x}) \doteq [\hat{g}_1(\mathbf{x}), \dots, \hat{g}_L(\mathbf{x})] \end{aligned}$$

LPA, Fast-LPA: penalty methods

PPGD: Projected gradient descent

PWCF, an optimizer with a principled stopping criterion on **stationarity** & **feasibility**

Penalty methods tend to encounter

large constraint violation (i.e., infeasible solution, known in optimization theory) or **suboptimal solution**

Unreliable optimization
=Unreliable RE

Issues and answers

projected gradient descent

$$\min_{\mathbf{x} \in \mathcal{Q}} f(\mathbf{x})$$

$$\mathbf{x}_{k+1} = P_{\mathcal{Q}}\left(\mathbf{x}_k - \alpha_k \nabla f(\mathbf{x}_k)\right)$$

Issue: no principled stopping criterion
/step size rules

penalty methods

$$\min_{\mathbf{x}} f(\mathbf{x}) \quad \text{s. t. } g(\mathbf{x}) \leq 0$$

$$\min_{\mathbf{x}} f(\mathbf{x}) + \lambda \max(0, g(\mathbf{x}))$$

Solved with increasing λ sequence

Issue: infeasible solution

- Feasible & stationary solution **Stationarity and feasibility check:**
KKT condition
- Reasonable speed **Line search & 2nd order methods**
- A hidden problem: nonsmoothness

A principled solver for
constrained, nonconvex,
nonsmooth problems



Nonconvex, nonsmooth, constrained

$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x}), \quad \text{s.t. } c_i(\mathbf{x}) \leq 0, \forall i \in \mathcal{I}; \quad c_i(\mathbf{x}) = 0, \forall i \in \mathcal{E}.$$

**Penalty sequential quadratic programming
(P-SQP)**

$$\begin{aligned} \min_{d \in \mathbb{R}^n, s \in \mathbb{R}^p} \quad & \mu(f(x_k) + \nabla f(x_k)^\top d) + e^\top s + \frac{1}{2} d^\top H_k d \\ \text{s.t.} \quad & c(x_k) + \nabla c(x_k)^\top d \leq s, \quad s \geq 0, \end{aligned}$$

Advantage: 2nd order method (BFGS) $\rightarrow$ high-precision solution

Principled line search, stationarity/feasibility check

Ref Curtis, Frank E., Tim Mitchell, and Michael L. Overton. "A BFGS-SQP method for nonsmooth, nonconvex, constrained optimization and its evaluation using relative minimization profiles." Optimization Methods and Software 32.1 (2017): 148-181.

Our PyGranso (and NCVX framework)

<https://ncvx.org/>

GRAN~~SO~~SO + PyTorch



NCVX PyGRANSO
Documentation

Search the docs ...

Introduction

Installation

Sett

Exa

Home

<https://ncvx.org/>



$$\min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x}), \text{ s.t. } c_i(\mathbf{x}) \leq 0, \forall i \in \mathcal{I}; c_i(\mathbf{x}) = 0, \forall i \in \mathcal{E}$$

- First general-purpose solver for hard-constrained DL problems
- Recently updated to be compatible with PyTorch 2.6

Strategies to speed up PyGranso for RE

$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}')) \\ \text{s. t. } & d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

Constraint folding: many constraints into few

$$\begin{aligned} h_j(\mathbf{x}) = 0 & \iff |h_j(\mathbf{x})| \leq 0, \\ c_i(\mathbf{x}) \leq 0 & \iff \max\{c_i(\mathbf{x}), 0\} \leq 0, \\ \mathcal{F}(|h_1(\mathbf{x})|, \dots, |h_i(\mathbf{x})|, \max\{c_1(\mathbf{x}), 0\}, \\ & \dots, \max\{c_j(\mathbf{x}), 0\}) \leq 0, \end{aligned}$$

$$\begin{aligned} & \min_{\mathbf{x}'} d(\mathbf{x}, \mathbf{x}') \\ \text{s. t. } & \max_{i \neq y} f_{\theta}^i(\mathbf{x}') \geq f_{\theta}^y(\mathbf{x}'), \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

Two-stage optimization

1. **Stage 1 (selecting the best initialization):** Optimize the problems by PWCF with R different random initialization $\mathbf{x}^{(r,0)}$ for k iterations, where $r = 1, \dots, R$, and collect the final first-stage solution $\mathbf{x}^{(r,k)}$ for each run. Determine the best intermediate result $\mathbf{x}^{(*,k)}$ following [Algorithm 1](#).
2. **Stage 2 (optimization):** Warm start the optimization process with $\mathbf{x}^{*,k}$ until the stopping criterion is met [7](#) (i.e., reaching both the stationarity and feasibility tolerance, or reaching the MaxIter K).

First general-purpose, reliable solver for RE

$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}')) \\ \text{s. t. } & d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

$$\begin{aligned} & \min_{\mathbf{x}'} d(\mathbf{x}, \mathbf{x}') \\ \text{s. t. } & \max_{i \neq y} f_{\theta}^i(\mathbf{x}') \geq f_{\theta}^y(\mathbf{x}'), \quad \mathbf{x}' \in [0, 1]^n \end{aligned}$$

Reliability

- SOTA methods  **ROBUSTBENCH**
A standardized benchmark for adversarial robustness
No stopping criterion (only use maxit); step size scheduler
- PWCF (ours)
Principled line-search criterion and termination condition

Generality

- SOTA methods
Can mostly only handle several lp metrics (l1, l2, linf)
- PWCF (ours)
Any differentiable metrics and both min and max forms
E.g., perceptual distance $d(\mathbf{x}, \mathbf{x}') \doteq \|\phi(\mathbf{x}) - \phi(\mathbf{x}')\|_2$
where $\phi(\mathbf{x}) \doteq [\hat{g}_1(\mathbf{x}), \dots, \hat{g}_L(\mathbf{x})]$

A quick example

$$\begin{aligned} & \max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}')) \\ \text{s. t. } & d(\mathbf{x}, \mathbf{x}') \leq \varepsilon, \quad \mathbf{x}' \in [0, 1]^n \end{aligned} \quad \text{where} \quad \begin{aligned} & d(\mathbf{x}, \mathbf{x}') \doteq \|\phi(\mathbf{x}) - \phi(\mathbf{x}')\|_2 \\ & \phi(\mathbf{x}) \doteq [\hat{g}_1(\mathbf{x}), \dots, \hat{g}_L(\mathbf{x})] \end{aligned}$$

Method	cross-entropy loss		margin loss	
	Viol. (%) ↓	Att. Succ. (%) ↑	Viol. (%) ↓	Att. Succ. (%) ↑
Fast-LPA	73.8	3.54	41.6	56.8
LPA	0.00	80.5	0.00	97.0
PPGD	5.44	25.5	0.00	38.5
PWCF (ours)	0.62	93.6	0.00	100

More details

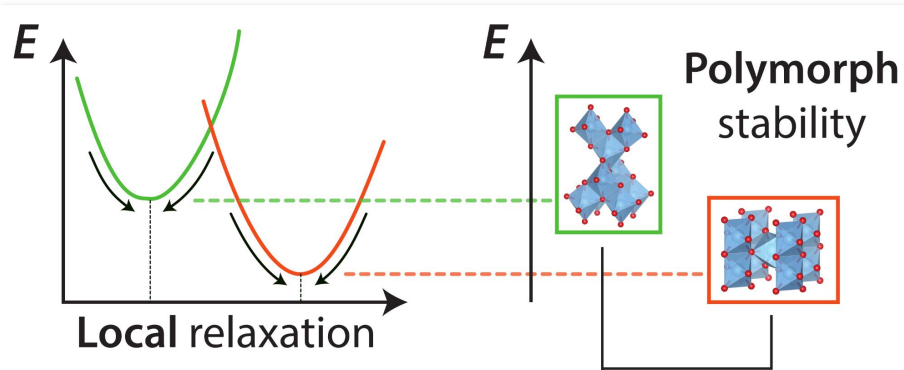
[Submitted on 23 Mar 2023]

Optimization and Optimizers for Adversarial Robustness

Hengyue Liang, Buyun Liang, Le Peng, Ying Cui, Tim Mitchell, Ju Sun

Empirical robustness evaluation (RE) of deep learning models against adversarial perturbations entails solving nontrivial constrained optimization problems. Existing numerical algorithms that are commonly used to solve them in practice predominantly rely on projected gradient, and mostly handle perturbations modeled by the ℓ_1 , ℓ_2 and ℓ_∞ distances. In this paper, we introduce a novel algorithmic framework that blends a general-purpose constrained-optimization solver PyGRANSO with Constraint Folding (PWCF), which can add more reliability and generality to the state-of-the-art RE packages, e.g., AutoAttack. Regarding reliability, PWCF provides solutions with stationarity measures and feasibility tests to assess the solution quality. For generality, PWCF can handle perturbation models that are typically inaccessible to the existing projected gradient methods; the main requirement is the distance metric to be almost everywhere differentiable. Taking advantage of PWCF and other existing numerical algorithms, we further explore the distinct patterns in the solutions found for solving these optimization problems using various combinations of losses, perturbation models, and optimization algorithms. We then discuss the implications of these patterns on the current robustness evaluation and adversarial training.

Subjects: Machine Learning (cs.LG); Computer Vision and Pattern Recognition (cs.CV)



Enabling

Robustness evaluation

$$\max_{\mathbf{x}'} \ell(\mathbf{y}, f_{\theta}(\mathbf{x}'))$$

$$\text{s. t. } \mathbf{x}' \in \Delta(\mathbf{x}) = \{\mathbf{x}' \in [0, 1]^n : d(\mathbf{x}, \mathbf{x}') \leq \varepsilon\}$$

$$\min_{\mathbf{x}' \in [0, 1]^n} d(\mathbf{x}, \mathbf{x}') \quad \text{s. t. } \max_{\ell \neq y} f_{\theta}^{\ell}(\mathbf{x}') \geq f_{\theta}^y(\mathbf{x}')$$

Stability
respe
phase se

- **(W-CSTR-T) Weak explicit constraints informed by phase transition:** For each composition, the energy/atom of any non-ground-state polymorph is greater than that of the ground-state polymorph
- **(W-CSTR-S) Weak explicit constraints informed by phase separation:** For each chemical space, the energy/atom of any TUS material is above the lower convex envelope in the composition-energy space

Settings

Examples

$$\min_{\theta} \frac{1}{N} \sum_{i=1}^N \ell(\mathbf{y}_i, f_{\theta}(\mathbf{x}_i)) + \Omega(\theta) \quad \text{s. t. } f_{\theta}(\mathbf{x}_p) \geq f_{\theta}(\mathbf{x}_q) \quad \forall (p, q) \in \mathcal{O}_{\text{W-CSTR-T}},$$

NCVX: A General-Purpose Optimization Solver for Constrained Machine and Deep Learning

Buyun Liang, Tim Mitchell, Ju Sun

$$\text{s. t. } \frac{\sum_{i=1}^N \mathbb{1}\{\mathbf{y}_i = +1\} \mathbb{1}\{f_{\theta}(\mathbf{x}_i) > t\}}{\sum_{i=1}^N \mathbb{1}\{\mathbf{y}_i = +1\}} \geq \alpha$$

Yash, Le, Zhong, Buyun

ML models are not perfect

$(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \sim_{iid} \mathcal{D}_{\mathcal{X} \times \mathcal{Y}}$ on $\mathcal{X} \times \mathcal{Y}$.

$$\hat{R}_S(f) \doteq \frac{1}{N} \sum_{i \in [N]} \mathbb{1} \{f(\mathbf{x}_i) \neq y_i\}$$

$$R(f) \doteq \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}_{\mathcal{X} \times \mathcal{Y}}} \mathbb{1} \{f(\mathbf{x}) \neq y\}$$

$$h_* \in \arg \min_{h \text{ "reasonable"}} R(h).$$

Bayes optimal classifier for $\arg \max_{y \in \{+1, -1\}} \mathbb{P}[y|\mathbf{x}]$

binary classification

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathcal{X}}} \min(\mathbb{P}[1|\mathbf{x}], \mathbb{P}[-1|\mathbf{x}]) \in [0, 1/2]$$

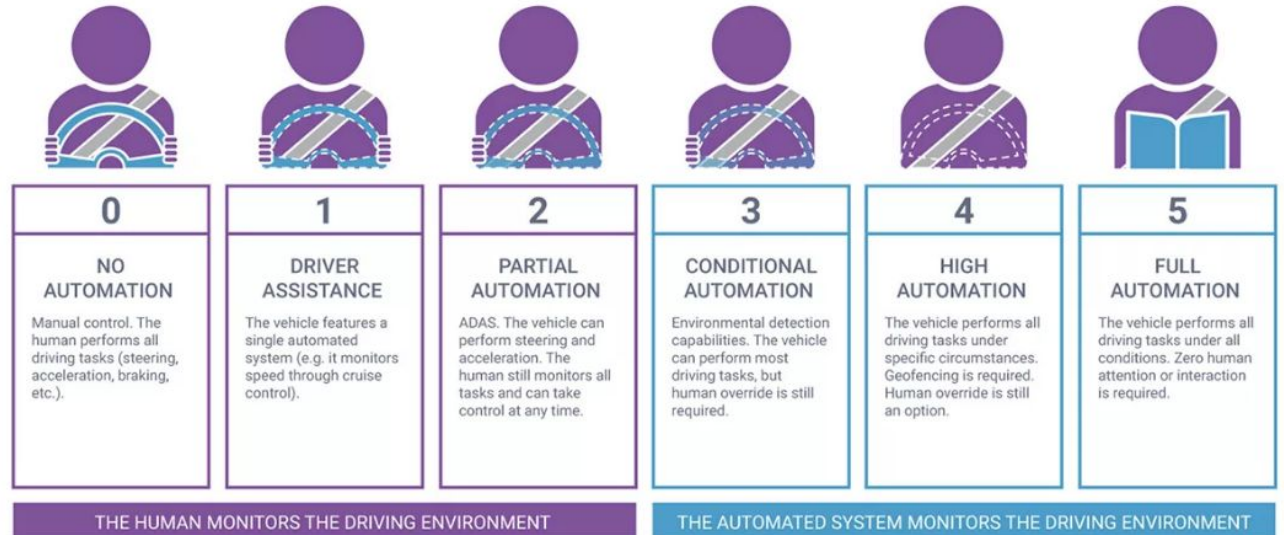
100% often not achievable even if with infinite amount of training data

Imperfect AI models can still be deployed



SYNOPSYS®

LEVELS OF DRIVING AUTOMATION



A crucial component: allowing AI to restrain itself

predictor $f : \mathcal{X} \rightarrow \mathbb{R}^K$ selector $g : \mathcal{X} \rightarrow \{0, 1\}$

$$(f, g)(\mathbf{x}) \triangleq \begin{cases} f(\mathbf{x}) & \text{if } g(\mathbf{x}) = 1; \\ \text{abstain} & \text{if } g(\mathbf{x}) = 0. \end{cases}$$

No prediction on uncertain samples and defer them to humans

$$g_\gamma(\mathbf{x}) = \mathbb{1}[s(\mathbf{x}) > \gamma]$$

Typically, selection by thresholding prediction confidence

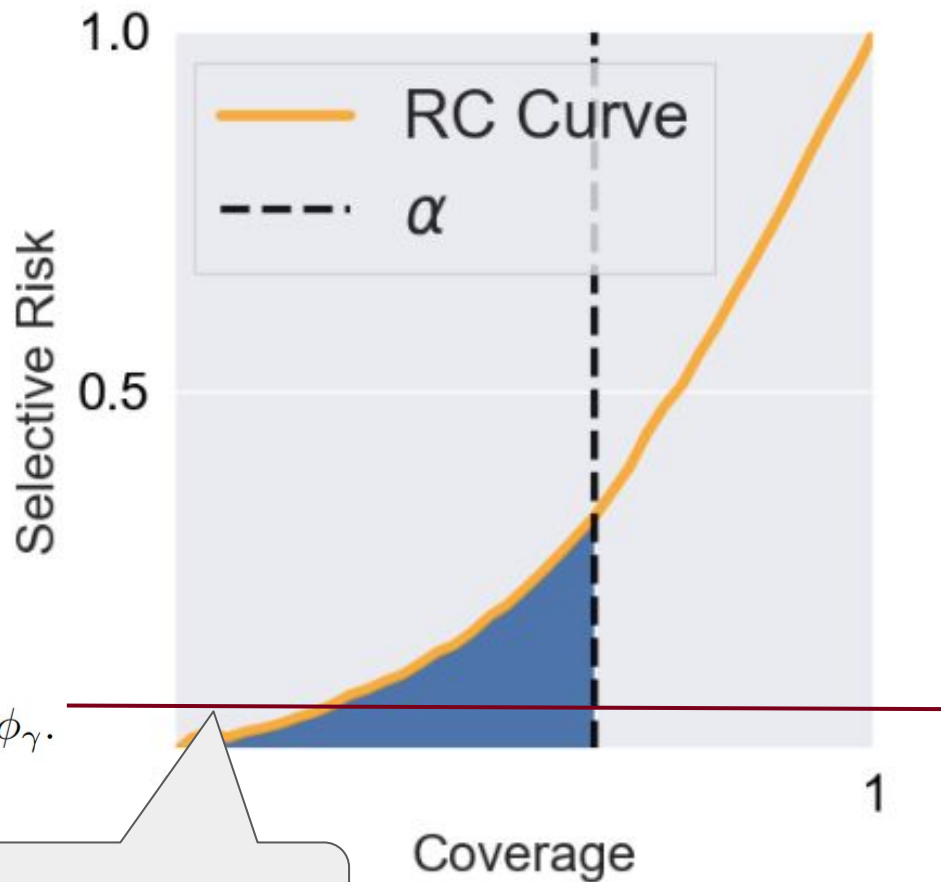
Risk-coverage tradeoff

$$(f, g)(\mathbf{x}) \triangleq \begin{cases} f(\mathbf{x}) & \text{if } g(\mathbf{x}) = 1; \\ \text{abstain} & \text{if } g(\mathbf{x}) = 0. \end{cases}$$

$$g_\gamma(\mathbf{x}) = \mathbb{1}[s(\mathbf{x}) > \gamma]$$

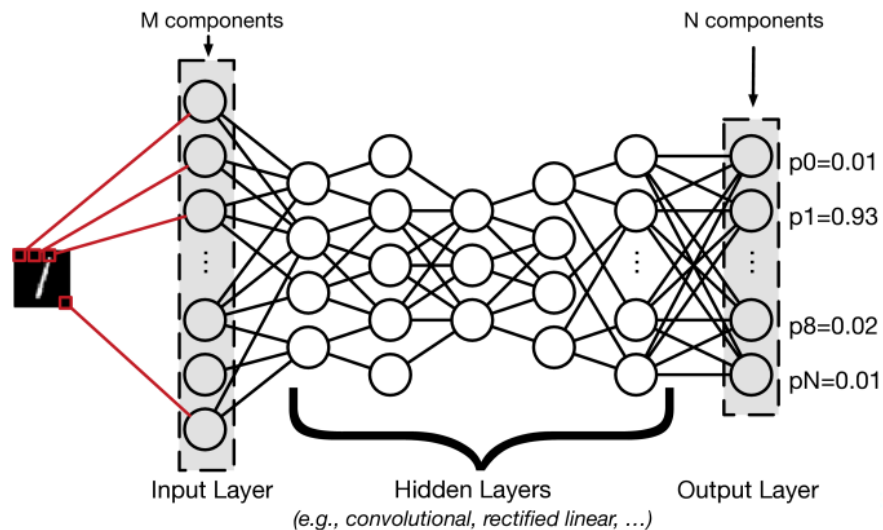
(coverage) $\phi_\gamma = \mathbb{E}_{\mathcal{D}}[g_\gamma(\mathbf{x})]$,

(selection risk) $R_\gamma = \mathbb{E}_{\mathcal{D}}[\ell(f(\mathbf{x}), y)g_\gamma(\mathbf{x})]/\phi_\gamma$.



High-stakes corner

Which confidence score?



$$\approx p(y = 1|x)$$

$\mathbf{z} \in \mathbb{R}^K$ contains the raw logits (RLs)

$$SR_{\max} \triangleq \max_i \sigma(z^i),$$

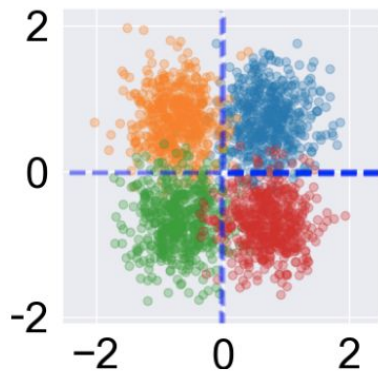
$$SR_{\text{doctor}} \triangleq (\|\sigma(\mathbf{z})\|_2^2 - 1) / \|\sigma(\mathbf{z})\|_2^2 = 1 - \|\sigma(\mathbf{z})\|_1 / \|\sigma(\mathbf{z})\|_2^2,$$

$$SR_{\text{ent}} \triangleq \sum_i \sigma(z^i) \log \sigma(z^i),$$

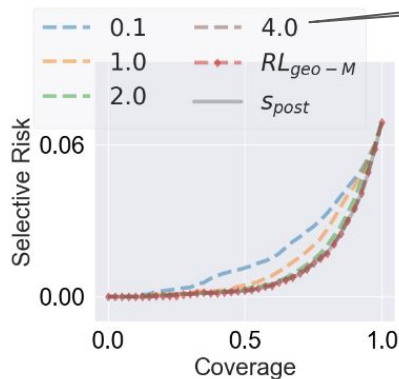
But are they good scores?

$z \in \mathbb{R}^K$ contains the raw logits (RLs)

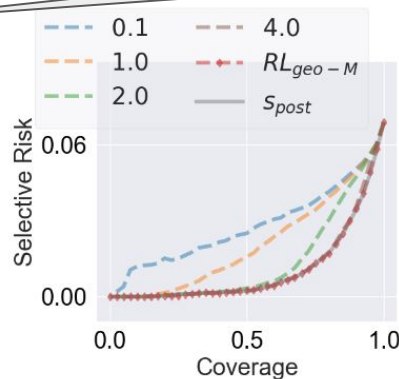
Scale factor applied to RLs



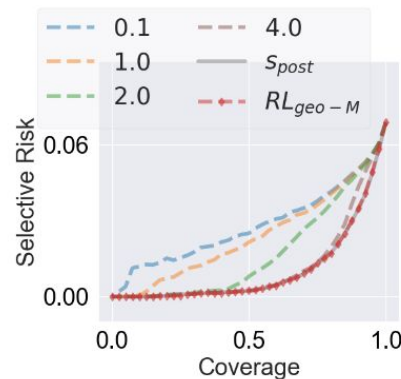
(a) Sample visualization



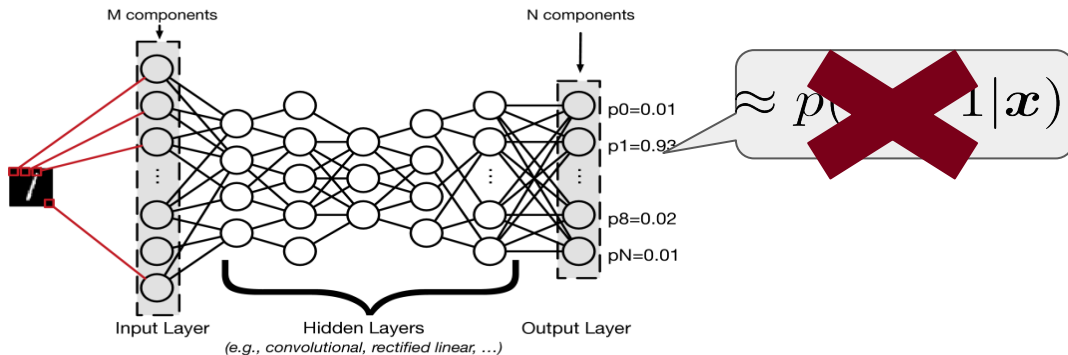
(b) $SR_{\max}$



(c) SR_{doctor}



(d) SR_{ent}



Calibration: align the outputs with the true posterior probs

Our margin-based scores

Signed dist to the separating hyperplane

Binary SVMs: $f(x) = w^\top x + b$

Geometric margin: $y(w^\top x + b) / \|w\|_2$

Multiclass SVMs: $f(x) = W^\top x + b$

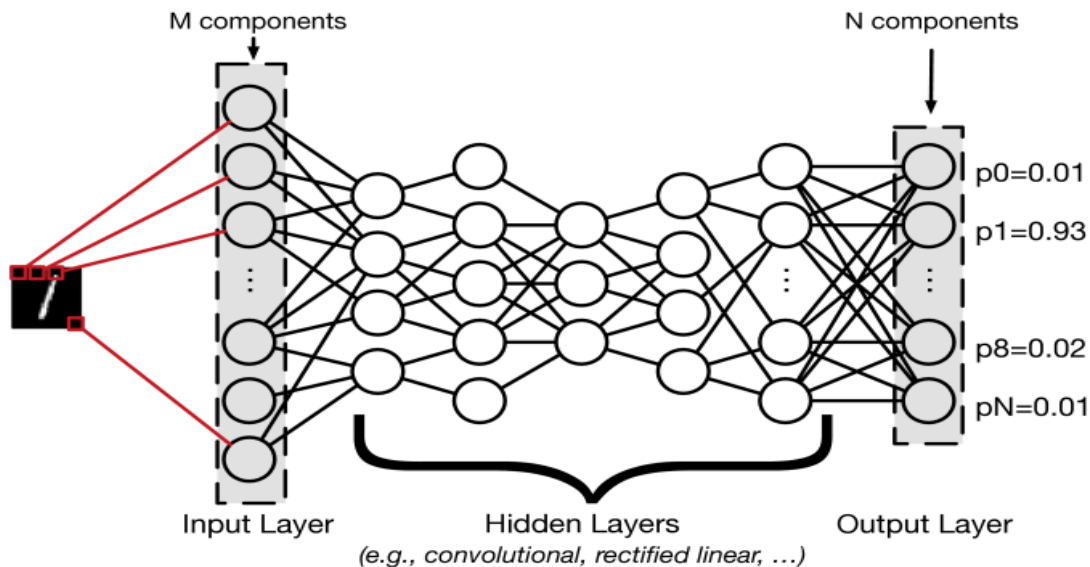
Geometric margin: $\frac{w_{y'}^\top x + b_{y'}}{\|w_{y'}\|_2} - \max_{j \in \{1, \dots, K\} \setminus y'} \frac{w_j^\top x + b_j}{\|w_j\|_2}$

Confidence margin: $(w_{y'}^\top x + b_{y'}) - \max_{i \in \{1, \dots, K\} \setminus y'} (w_i^\top x + b_i)$

Difference of dists between the two nearest hyperplanes

These scores are not affected by the logit scaling

Our margin-based scores



Geometric margin:

$$\frac{\mathbf{w}_{y'}^T \mathbf{x} + b_{y'}}{\|\mathbf{w}_{y'}\|_2} - \max_{j \in \{1, \dots, K\} \setminus y'} \frac{\mathbf{w}_j^T \mathbf{x} + b_j}{\|\mathbf{w}_j\|_2}$$

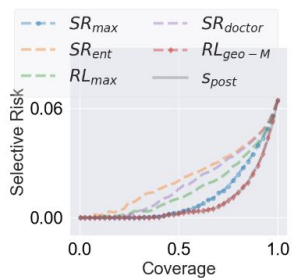
Confidence margin:

$$(\mathbf{w}_{y'}^T \mathbf{x} + b_{y'}) - \max_{i \in \{1, \dots, K\} \setminus y'} (\mathbf{w}_i^T \mathbf{x} + b_i)$$

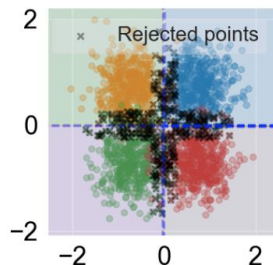
Apply them to the RLs $\mathcal{Z}$

Benefit: We don't need to worry about the scale of $\mathcal{Z}$

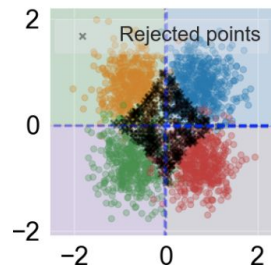
Additional benefit: robustness



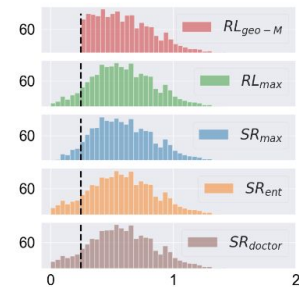
(a-1) RC curves



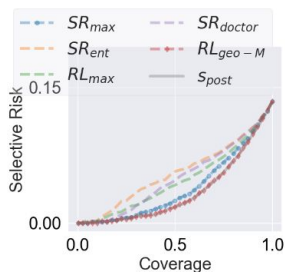
(b-1) RL_{geo-M}



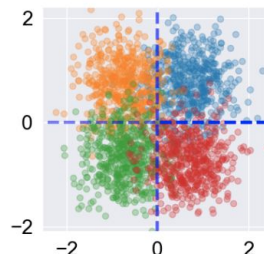
(c-1) SR_{max}



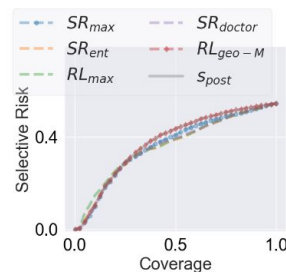
(d-1) Robustness radius



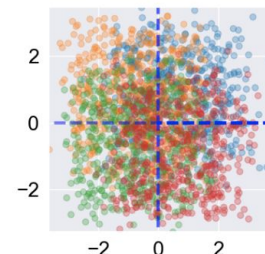
(a-2) RC curves



(b-2) Sample visualization



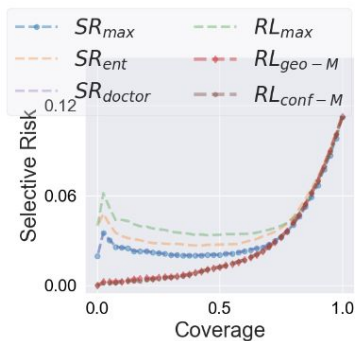
(a-3) RC curves



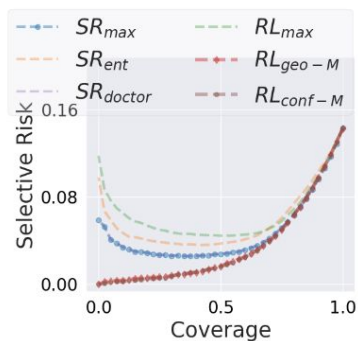
(b-3) Sample visualization

On real data

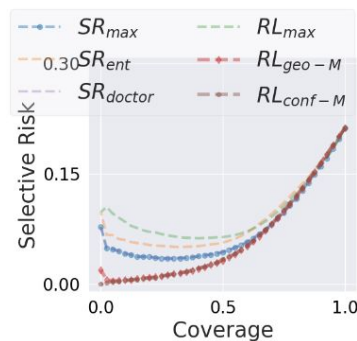
ImageNet vs ImageNet-C



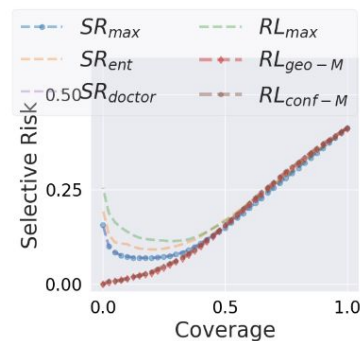
(a) IN (Clean)



(b) Gaussian blur Lv.1



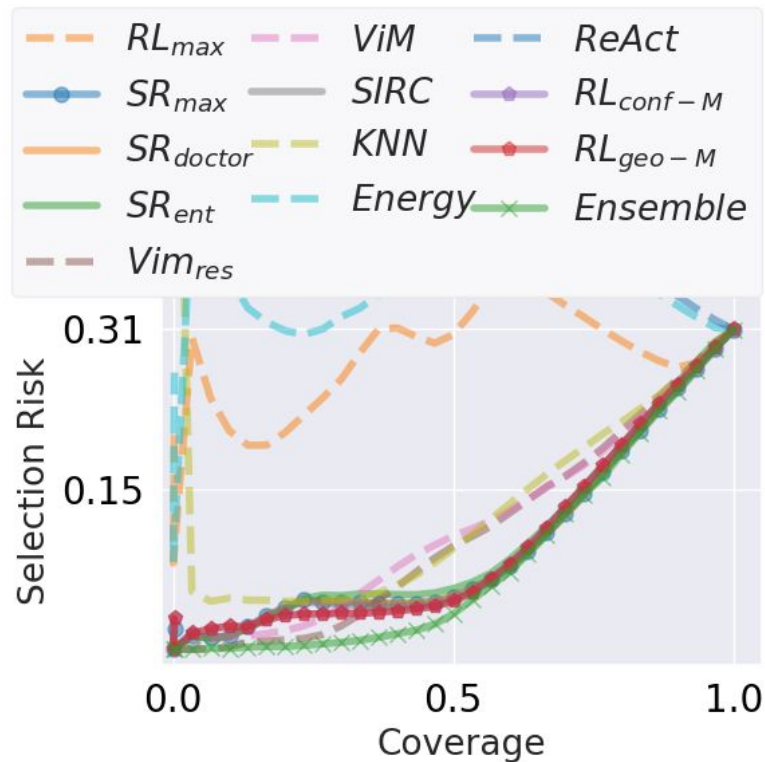
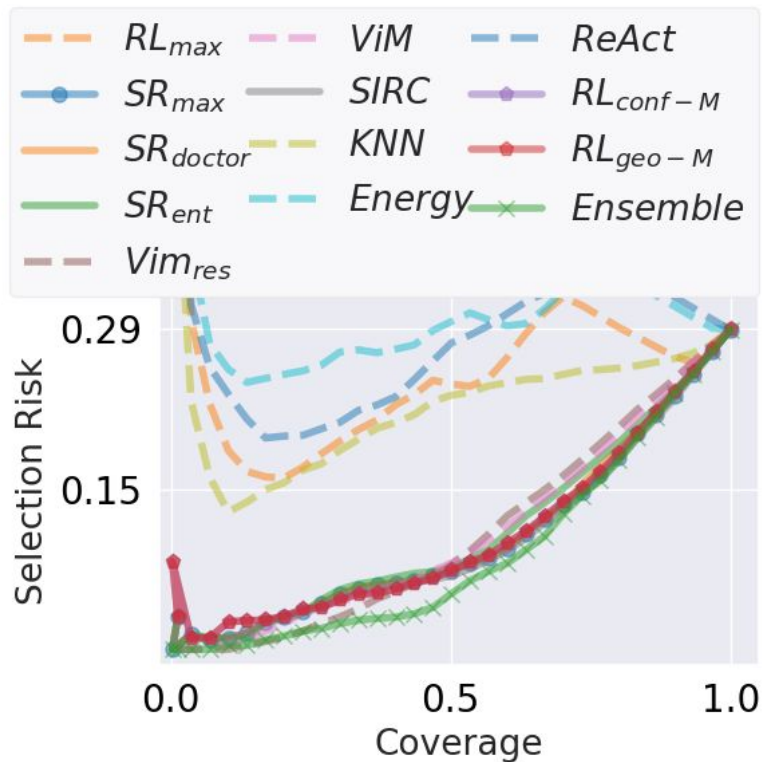
(c) Gaussian blur Lv.3



(d) Gaussian blur Lv.5

	IN (Clean)			Gaussian Blur			Brightness			Fog			Snow		
α	0.1	0.5	1	0.1	0.5	1	0.1	0.5	1	0.1	0.5	1	0.1	0.5	1
$RL_{\text{conf-M}}$	0.16	0.53	2.39	0.37	1.31	<u>6.05</u>	0.21	0.72	3.35	0.14	0.79	4.21	0.17	0.95	4.80
$RL_{\text{geo-M}}$	<u>0.27</u>	<u>0.59</u>	<u>2.43</u>	<u>0.57</u>	<u>1.36</u>	6.04	<u>0.33</u>	<u>0.79</u>	<u>3.39</u>	<u>0.25</u>	<u>0.86</u>	<u>4.22</u>	<u>0.34</u>	<u>1.02</u>	<u>4.81</u>
RL_{max}	5.54	4.05	4.57	9.74	7.38	9.52	7.38	5.17	6.06	7.74	5.77	7.01	9.44	6.44	7.90
SR_{max}	3.19	2.40	3.38	5.02	4.02	7.39	4.07	2.90	4.53	3.92	3.07	5.37	5.35	3.67	6.13
SR_{ent}	4.28	3.13	4.04	6.80	5.63	8.71	5.51	4.01	5.48	5.56	4.37	6.42	7.29	5.07	7.27
SR_{doctor}	3.21	2.38	3.40	5.05	4.05	7.47	4.10	2.93	4.58	3.95	3.10	5.42	5.39	3.71	6.20

Boosting the confidence?



More details

[Submitted on 8 May 2024 (v1), last revised 27 Nov 2024 (this version, v2)]

Selective Classification Under Distribution Shifts

Hengyue Liang, Le Peng, Ju Sun

In selective classification (SC), a classifier abstains from making predictions that are likely to be wrong to avoid excessive errors. To deploy imperfect classifiers -- either due to intrinsic statistical noise of data or for robustness issue of the classifier or beyond -- in high-stakes scenarios, SC appears to be an attractive and necessary path to follow. Despite decades of research in SC, most previous SC methods still focus on the ideal statistical setting only, i.e., the data distribution at deployment is the same as that of training, although practical data can come from the wild. To bridge this gap, in this paper, we propose an SC framework that takes into account distribution shifts, termed generalized selective classification, that covers label-shifted (or out-of-distribution) and covariate-shifted samples, in addition to typical in-distribution samples, the first of its kind in the SC literature. We focus on non-training-based confidence-score functions for generalized SC on deep learning (DL) classifiers, and propose two novel margin-based score functions. Through extensive analysis and experiments, we show that our proposed score functions are more effective and reliable than the existing ones for generalized SC on a variety of classification tasks and DL classifiers. Code is available at [this https URL](#).

Today's talk: **toward trustworthy and efficient AI**

- **Robustness evaluation and selective classification**
- **Inverse problems with pretrained diffusion models**
- **AI for healthcare: handling data imbalance**

Inverse problems

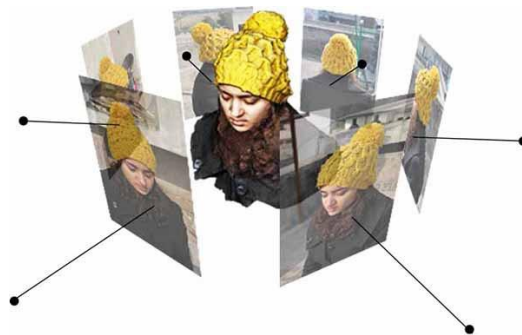
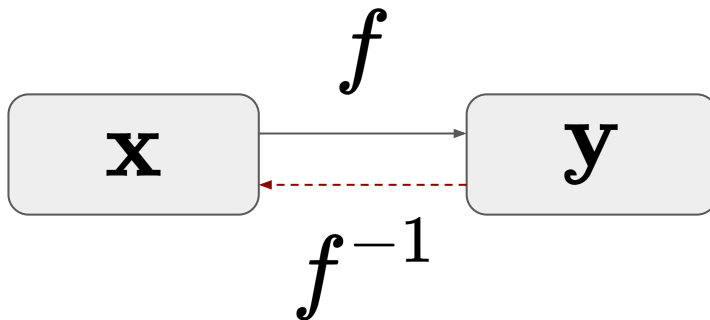
Inverse problem: given $\mathbf{y} = f(\mathbf{x})$, recover $\mathbf{x}$



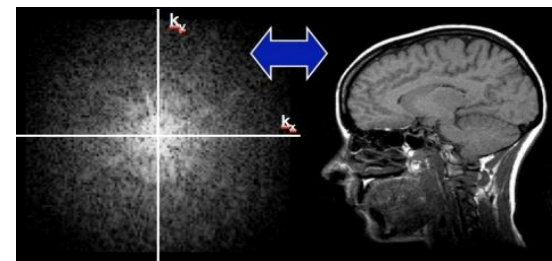
Image denoising



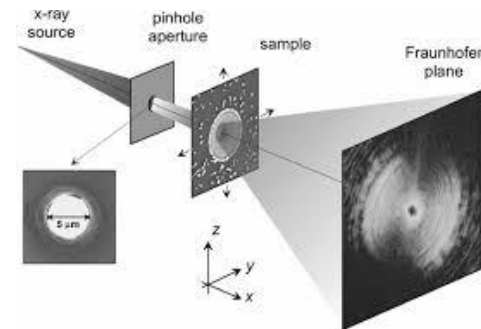
Image super-resolution



3D reconstruction



MRI reconstruction



Coherent diffraction imaging (CDI)

Traditional methods

Inverse problem: given $\mathbf{y} = f(\mathbf{x})$, recover $\mathbf{x}$

$$\min_{\mathbf{x}} \underbrace{\ell(\mathbf{y}, f(\mathbf{x}))}_{\text{data fitting}} + \lambda \underbrace{R(\mathbf{x})}_{\text{regularizer}} \quad \text{RegFit}$$

Questions

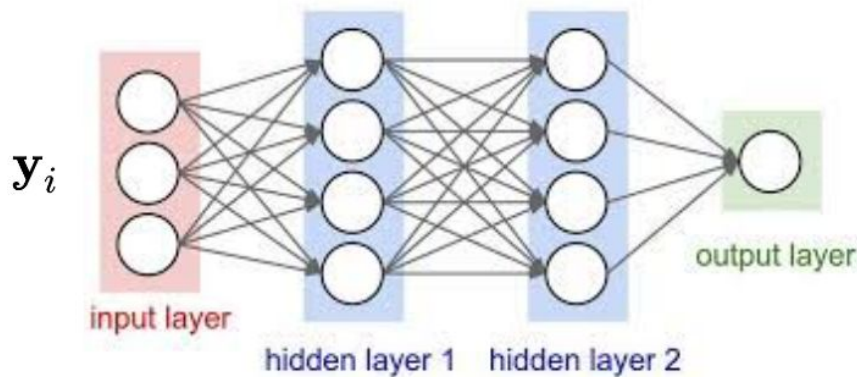
- Which ℓ ? (e.g., unknown/compound noise)
- Which R ? (e.g., structures not amenable to math description)
- Speed

Deep learning has changed everything

With paired datasets $\{(\mathbf{y}_i, \mathbf{x}_i)\}_{i=1,\dots,N}$

Direct inversion

Learn f^{-1} from $\{(\mathbf{y}_i, \mathbf{x}_i)\}_{i=1,\dots,N}$

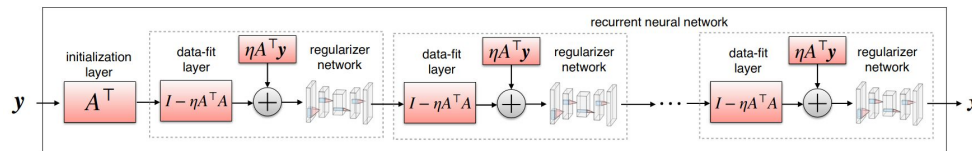


Algorithm unrolling

$$\min_{\mathbf{x}} \ell(\mathbf{y}, f(\mathbf{x})) + \lambda R(\mathbf{x})$$

$$\mathbf{x}^{k+1} = \mathcal{P}_R(\mathbf{x}^k - \eta \nabla^\top f(\mathbf{x}^k) \ell'(\mathbf{y}, f(\mathbf{x}^k)))$$

Idea: make $\mathcal{P}_R$ trainable



With paired datasets $\{(\mathbf{y}_i, \mathbf{x}_i)\}_{i=1,\dots,N}$

Conditional generation & regularization

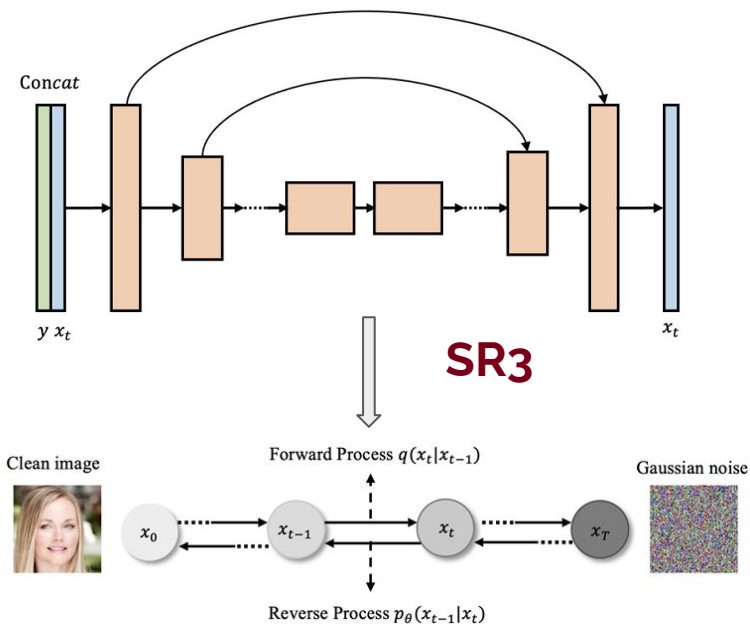
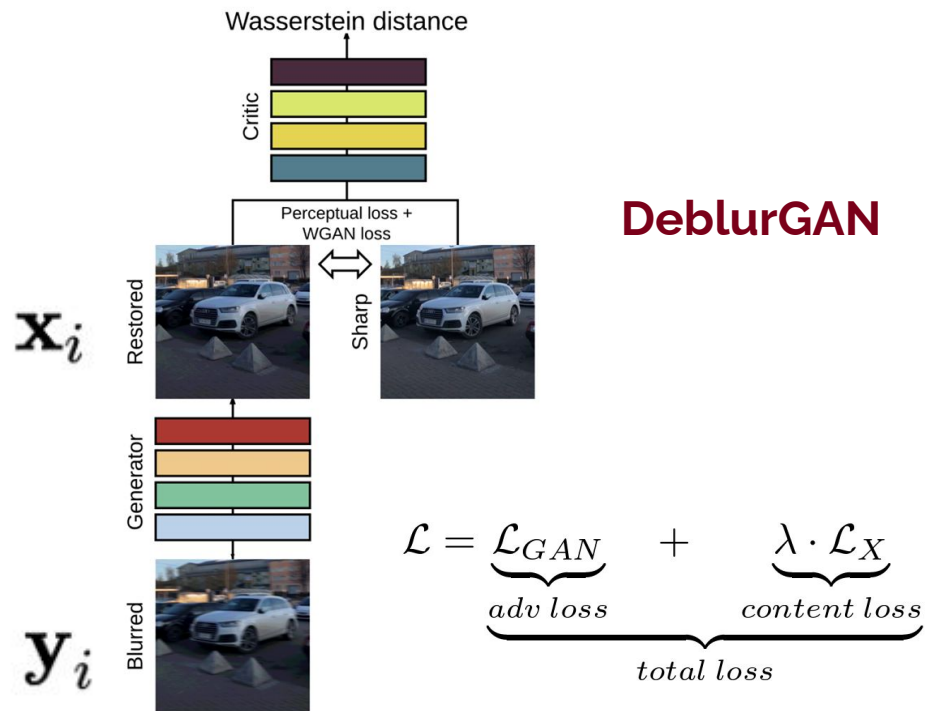


Image credit: <https://arxiv.org/abs/2308.09388>

With object datasets only $\{\mathbf{x}_i\}_{i=1,\dots,N}$

Model the distribution of the objects first, and then plug the prior in

GAN Inversion

Pretraining: $\mathbf{x}_i \approx G_\theta(\mathbf{z}_i) \quad \forall i$

Deployment: $\min_{\mathbf{z}} \ell(\mathbf{y}, f \circ G_\theta(\mathbf{z})) + \lambda R \circ G_\theta(\mathbf{z})$

Interleaving pretrained diffusion models

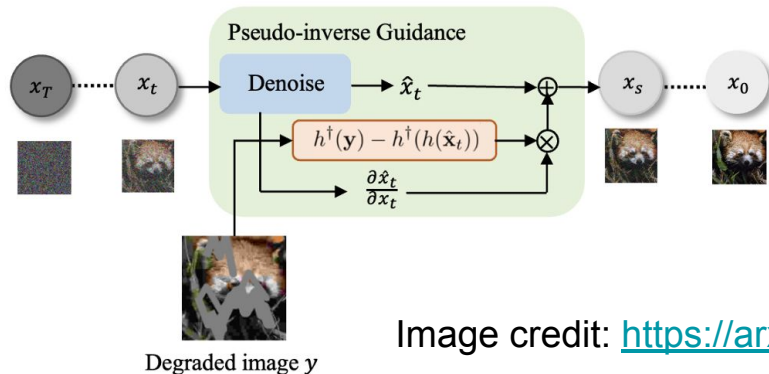


Image credit: <https://arxiv.org/abs/2308.09388>

Without datasets? Single-instance methods

Deep image prior (DIP) $\mathbf{x} \approx G_{\theta}(\mathbf{z})$ G_{θ} (and $\mathbf{z}$) trainable

$$\min_{\mathbf{x}} \underbrace{\ell(\mathbf{y}, f(\mathbf{x}))}_{\text{data fitting}} + \lambda \underbrace{R(\mathbf{x})}_{\text{regularizer}}$$

No extra training data!

$$\min_{\mathbf{z}} \ell(\mathbf{y}, f \circ G_{\theta}(\mathbf{z})) + \lambda R \circ G_{\theta}(\mathbf{z})$$

Neural implicit representation (NIR)

$$\mathbf{x} \approx \mathcal{D} \circ \bar{\mathbf{x}} \quad \mathcal{D} : \text{discretization} \quad \bar{\mathbf{x}} : \text{continuous function}$$

Physics-informed neural networks (PINN)

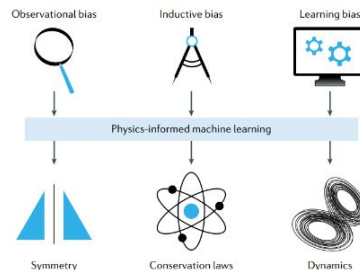
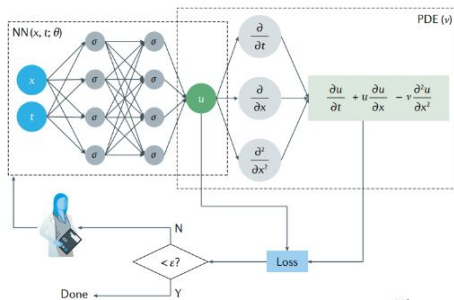


Table 2: Major categories of methods learning to solve inverse problems based on what is known about the forward model $\mathcal{A}$ and the nature of the training data, with examples for each. Details are described throughout Section 4.

	Supervised with unmatched (x, y) pairs	Train from unpaired x 's and y 's (Unpaired ground truths and Measurements)	Train from x 's only (Ground truth only)	Train from y 's only (Measurements only)
$\mathcal{A}$ fully known during training and testing (§4.1)	§4.1.1: Denoising auto-encoders [16], U-Net [78], Deep convolutional framelets [79] Unrolled optimization [80–83], Neumann networks [84]	<i>amounts to training from (x, y) pairs</i>	<i>amounts to training from (x, y) pairs</i>	§4.1.2: SURE LDAMP [85, 86], Deep Basis Pursuit [87]
$\mathcal{A}$ known only at test time (§4.2)	§4.2.2	§4.2.2	§4.2.1: CSGM [25], LDAMP [88], OneNet [22], Plug-and-play [89], RED [90]	§4.2.2
$\mathcal{A}$ partially known (§4.3)	§4.3.1	§4.3.2: CycleGAN [91]	§4.3.3: Blind deconvolution with GAN's [92–94]	§4.3.4: AmbientGAN [76], Noise2Noise [95], UAIR [96]
$\mathcal{A}$ unknown (§4.4)	§4.4.1: AUTOMAP [97]	§4.4.2	§4.4.2	§4.4.2

Deep Learning Techniques for Inverse Problems in Imaging

Gregory Ongie^{*}, Ajil Jalal[†], Christopher A. Metzler[‡]
 Richard G. Baraniuk[§], Alexandros G. Dimakis[¶], Rebecca Willett^{||}

<https://arxiv.org/abs/2005.06001>

But focused on linear IPs

Other specialized surveys

Algorithm Unrolling: Interpretable, Efficient Deep Learning for Signal and Image Processing

Vishal Monga, *Senior Member, IEEE*, Yuelong Li, *Member, IEEE*, and Yonina C. Eldar, *Fellow, IEEE*

Focused on alg. unrolling

Untrained Neural Network Priors for Inverse Imaging Problems: A Survey

Deep Internal Learning:

Deep Learning from a Single Input

Understanding Untrained Deep Models for Inverse Problems: Algorithms and Theory

Tom Tirer *Member,*

**Focused on
single-instance methods**

Ismail Alkhouri, Evan Bell, Avrajit Ghosh, Shijun Liang, Rongrong Wang,

Theoretical Perspectives on Deep Learning Methods in Inverse Problems

Jonathan Scarlett, Reinhard Heckel, Miguel R. D. Rodrigues, Paul Hand, and Yonina C. Eldar

**Focused on theories for
linear IPs**

Focus here:

Solving Inverse Problems (IPs) Using Pretrained Flow-Based Models

[Submitted on 30 Sep 2024]

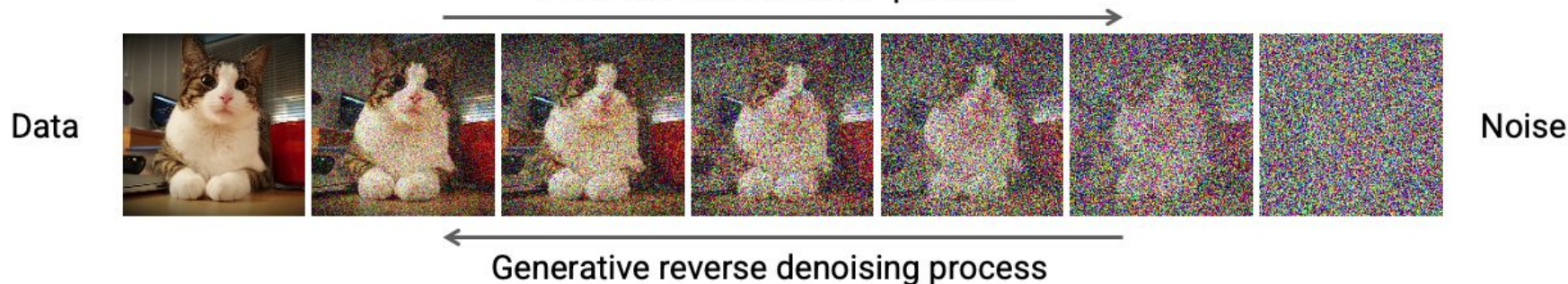
A Survey on Diffusion Models for Inverse Problems

Giannis Daras, Hyungjin Chung, Chieh-Hsin Lai, Yuki Mitsufuji, Jong Chul Ye, Peyman Milanfar,
Alexandros G. Dimakis, Mauricio Delbracio

Diffusion models

$$d\mathbf{x} = -\beta_t/2 \cdot \mathbf{x}dt + \sqrt{\beta_t}d\mathbf{w},$$

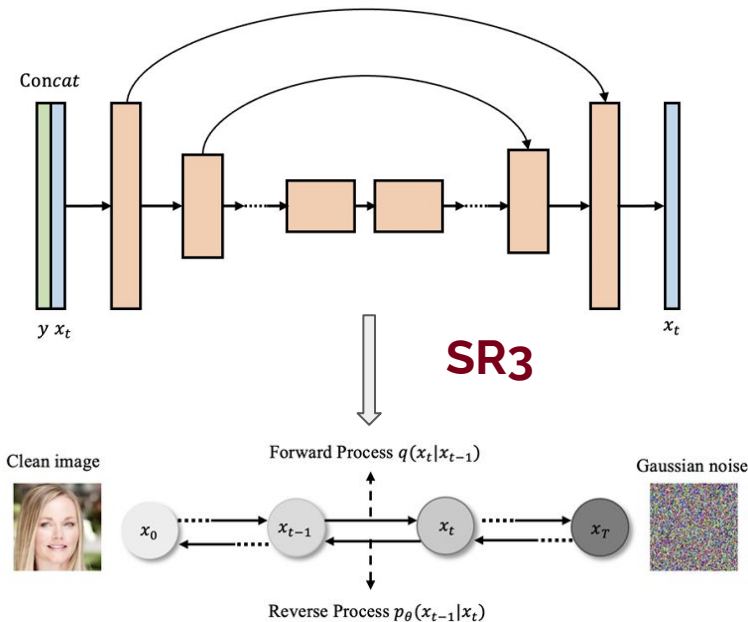
Fixed forward diffusion process



$$d\mathbf{x} = -\beta_t \left[\mathbf{x}/2 + \nabla_{\mathbf{x}} \log p_t(\mathbf{x}) \right] dt + \sqrt{\beta_t}d\overline{\mathbf{w}}.$$

Diffusion models for inverse problems (IPs)

Supervised



Zero-shot

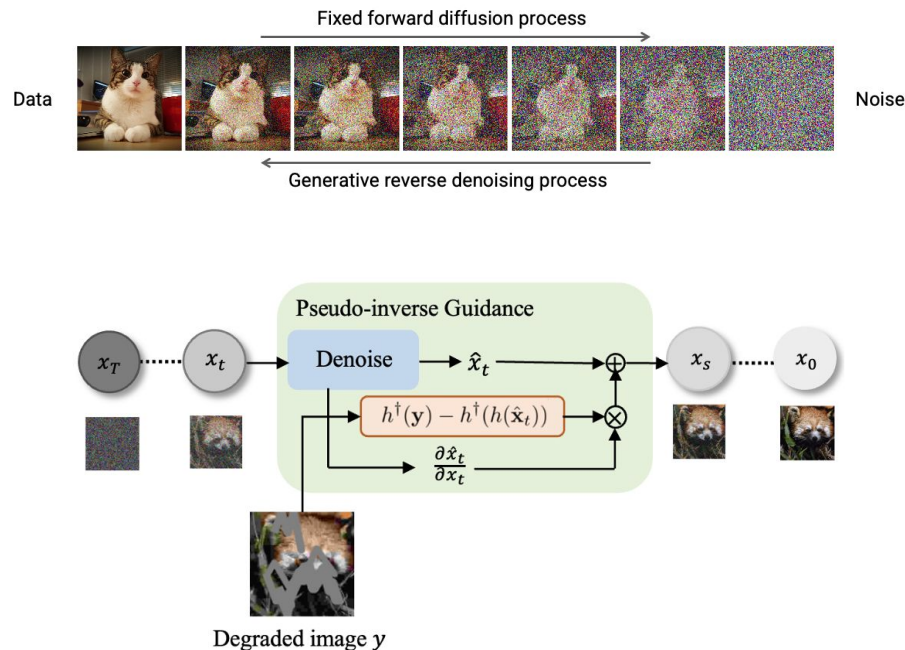


Image credit: <https://arxiv.org/abs/2308.09388>

Focus: IPs with pretrained diffusion models

(Reverse SDE for DDPM)
$$d\mathbf{x} = -\beta_t [\mathbf{x}/2 + \nabla_{\mathbf{x}} \log p_t(\mathbf{x})] dt + \sqrt{\beta_t} d\bar{\mathbf{w}}$$



Think of **conditional score function**

$$\nabla_{\mathbf{x}} \log p_t(\mathbf{x}|\mathbf{y}) = \nabla_{\mathbf{x}} \log p_t(\mathbf{x}) + \nabla_{\mathbf{x}} \log p_t(\mathbf{y}|\mathbf{x})$$



Conditional reverse SDE

$$d\mathbf{x} = [-\beta_t/2 \cdot \mathbf{x} - \beta_t(\nabla_{\mathbf{x}} \log p_t(\mathbf{x}) + \nabla_{\mathbf{x}} \log p_t(\mathbf{y}|\mathbf{x}))] dt + \sqrt{\beta_t} d\bar{\mathbf{w}}$$

Coping with conditional score function

$$\nabla_{\mathbf{x}} \log p_t(\mathbf{x}|\mathbf{y}) = \boxed{\nabla_{\mathbf{x}} \log p_t(\mathbf{x})} + \boxed{\nabla_{\mathbf{x}} \log p_t(\mathbf{y}|\mathbf{x})}$$

$$p_t(\mathbf{y}|\mathbf{x}(t))$$

$$\approx \underline{p}_t(\mathbf{y}|\hat{\mathbf{x}}(0)[\mathbf{x}(t)])$$

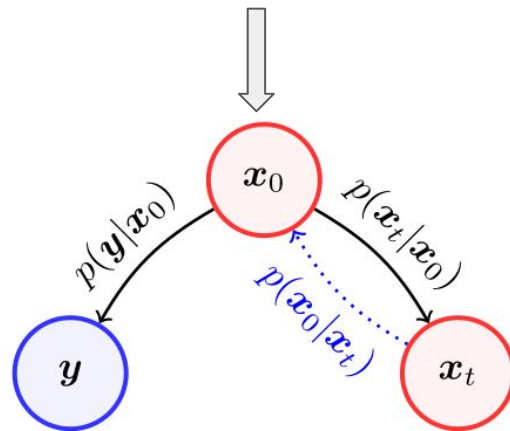


Figure 2: Probabilistic graph. Black solid line: tractable, blue dotted line: intractable in general.

Interleaving methods

Algorithm 1 Template for interleaving methods

Input: # Diffusion steps T , measurement $\mathbf{y}$

1: $\mathbf{x}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

2: **for** $i = T - 1$ to 0 **do**

3: $\hat{\mathbf{s}} \leftarrow \boldsymbol{\varepsilon}_{\theta}^{(i)}(\mathbf{x}_i)$

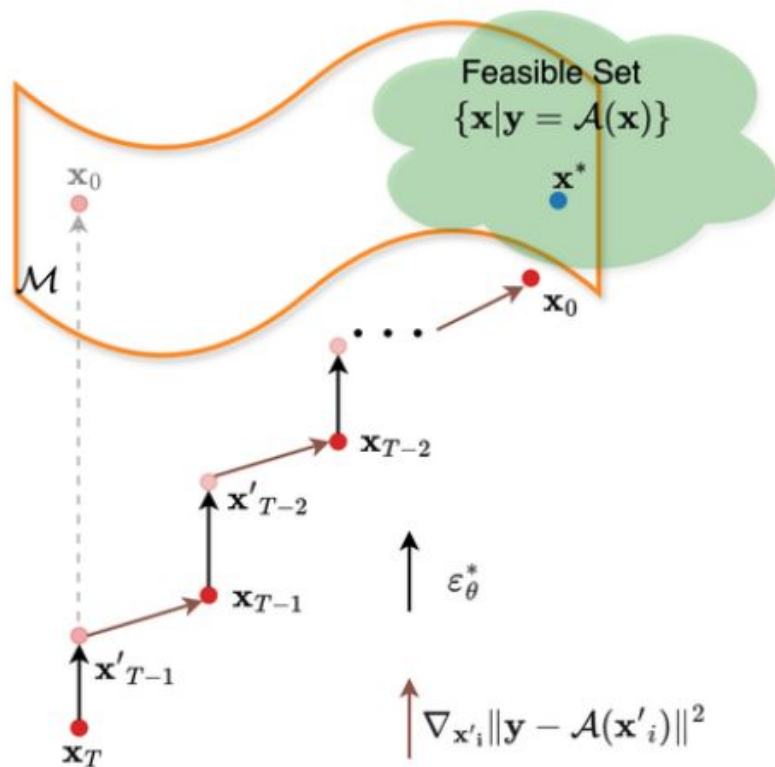
4: $\hat{\mathbf{x}}_0 \leftarrow \frac{1}{\sqrt{\bar{\alpha}_i}}(\mathbf{x}_i - \sqrt{1 - \bar{\alpha}_i}\hat{\mathbf{s}})$

5: $\mathbf{x}'_{i-1} \leftarrow$ DDIM reverse with $\hat{\mathbf{x}}_0$ and $\hat{\mathbf{s}}$

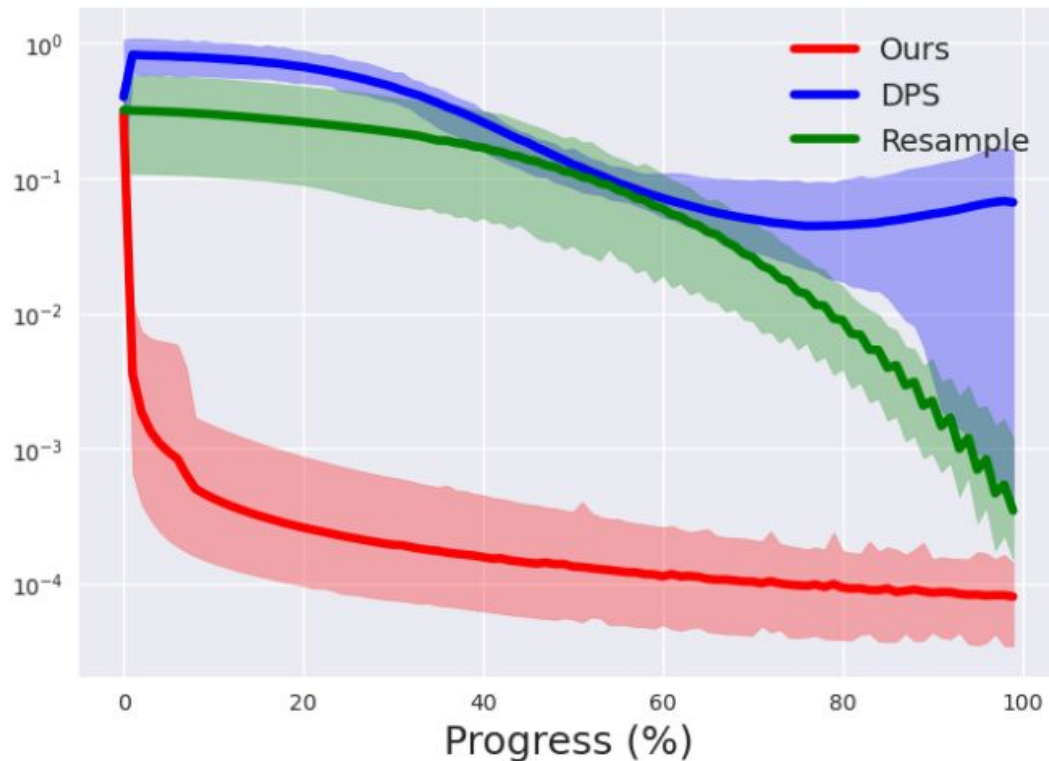
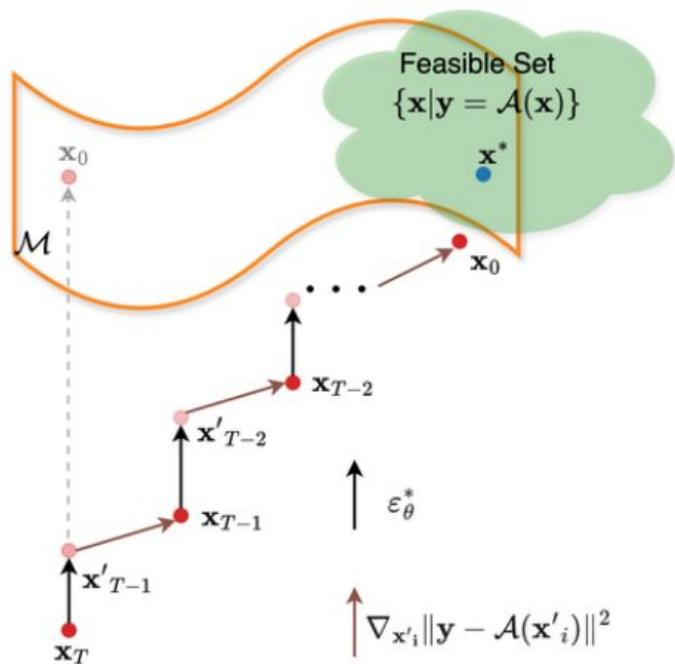
6: $\mathbf{x}_{i-1} \leftarrow$ (Approximately) Projection [39, 30, 33, 32, 40, 41, 34] or gradient update [20, 28, 19, 21, 29, 27, 26] with $\hat{\mathbf{x}}_0$ and $\mathbf{x}'_{i-1}$ to get closer to $\{\mathbf{x} | \mathbf{y} = \mathcal{A}(\mathbf{x})\}$

7: **end for**

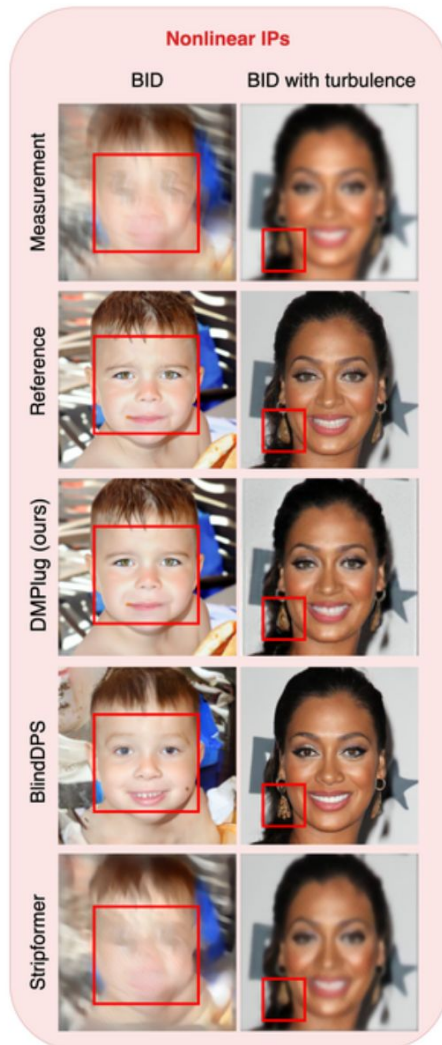
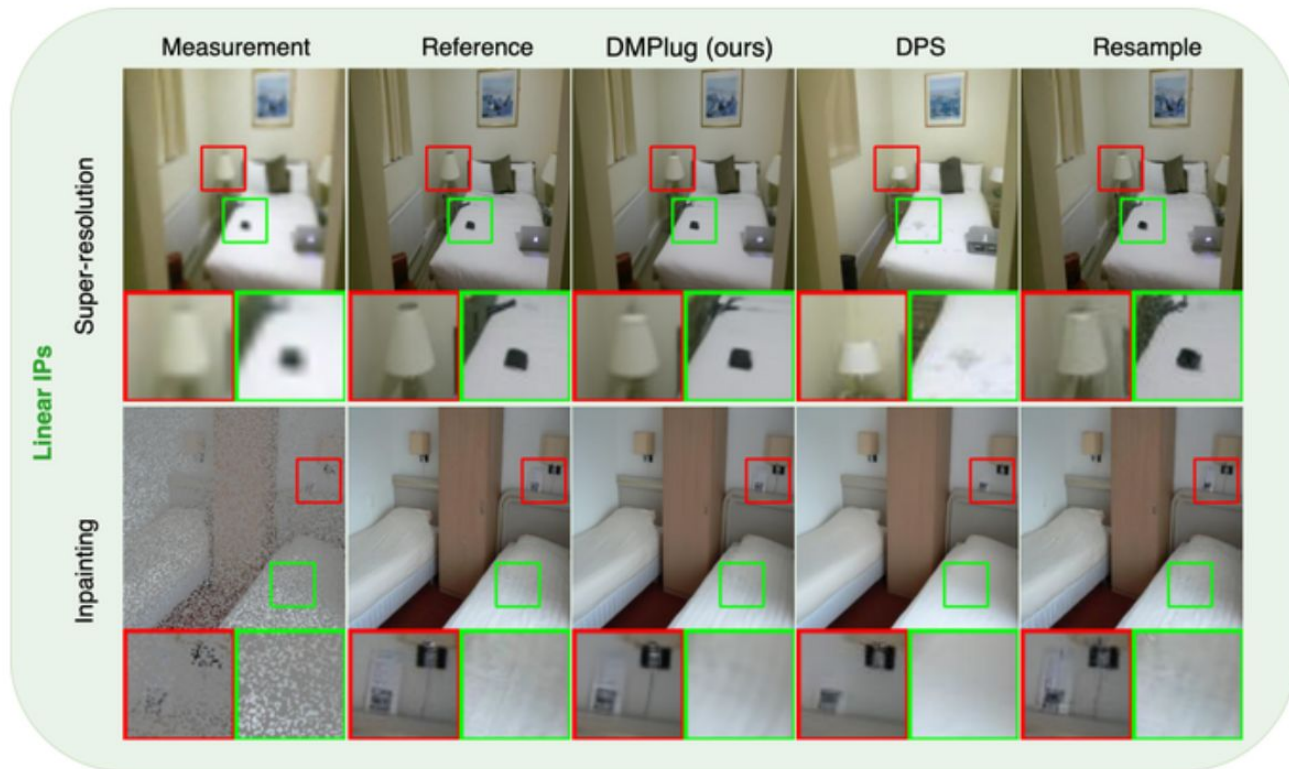
Output: Recovered object $\mathbf{x}_0$



Issue I: Measurement feasibility



Issue 2: Manifold feasibility



Issue 3: Robustness to unknown noise

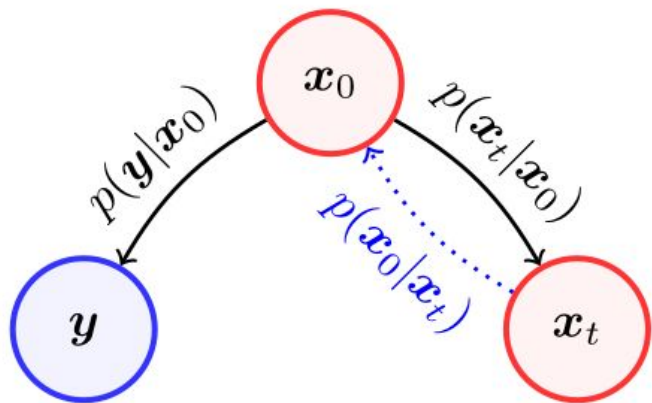


Figure 2: Probabilistic graph. Black solid line: tractable, blue dotted line: intractable in general.

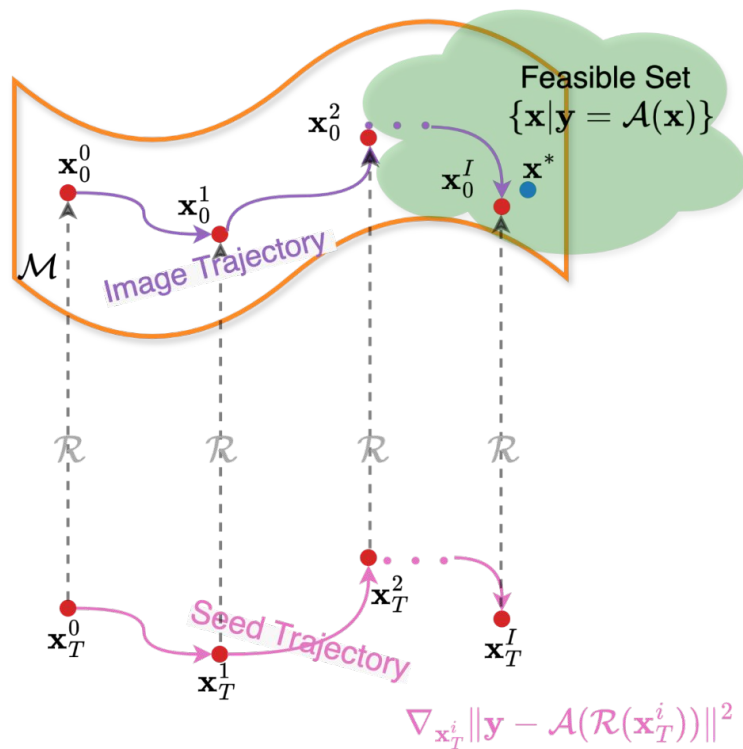
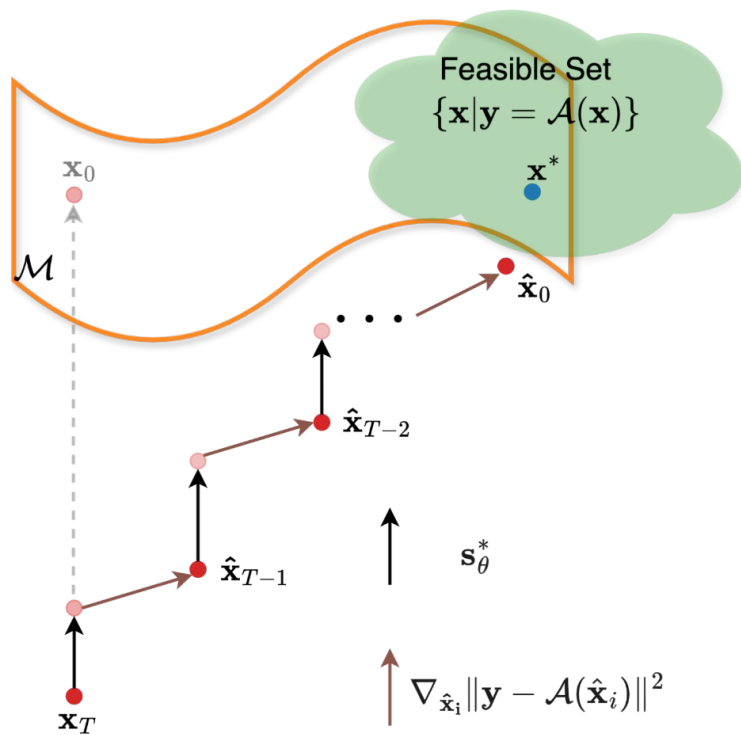
Algorithm 1 DPS - Gaussian

Require: $N, \mathbf{y}, \{\zeta_i\}_{i=1}^N, \{\tilde{\sigma}_i\}_{i=1}^N$

- 1: $\mathbf{x}_N \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - 2: **for** $i = N - 1$ **to** 0 **do**
 - 3: $\hat{\mathbf{s}} \leftarrow \mathbf{s}_\theta(\mathbf{x}_i, i)$
 - 4: $\hat{\mathbf{x}}_0 \leftarrow \frac{1}{\sqrt{\bar{\alpha}_i}}(\mathbf{x}_i + (1 - \bar{\alpha}_i)\hat{\mathbf{s}})$
 - 5: $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - 6: $\mathbf{x}'_{i-1} \leftarrow \frac{\sqrt{\bar{\alpha}_i}(1 - \bar{\alpha}_{i-1})}{1 - \bar{\alpha}_i}\mathbf{x}_i + \frac{\sqrt{\bar{\alpha}_{i-1}}\beta_i}{1 - \bar{\alpha}_i}\hat{\mathbf{x}}_0 + \tilde{\sigma}_i\mathbf{z}$
 - 7: $\mathbf{x}_{i-1} \leftarrow \mathbf{x}'_{i-1} - \zeta_i \nabla_{\mathbf{x}_i} \|\mathbf{y} - \mathcal{A}(\hat{\mathbf{x}}_0)\|_2^2$
 - 8: **end for**
 - 9: **return** $\hat{\mathbf{x}}_0$
-

depending on noise level

Our solution: DMPlug



Our solution: DMPlug

Viewing the reverse process as a function $\mathcal{R}$

$$\mathcal{R} = g_{\epsilon_{\theta}^{(0)}} \circ g_{\epsilon_{\theta}^{(1)}} \circ \cdots \circ g_{\epsilon_{\theta}^{(T-2)}} \circ g_{\epsilon_{\theta}^{(T-1)}}. \quad (\circ \text{ means function composition})$$

$$(\mathbf{DMPlug}) \quad z^* \in \arg \min_z \ell(y, \mathcal{A}(\mathcal{R}(z))) + \Omega(\mathcal{R}(z)), \quad x^* = \mathcal{R}(z^*).$$

Measurement
feasibility

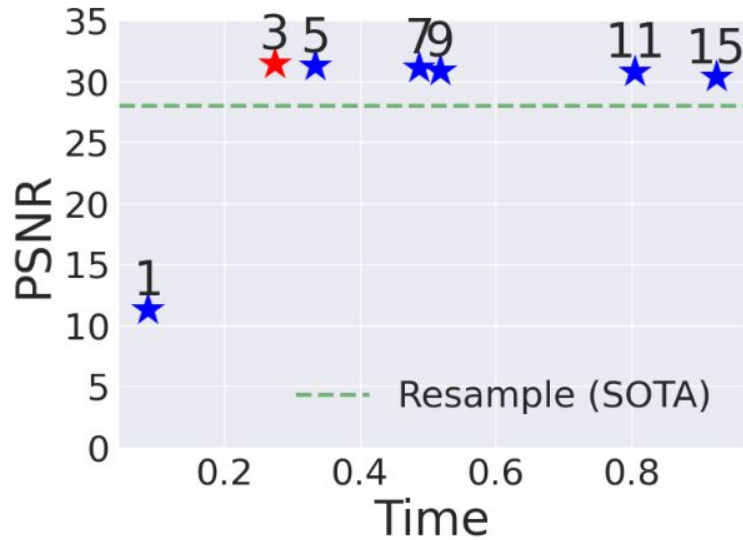
Manifold
feasibility

Overcoming the computational bottleneck

$$\mathcal{R} = g_{\epsilon_{\theta}^{(0)}} \circ g_{\epsilon_{\theta}^{(1)}} \circ \cdots \circ g_{\epsilon_{\theta}^{(T-2)}} \circ g_{\epsilon_{\theta}^{(T-1)}}. \quad (\circ \text{ means function composition})$$

$$(\mathbf{DMPlug}) \quad z^* \in \arg \min_z \ell(\mathbf{y}, \mathcal{A}(\mathcal{R}(z))) + \Omega(\mathcal{R}(z)), \quad x^* = \mathcal{R}(z^*).$$

Issue: T blocks of DNNs involved, and we have to back-propagate through it



On linear IPs

Table 1: (**Linear IPs**) **Super-resolution** and **inpainting** with additive Gaussian noise ($\sigma = 0.01$). (**Bold**: best, under: second best, **green**: performance increase, **red**: performance decrease)

	Super-resolution ($4\times$)						Inpainting (Random 70%)					
	CelebA [65] (256×256)			FFHQ [66] (256×256)			CelebA [65] (256×256)			FFHQ [66] (256×256)		
	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑
ADMM-PnP [68]	0.217	26.99	0.808	0.229	26.25	0.794	0.091	31.94	0.923	0.104	30.64	0.901
DMPS [29]	<u>0.070</u>	<u>28.89</u>	<u>0.848</u>	0.076	<u>28.03</u>	<u>0.843</u>	0.297	24.52	0.693	0.326	23.31	0.664
DDRM [32]	0.226	26.34	0.754	0.282	25.11	0.731	0.185	26.10	0.712	0.201	25.44	0.722
MCG [30]	0.725	19.88	0.323	0.786	18.20	0.271	1.283	10.16	0.049	1.276	10.37	0.050
ILVR [41]	0.322	21.63	0.603	0.360	20.73	0.570	0.447	15.82	0.484	0.483	15.10	0.450
DPS [19]	0.087	28.32	0.823	0.098	27.44	0.814	0.043	<u>32.24</u>	<u>0.924</u>	0.046	<u>30.95</u>	<u>0.913</u>
ReSample [20]	0.080	28.29	0.819	0.108	25.22	0.773	0.039	30.12	0.904	<u>0.044</u>	27.91	0.884
DMPlug (ours)	0.067	31.25	0.878	<u>0.079</u>	30.25	0.871	0.039	34.03	0.936	0.038	33.01	0.931
Ours vs. Best compe.	-0.003	+2.36	+0.030	+0.003	+2.22	+0.028	-0.000	+1.79	+0.012	-0.006	+2.06	+0.018

On nonlinear IPs

Table 2: (**Nonlinear IP**) **Nonlinear deblurring** with additive Gaussian noise ($\sigma = 0.01$). (**Bold**: best, under: second best, **green**: performance increase, **red**: performance decrease)

	CelebA 65 (256×256)			FFHQ 66 (256×256)			LSUN 67 (256×256)		
	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑
BKS-styleGAN 69	1.047	22.82	0.653	1.051	22.07	0.620	0.987	20.90	0.538
BKS-generic 69	1.051	21.04	0.591	1.056	20.76	0.583	0.994	18.55	0.481
MCG 30	0.705	13.18	0.135	0.675	13.71	0.167	0.698	14.28	0.188
ILVR 41	0.335	21.08	0.586	0.374	20.40	0.556	0.482	18.76	0.444
DPS 19	0.149	24.57	0.723	0.130	25.00	0.759	0.244	23.46	0.684
ReSample 20	<u>0.104</u>	<u>28.52</u>	<u>0.839</u>	<u>0.104</u>	<u>27.02</u>	<u>0.834</u>	<u>0.143</u>	<u>26.03</u>	<u>0.803</u>
DMPlug (ours)	0.073	31.61	0.882	0.057	32.83	0.907	0.083	30.74	0.882
Ours vs. Best compe.	-0.031	+3.09	+0.043	-0.047	+5.79	+0.073	-0.060	+4.71	+0.079

More on nonlinear IPs

Table 4: (**Nonlinear IP**) **BID** with additive Gaussian noise ($\sigma = 0.01$). (**Bold**: best, under: second best, **green**: performance increase, **red**: performance decrease)

	CelebA [65] (256 × 256)						FFHQ [66] (256 × 256)					
	Motion blur			Gaussian blur			Motion blur			Gaussian blur		
	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑	LPIPS↓	PSNR↑	SSIM↑
SelfDeblur [75]	0.568	16.59	0.417	0.579	16.55	0.423	0.628	16.33	0.408	0.604	16.22	0.410
DeBlurGANv2 [5]	0.313	20.56	0.613	0.350	24.29	0.743	0.353	19.67	0.581	0.374	23.58	0.726
Stripformer [6]	0.287	22.06	0.644	0.316	25.03	<u>0.747</u>	0.324	21.31	0.613	0.339	<u>24.34</u>	<u>0.728</u>
MPRNet [7]	0.332	20.53	0.620	0.375	22.72	0.698	0.373	19.70	0.590	0.394	22.33	0.685
Pan-DCP [73]	0.606	15.83	0.483	0.653	20.57	0.701	0.616	15.59	0.464	0.667	20.69	0.698
Pan- ℓ_0 [74]	0.631	15.16	0.470	0.654	20.49	0.675	0.642	14.43	0.443	0.669	20.34	0.671
ILVR [41]	0.398	19.23	0.520	0.338	21.20	0.588	0.445	18.33	0.484	0.375	20.45	0.555
BlindDPS [21]	<u>0.164</u>	<u>23.60</u>	<u>0.682</u>	<u>0.173</u>	<u>25.15</u>	0.721	<u>0.185</u>	<u>21.77</u>	<u>0.630</u>	<u>0.193</u>	23.83	0.693
DMPlug (ours)	0.104	29.61	0.825	0.140	28.84	0.795	0.135	27.99	0.794	0.169	28.26	0.811
Ours vs. Best compe.	-0.060	+6.01	+0.143	-0.033	+3.69	+0.048	-0.050	+6.22	+0.164	-0.024	+3.92	+0.083

More details

[Submitted on 27 May 2024 (v1), last revised 6 Nov 2024 (this version, v2)]

DMPlug: A Plug-in Method for Solving Inverse Problems with Diffusion Models

Hengkang Wang, Xu Zhang, Taihui Li, Yuxiang Wan, Tiancong Chen, Ju Sun

Pretrained diffusion models (DMs) have recently been popularly used in solving inverse problems (IPs). The existing methods mostly interleave iterative steps in the reverse diffusion process and iterative steps to bring the iterates closer to satisfying the measurement constraint. However, such interleaving methods struggle to produce final results that look like natural objects of interest (i.e., manifold feasibility) and fit the measurement (i.e., measurement feasibility), especially for nonlinear IPs. Moreover, their capabilities to deal with noisy IPs with unknown types and levels of measurement noise are unknown. In this paper, we advocate viewing the reverse process in DMs as a function and propose a novel plug-in method for solving IPs using pretrained DMs, dubbed DMPlug. DMPlug addresses the issues of manifold feasibility and measurement feasibility in a principled manner, and also shows great potential for being robust to unknown types and levels of noise. Through extensive experiments across various IP tasks, including two linear and three nonlinear IPs, we demonstrate that DMPlug consistently outperforms state-of-the-art methods, often by large margins especially for nonlinear IPs. The code is available at [this https URL](#).

DMPlug for video restoration

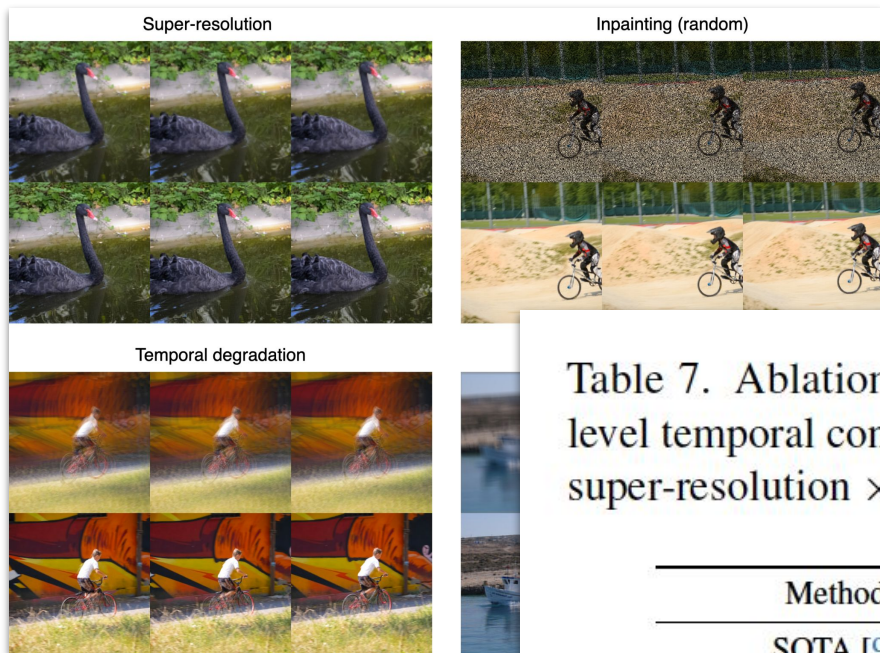


Table 7. Ablation study on two essential components for multi-level temporal consistency, performed on DAVIS dataset for video super-resolution $\times 4$. (**Bold**: best, under: second best)

Method	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	WE(10^{-2}) $\downarrow$
SOTA [9]	26.037	0.717	0.339	1.411
Base	24.701	0.612	0.366	1.398
Base + Semantic	26.098	0.703	0.410	1.057
Base + Pixel	<u>27.141</u>	<u>0.736</u>	0.301	<u>0.943</u>
Base + Semantic + Pixel	27.959	0.790	<u>0.321</u>	0.725

More details

[Submitted on 19 Mar 2025]

Temporal-Consistent Video Restoration with Pre-trained Diffusion Models

Hengkang Wang, Yang Liu, Huidong Liu, Chien-Chih Wang, Yanhui Guo, Hongdong Li, Bryan Wang, Ju Sun

Video restoration (VR) aims to recover high-quality videos from degraded ones. Although recent zero-shot VR methods using pre-trained diffusion models (DMs) show good promise, they suffer from approximation errors during reverse diffusion and insufficient temporal consistency. Moreover, dealing with 3D video data, VR is inherently computationally intensive. In this paper, we advocate viewing the reverse process in DMs as a function and present a novel Maximum a Posterior (MAP) framework that directly parameterizes video frames in the seed space of DMs, eliminating approximation errors. We also introduce strategies to promote bilevel temporal consistency: semantic consistency by leveraging clustering structures in the seed space, and pixel-level consistency by progressive warping with optical flow refinements. Extensive experiments on multiple virtual reality tasks demonstrate superior visual quality and temporal consistency achieved by our method compared to the state-of-the-art.

Subjects: **Computer Vision and Pattern Recognition (cs.CV)**

Cite as: [arXiv:2503.14863](https://arxiv.org/abs/2503.14863) [cs.CV]

Surprise?



	PSNR↑	SSIM↑	LPIPS↓	CLIPQA↑
DIP	<u>27.5854</u>	<u>0.7179</u>	0.3898	0.2396
D-Flow (DS)	28.1389	0.7628	0.2783	0.5871
D-Flow (FD)	25.0120	0.7084	0.5335	0.3607
D-Flow (FD-S)	25.1453	0.6829	0.5213	0.3228
FlowDPS (DS)	22.1191	0.5603	<u>0.3850</u>	<u>0.5417</u>
FlowDPS (FD)	22.1404	0.5930	0.5412	0.2906
FlowDPS (FD-S)	22.0538	0.5920	0.5408	0.2913

Table 1: Comparison between foundation FM, domain-specific FM, and untrained priors for Gaussian deblurring the on AFHQ-Cat dataset (resolution: 256×256). DS: domain-specific FM; FD: foundation FM; FD-S: strengthened foundation FM; DIP: deep image prior. **Bold**: best, & underline: second best, for each metric/column. The foundation model is Stable Diffusion V3 here.

Foundation/Universal priors <<
Domain-specific, and even untrained priors

More details

[Submitted on 1 Aug 2025]

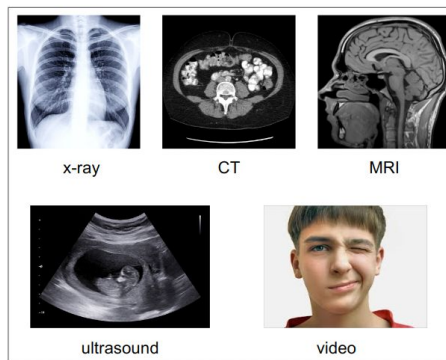
FMPlug: Plug-In Foundation Flow-Matching Priors for Inverse Problems

Yuxiang Wan, Ryan Devera, Wenjie Zhang, Ju Sun

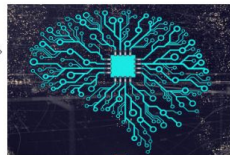
We present FMPlug, a novel plug-in framework that enhances foundation flow-matching (FM) priors for solving ill-posed inverse problems. Unlike traditional approaches that rely on domain-specific or untrained priors, FMPlug smartly leverages two simple but powerful insights: the similarity between observed and desired objects and the Gaussianity of generative flows. By introducing a time-adaptive warm-up strategy and sharp Gaussianity regularization, FMPlug unlocks the true potential of domain-agnostic foundation models. Our method beats state-of-the-art methods that use foundation FM priors by significant margins, on image super-resolution and Gaussian deblurring.

Today's talk: **toward trustworthy and efficient AI**

- **Robustness evaluation and selective classification**
- **Inverse problems with pretrained diffusion models**
- **AI for healthcare: handling data imbalance**

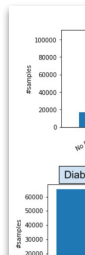


AI models



AI-based
diagnostic

**Supported by 3+
(4th forthcoming)
active NIH R01
grants**



Risk-coverage tradeoff

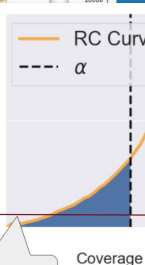
$$(f, g)(\mathbf{x}) \triangleq \begin{cases} f(\mathbf{x}) & \text{if } g(\mathbf{x}) = 1; \\ \text{abstain} & \text{if } g(\mathbf{x}) = 0. \end{cases}$$

$$g_\gamma(\mathbf{x}) = \mathbb{1}[s(\mathbf{x}) > \gamma]$$

$$(\text{coverage}) \phi_\gamma = \mathbb{E}_D[g_\gamma(\mathbf{x})],$$

$$(\text{selection risk}) R_\gamma = \mathbb{E}_D[\ell(f(\mathbf{x}), y)g_\gamma(\mathbf{x})]/\phi_\gamma.$$

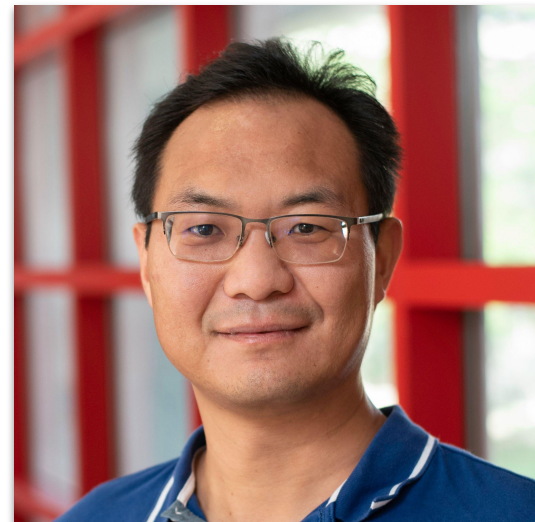
Selective Risk



High-stakes corner

**Federated learning
Imbalanced learning
Selective prediction**

data
imbalance,
and AI safety
in healthcare
(& beyond)



Professor Ju Sun

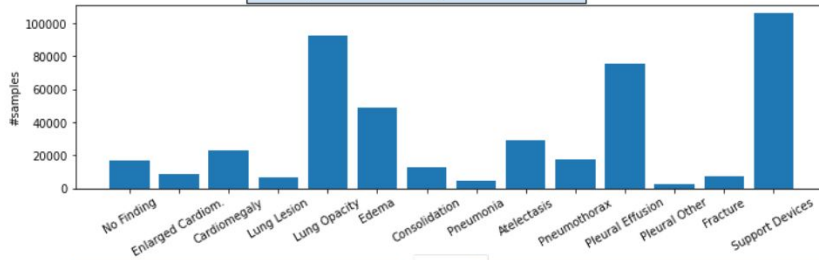
McKnight Land-Grant Professor
Computer Science & Engineering

<https://sunju.org/>

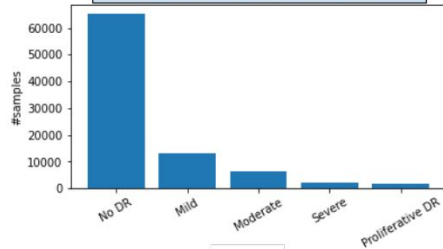
Leader, Group of Learning,
Optimization, Vision, healthcare, and
X (GLOVEX; <https://glovex.umn.edu/>)
Leader of Computer Vision, UMN
Program for Clinical AI

Addressing data inequality—imbalanced learning

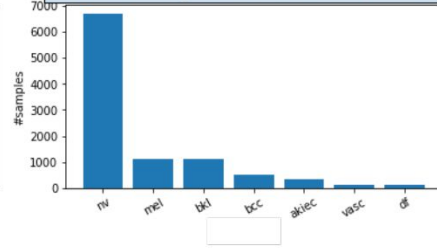
CheXpert Chest X-ray Classification



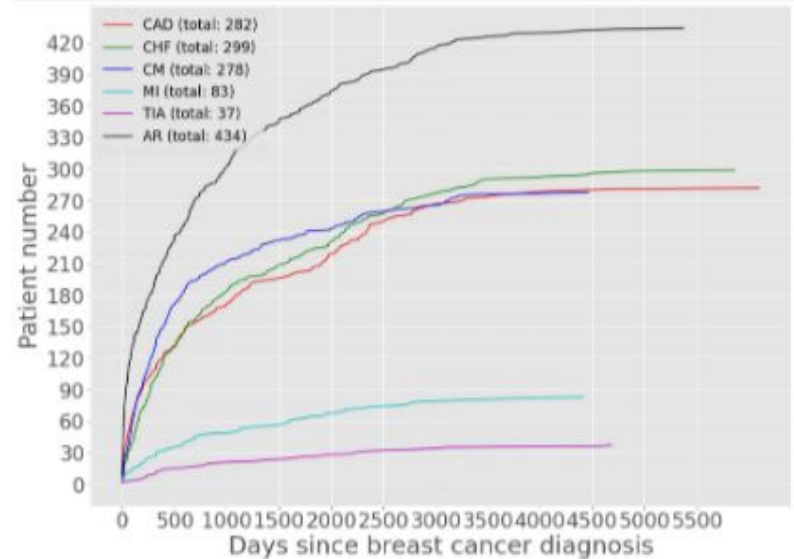
Diabetic Retinopathy Classification



HAM10000: Pigmented Skin Lesion Classification



Imbalanced classification (IC)



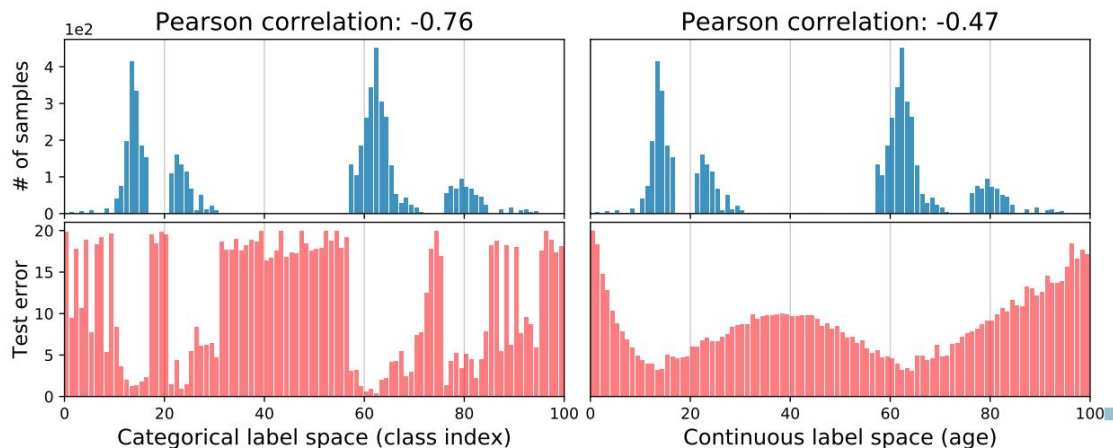
Imbalanced regression (IR)

Why imbalance learning is challenging?

	Predicted POS	Predicted NEG
POS	70	30
NEG	1000	9000

Accuracy: $9070/10100 = 0.898$
True Positive Rate (Sensitivity, Recall): 0.7
True Negative Rate (Specificity): 0.9
Balanced Accuracy: $(0.7 + 0.9)/2 = 0.80$
Precision (POS): $70/1070 = 0.065$
F1 Score: $2 * 0.065 * 0.7 / (0.065 + 0.7) = 0.119$

Figure 2: An example confusion table for binary classification, and the various associated performance metrics. POS: positive; NEG: negative.



(a) CIFAR-100 (subsampled) (b) IMDB-WIKI (subsampled)

Evaluation metrics $\Rightarrow$ Learning goals matter!

Principled learning goals

fix precision, optimize recall (FPOR): $\max_{\theta, t} \text{recall}(f_{\theta}, t)$ s. t. $\text{precision}(f_{\theta}, t) \geq \alpha$,

fix recall, optimize precision (FROP): $\max_{\theta, t} \text{precision}_t$ s. t. $\text{recall}(f_{\theta}, t) \geq \alpha$,

optimize F_{β} score (OFBS): $\max_{\theta, t} F_{\beta}(f_{\theta}, t)$,

optimize AP (OAP): $\max_{\theta} \text{AP}(f_{\theta})$.

optimize multiclass performance (OMCP): $\max_{\theta, t} \text{multiclass-metric}(f_{\theta}, t)$.

optimize regression performance (OREGP): $\max_{\theta} \text{regression-metric}(f_{\theta})$;

Brand-new ideas for dealing with indicator functions

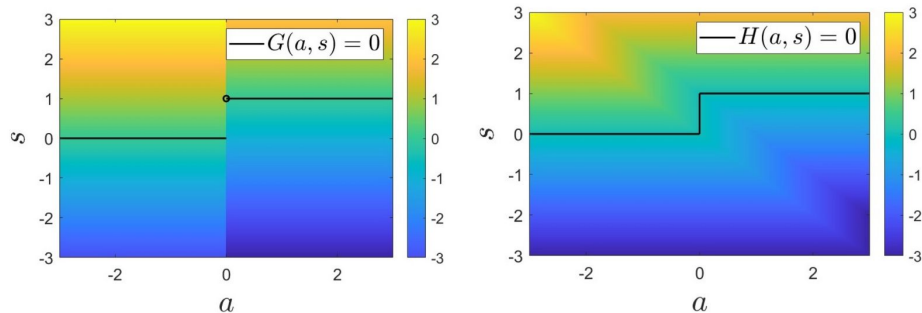
$$\max_{\boldsymbol{\theta}, t} \frac{1}{N_+} \sum_{i \in \mathcal{P}} \mathbf{1}\{f_{\boldsymbol{\theta}}(z_i) \leq t\}$$

Exact continuous reformulation of indicator function. Consider the following key observation:

$$(3.3) \quad s - \mathbf{1}\{a > 0\} = 0 \quad \stackrel{a \neq 0}{\iff} \quad s + [s + a - 1]_+ - [s + a]_+ = 0,$$

where $s \in \mathbb{R}$, $a \in \mathbb{R} \setminus \{0\}$, $[\cdot]_+ \doteq \max\{\cdot, 0\}$. To verify the validity of (3.3), we present in Figure 3 visualizations of the function values of $G(a, s)$ and $H(a, s)$, along with their level sets at 0, where

$$(3.4) \quad G(a, s) \doteq s - \mathbf{1}\{a > 0\}, \quad H(a, s) \doteq s + [s + a - 1]_+ - [s + a]_+.$$



Consistently (substantially) better results

Table 1: Objective values and feasibility for all compared methods on **FPOR**. Feasible solutions ($precision \geq 0.8$) are underlined, and among them, the highest objective values are **bolded**. Values in (parentheses) indicate results with optimized thresholds.

dataset	method	train		test	
		feasibility (precision)	objective (recall)	feasibility	objective
wilt	WCE	<u>0.872 ± 0.030</u> (0.886 ± 0.028)	1.000 ± 0.000 (1.000 ± 0.000)	0.776 ± 0.032 (0.7900 ± 0.023)	0.924 ± 0.026 (0.910 ± 0.010)
	TFCO	<u>0.867 ± 0.022</u> (0.874 ± 0.021)	1.000 ± 0.000 (1.000 ± 0.000)	0.811 ± 0.008 (0.825 ± 0.019)	0.924 ± 0.010 (0.917 ± 0.000)
	DMO	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.814 ± 0.023</u> (0.814 ± 0.023)	0.882 ± 0.049 (0.882 ± 0.049)
monks-3	WCE	<u>0.984 ± 0.000</u> (0.984 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.909 ± 0.009</u> (0.914 ± 0.002)	0.952 ± 0.022 (0.952 ± 0.022)
	TFCO	<u>0.984 ± 0.000</u> (0.984 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.954 ± 0.007</u> (0.954 ± 0.007)	0.982 ± 0.015 (0.982 ± 0.015)
	DMO	<u>0.984 ± 0.000</u> (0.984 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.874 ± 0.015</u> (0.916 ± 0.045)	0.946 ± 0.015 (0.744 ± 0.136)
breast-cancer-wisc	WCE	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.910 ± 0.012</u> (0.903 ± 0.000)	0.606 ± 0.028 (0.636 ± 0.000)
	TFCO	<u>0.954 ± 0.037</u> (0.955 ± 0.037)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.892 ± 0.019</u> (0.891 ± 0.018)	0.864 ± 0.074 (0.848 ± 0.084)
	DMO	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.858 ± 0.044</u> (0.858 ± 0.044)	0.765 ± 0.028 (0.765 ± 0.028)
eyepacs	WCE	0.680 ± 0.005 (0.800 ± 0.000)	0.186 ± 0.028 (0.035 ± 0.006)	0.651 ± 0.006 (0.7970 ± 0.014)	0.200 ± 0.026 (0.037 ± 0.007)
	TFCO	0.259 ± 0.008 (0.7060 ± 0.297)	0.527 ± 0.335 (0.001 ± 0.000)	0.262 ± 0.010 (0.4830 ± 0.103)	0.527 ± 0.344 (0.001 ± 0.001)
	DMO	<u>0.804 ± 0.004</u> (0.800 ± 0.000)	0.311 ± 0.002 (0.317 ± 0.007)	<u>0.775 ± 0.004</u> (0.7710 ± 0.001)	0.308 ± 0.001 (0.313 ± 0.006)
wildfire	WCE	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>0.973 ± 0.009</u> (0.966 ± 0.009)	1.000 ± 0.000 (1.000 ± 0.000)
	TFCO	0.236 ± 0.070 (0.7670 ± 0.206)	0.595 ± 0.288 (0.012 ± 0.008)	0.210 ± 0.091 (0.3330 ± 0.471)	0.549 ± 0.315 (0.021 ± 0.030)
	DMO	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)	<u>1.000 ± 0.000</u> (1.000 ± 0.000)	1.000 ± 0.000 (1.000 ± 0.000)
adc-v2	WCE	0.717 ± 0.007 (0.800 ± 0.000)	0.883 ± 0.002 (0.786 ± 0.013)	0.720 ± 0.006 (0.7940 ± 0.000)	0.886 ± 0.001 (0.772 ± 0.014)
	TFCO	0.285 ± 0.028 (0.4630 ± 0.094)	0.639 ± 0.256 (0.001 ± 0.001)	0.290 ± 0.027 (0.2540 ± 0.184)	0.652 ± 0.248 (0.001 ± 0.001)
	DMO	<u>0.800 ± 0.000</u> (0.800 ± 0.000)	0.837 ± 0.001 (0.809 ± 0.040)	<u>0.786 ± 0.002</u> (0.7870 ± 0.003)	0.823 ± 0.002 (0.792 ± 0.044)

More details

[Submitted on 21 Jul 2025]

Exact Reformulation and Optimization for Direct Metric Optimization in Binary Imbalanced Classification

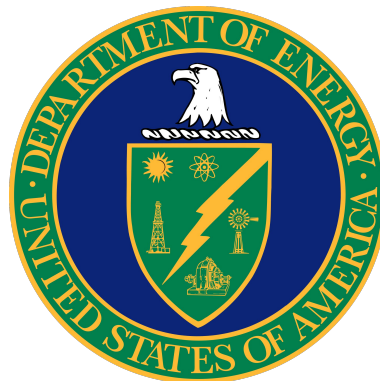
Le Peng, Yash Travadi, Chuan He, Ying Cui, Ju Sun

For classification with imbalanced class frequencies, i.e., imbalanced classification (IC), standard accuracy is known to be misleading as a performance measure. While most existing methods for IC resort to optimizing balanced accuracy (i.e., the average of class-wise recalls), they fall short in scenarios where the significance of classes varies or certain metrics should reach prescribed levels. In this paper, we study two key classification metrics, precision and recall, under three practical binary IC settings: fix precision optimize recall (FPOR), fix recall optimize precision (FROP), and optimize F_β -score (OFBS). Unlike existing methods that rely on smooth approximations to deal with the indicator function involved, we introduce, for the first time, exact constrained reformulations for these direct metric optimization (DMO) problems, which can be effectively solved by exact penalty methods. Experiment results on multiple benchmark datasets demonstrate the practical superiority of our approach over the state-of-the-art methods for the three DMO problems. We also expect our exact reformulation and optimization (ERO) framework to be applicable to a wide range of DMO problems for binary IC and beyond. Our code is available at [this https URL](#).

Today's talk: **toward trustworthy and efficient AI**

- General and reliable robustness evaluation and improved and robust selective classification
- Effective Inverse problems with pretrained flow-based models
- AI for healthcare: handling data imbalance via direct metric optimization

Thanks to



Thanks to

PhD's

- Yifan Wu (BCIB, co-advised with)
- Qiaozhi Huang (CS&E)
- Lingjie Su (CS&E)
- Harsh Hemant Shah (CS&E, co-
- Sinian Zhang (Biostats, co-advised)
- Guanchen Li (CS&E)
- Gaoxiang Luo (CS&E)
- Yuxiang Wan (CS&E)
- Corey Senger (CS&E, co-advised)
- Wenjie Zhang (CS&E)
- Jiandong Chen (IHI)
- Ryan de Vera (CS&E)

PhD & Postdoc Alumni

- Hengkang Wang (PhD'25 with thesis pending)
- Tiancong Chen (PhD'24 with thesis pending, @)
- Yash Travadi (PhD'24 with thesis pending)
- Le Peng (PhD'24 with thesis pending, @)
- Taihui Li (PhD'24 with thesis pending, @)
- Hengyue Liang (PhD'25, Applied Science)
- Chuan He (Postdoc'23–24, Assistant Professor)
- Zhong Zhuang (PhD'23 [thesis], Postdoc)