



THE UNIVERSITY OF
MELBOURNE

Improving Accuracy-robustness Trade-off via Pixel Reweighted Adversarial Training

Presenter: Jiacheng Zhang

School of Computing and Information Systems

The University of Melbourne

Date: 23 January 2025



What is an adversarial example (attack)?

☐ **Left-or-right challenge:** Guess which one is the adversarial example?



What is an adversarial example (attack)?

99% Guacamole



88% Tabby Cat



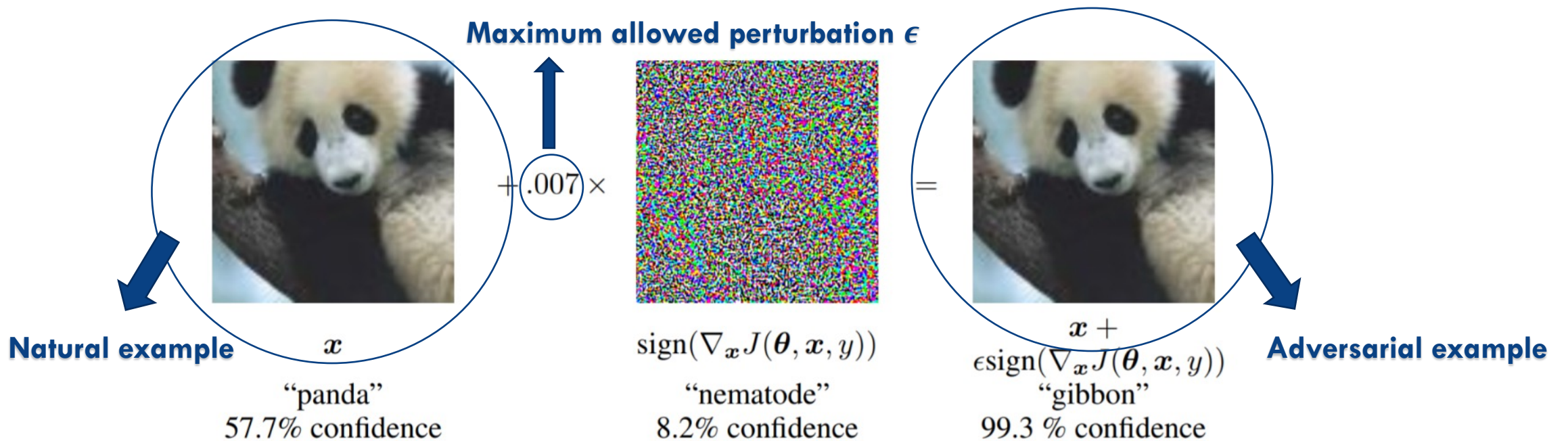
What is an adversarial example (attack)?

Adversarial examples can significantly drop the classification accuracy to **0%**.

How it works?

What is an adversarial example (attack)?

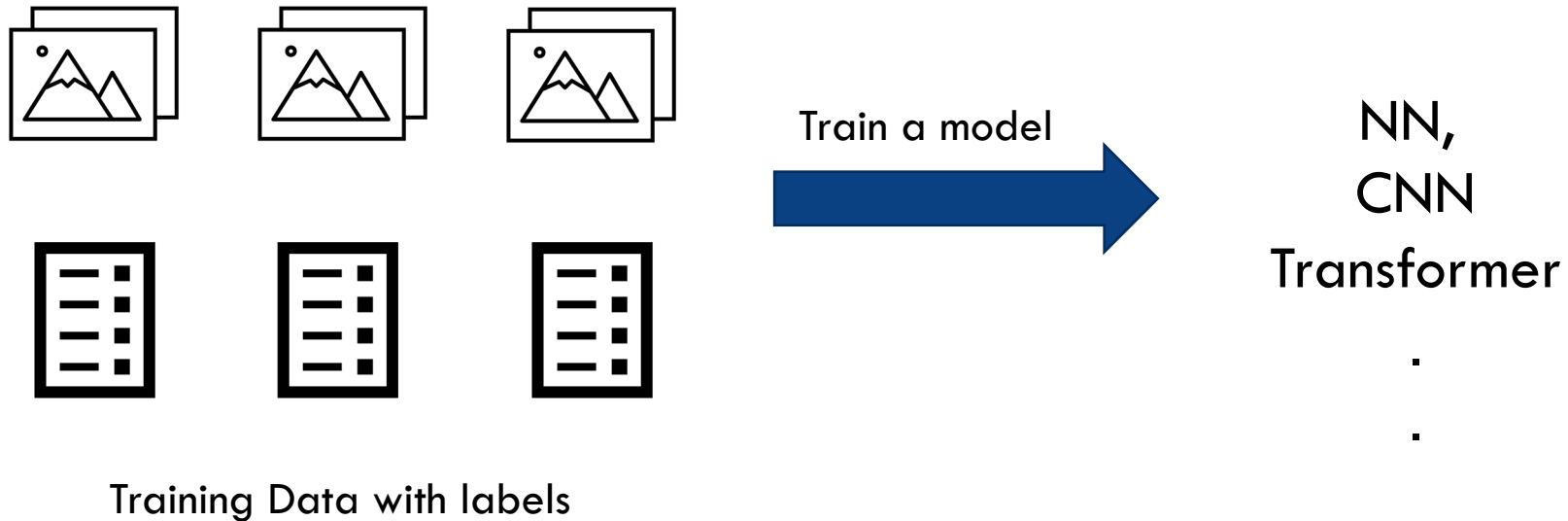
- Adding **imperceptible, non-random perturbations** to input data.



- Cannot fool human eyes, but **can easily fool** state-of-the-art neural networks.

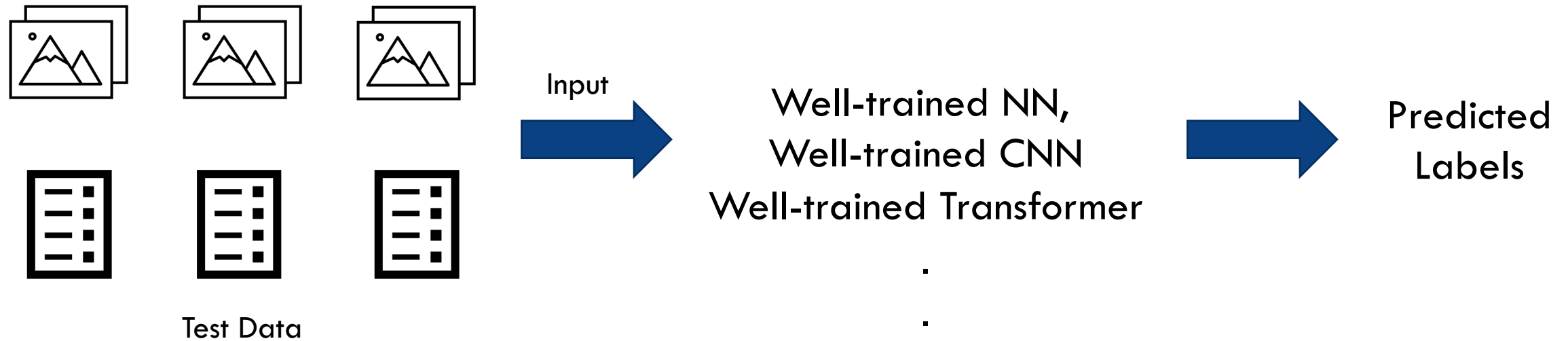
What is an adversarial example (attack)?

Conventional Machine Learning Pipeline (classification):



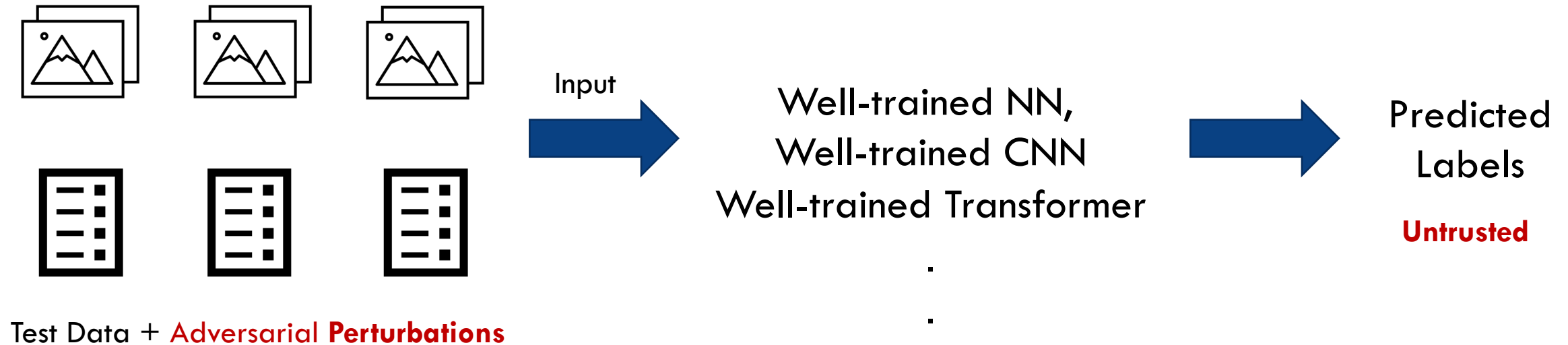
What is an adversarial example (attack)?

Conventional Machine Learning Pipeline (classification):



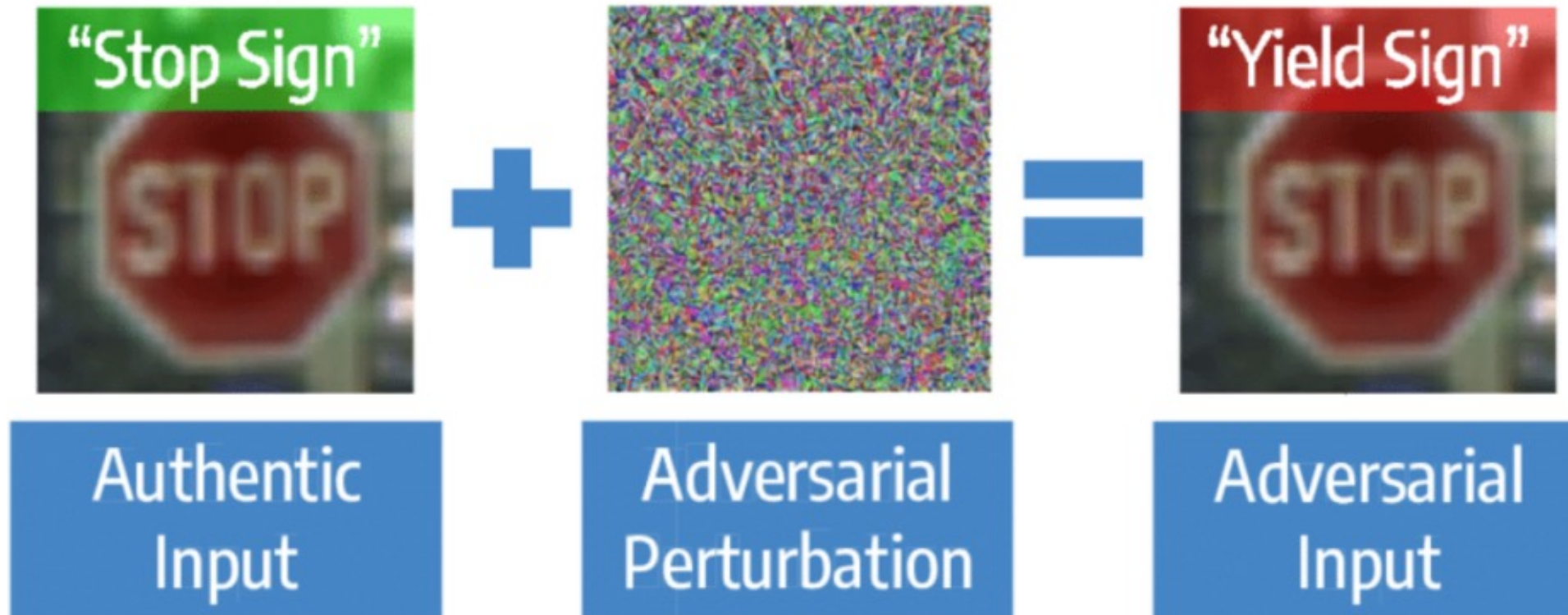
What is an adversarial example (attack)?

Adversarial attack happens:

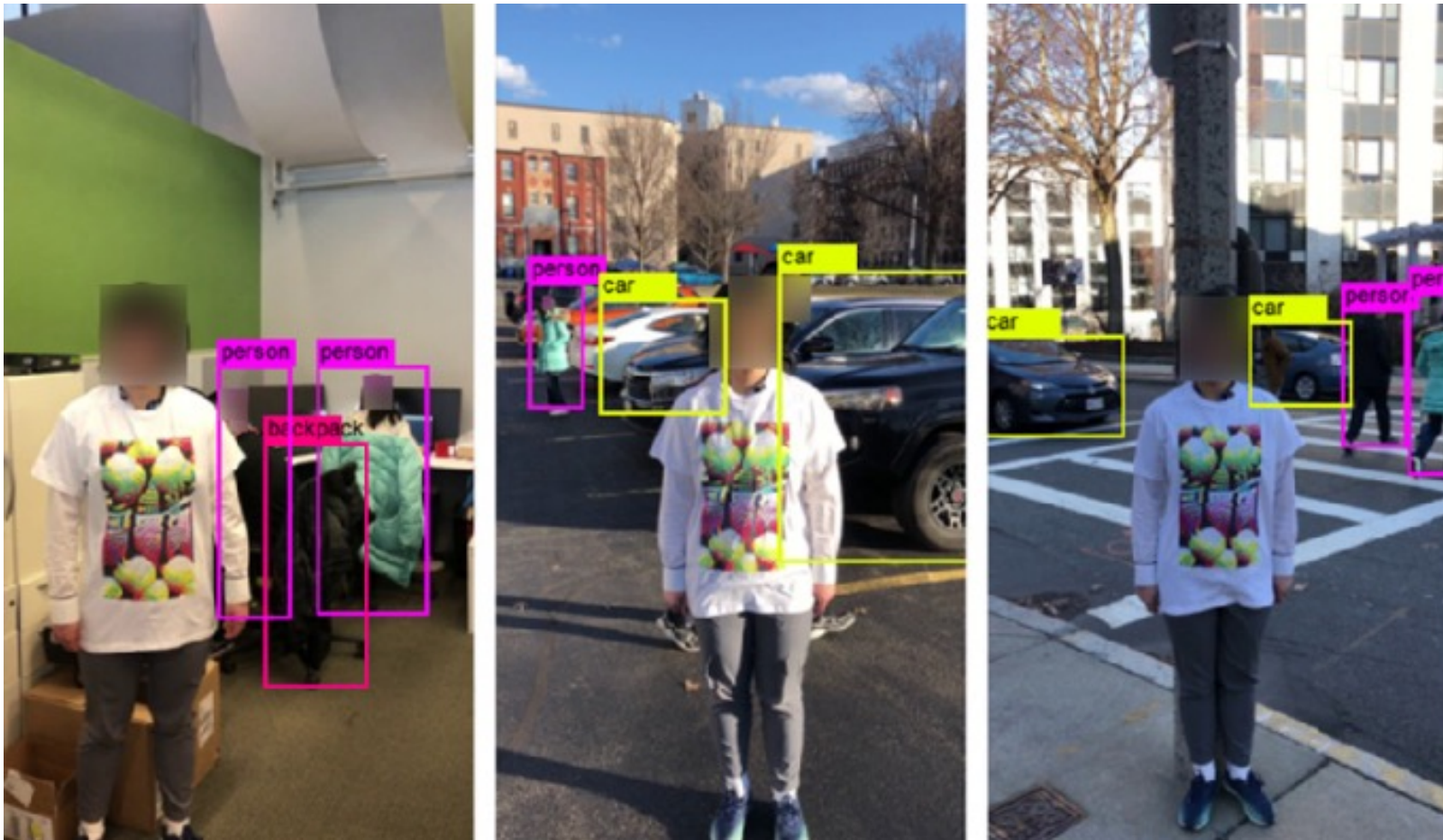


Why do we care?

- ❑ Cause **security and reliability issues** in the deployment of machine learning systems.
- ❑ E.g., mislead the autonomous driving system to recognize **a stop sign** into **something else**.

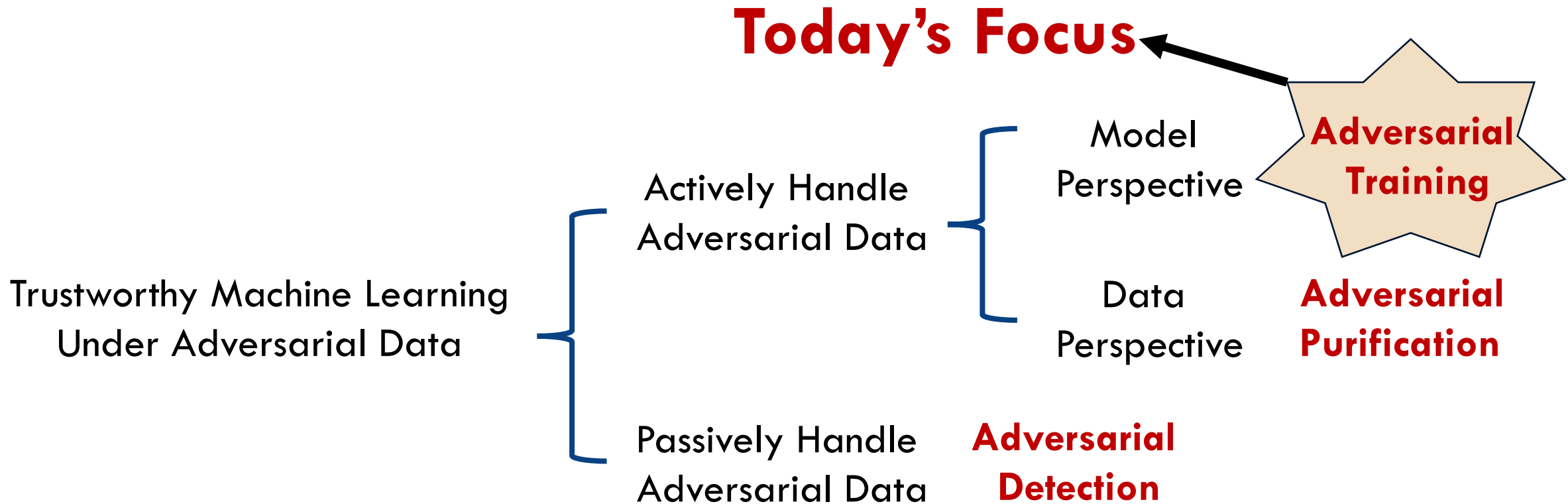


Why do we care?



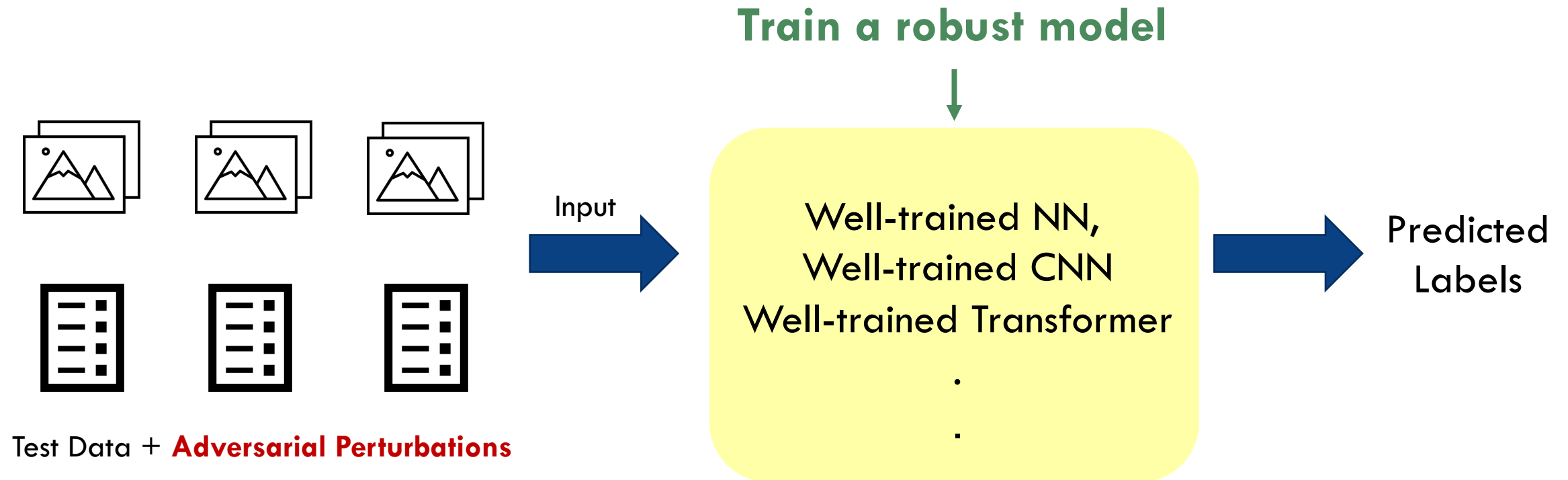
- ❑ Adding **adversarial examples** on T-shirts can bypass the AI detection system.
- ❑ Let you be invisible to the AI detection system!
- ❑ It's cool but it can cause **security and reliability issues.**

How to defend against adversarial attacks?



Adversarial training

- ❑ *Adversarial Training (AT)*: aims to train a robust model on *adversarial examples (AEs)*

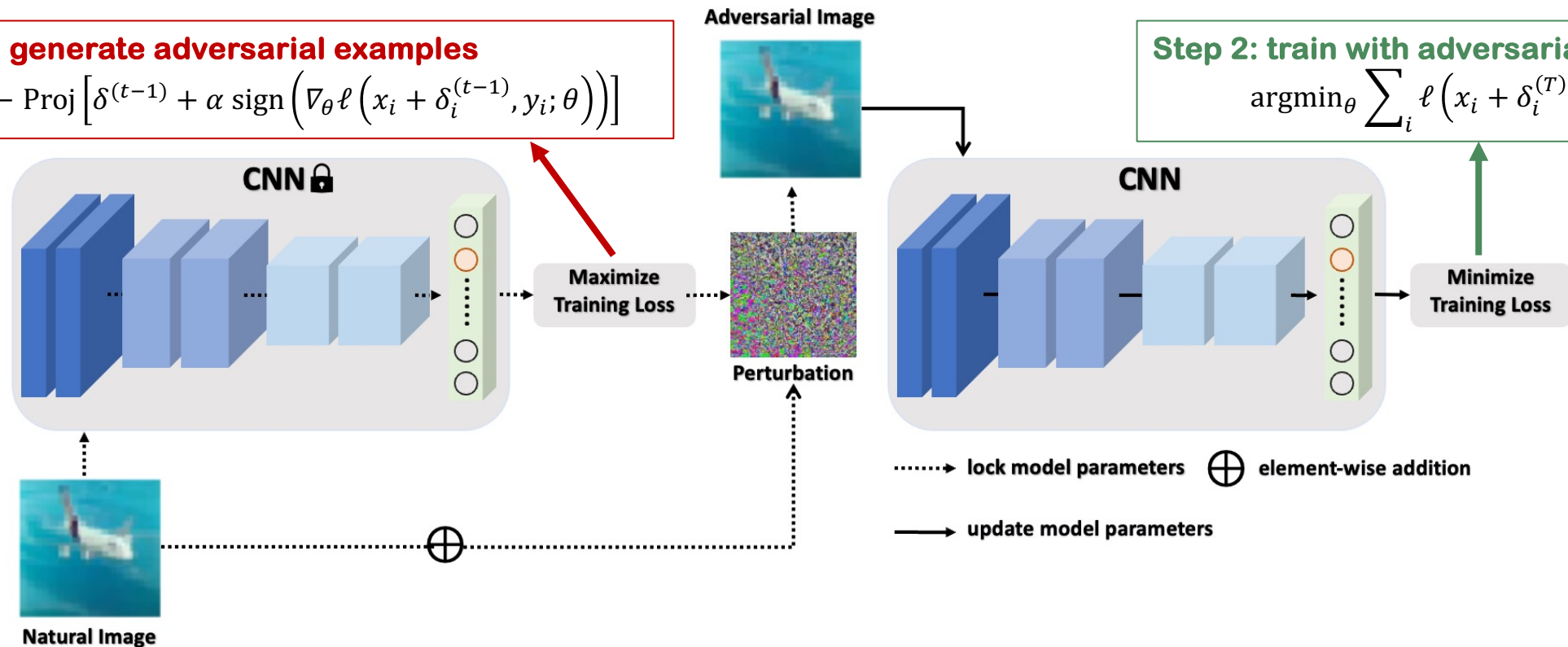


Adversarial training

□ *Adversarial Training (AT)*: aims to train a robust model on *adversarial examples (AEs)*

Step 1: generate adversarial examples

$$\delta^{(t)} \leftarrow \text{Proj} \left[\delta^{(t-1)} + \alpha \text{sign} \left(\nabla_{\theta} \ell \left(x_i + \delta_i^{(t-1)}, y_i; \theta \right) \right) \right]$$



What's the problem of adversarial training?

What's good about AT?

- ❑ AT improves robustness (i.e., accuracy on adversarial examples).

What's bad about AT?

- ❑ AT drops natural accuracy (i.e., accuracy on natural examples).

Problem: there is an accuracy-robustness trade-off!!!

- ❑ Higher robustness, lower accuracy; higher accuracy, lower robustness.

What we really want: a model that has high robustness without sacrificing natural accuracy.

How to mitigate this problem?

What will happen if we change ϵ ?

$$\begin{array}{ccc}
 \begin{array}{c} \text{Panda image} \\ x \\ \text{"panda"} \\ 57.7\% \text{ confidence} \end{array} & \begin{array}{c} \uparrow \\ + \textcircled{.007} \times \\ \text{Noise image} \\ \text{sign}(\nabla_x J(\theta, x, y)) \\ \text{"nematode"} \\ 8.2\% \text{ confidence} \end{array} & = \\
 & & \begin{array}{c} \text{Panda image} \\ x + \epsilon \text{sign}(\nabla_x J(\theta, x, y)) \\ \text{"gibbon"} \\ 99.3\% \text{ confidence} \end{array}
 \end{array}$$

- ❑ **Extreme case:** we decrease ϵ to 0, AT will converge to natural training.
- ❑ **Conclusion:** decrease the budget can improve natural accuracy.

Research gap and research question

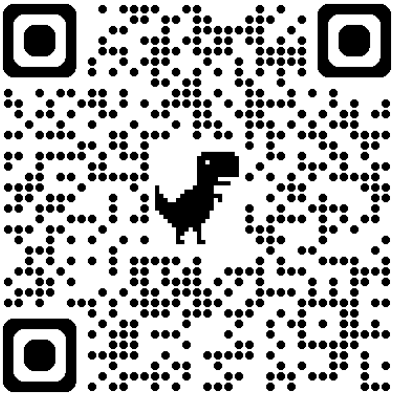
- ❑ **Research gap:** Existing AT methods apply a fixed ϵ for all pixels in an image. Therefore, changing ϵ must sacrifice natural accuracy or robustness.
- ❑ **Research question:** Can we design an adaptive method to **reweight ϵ for only partial pixels in an image** so that we can increase natural accuracy without sacrificing robustness?
- ❑ In our recent work, we show that the answer to this question is **YES**.

ICML 2024

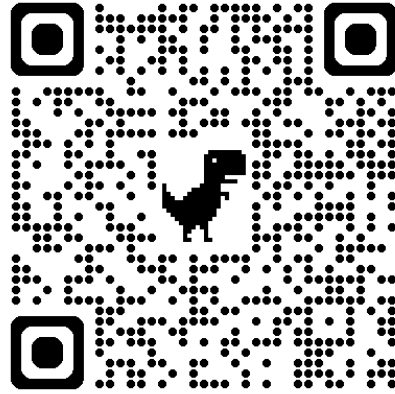
Improving Accuracy-robustness Trade-off via Pixel Reweighted Adversarial Training

Jiacheng Zhang, Feng Liu*, Dawei Zhou, Jingfeng Zhang, Tongliang Liu*

Paper

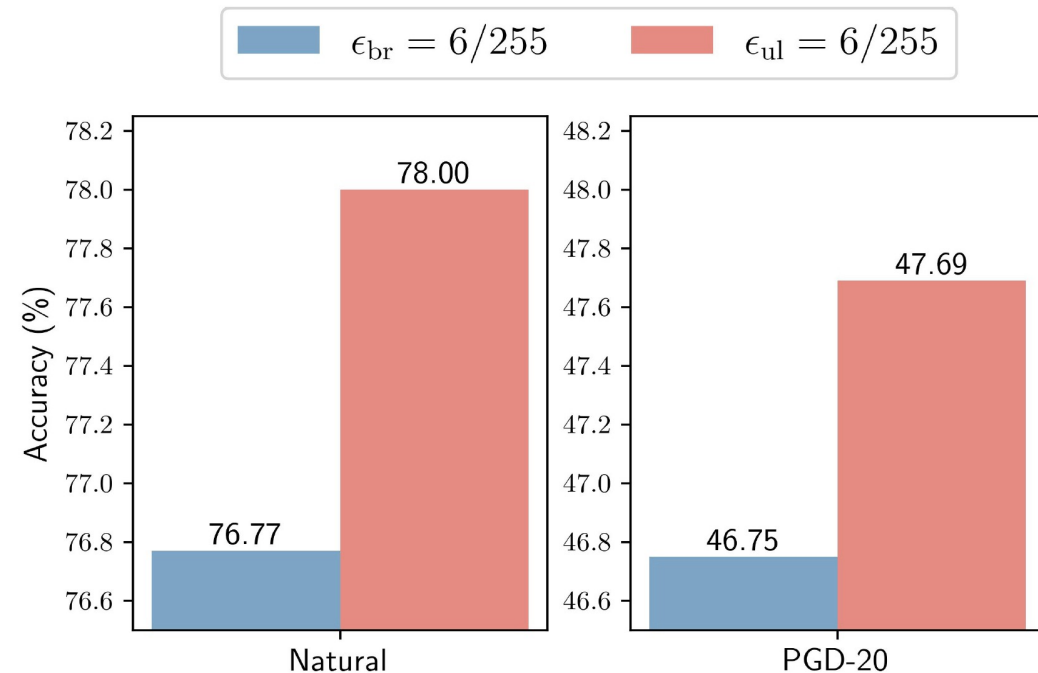
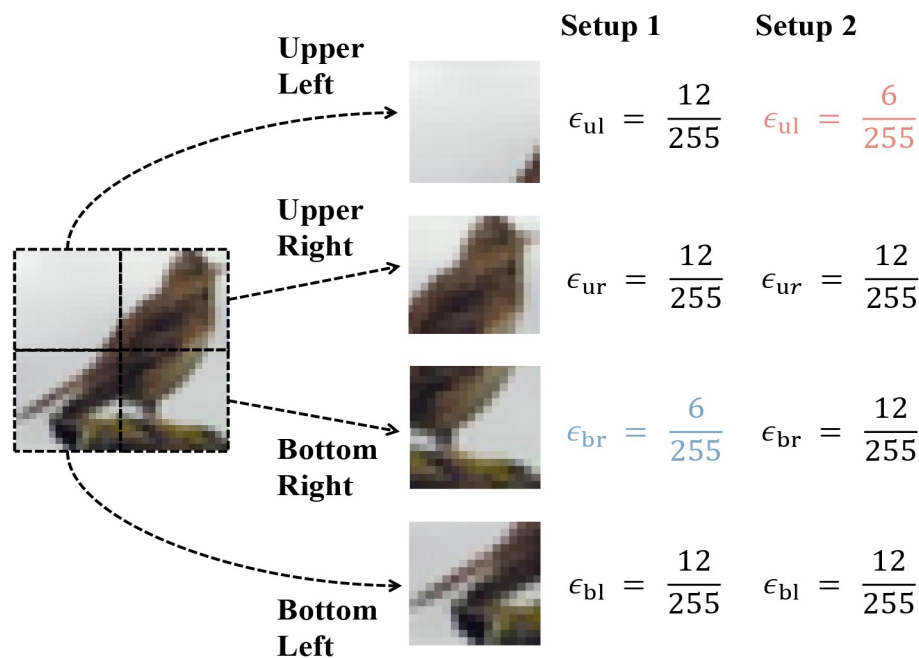


Code



Are all pixels equally important in robust classification?

- ❑ **Proof-of-concept experiment:** changing the perturbation budgets for different parts of an image has the potential to boost robustness and accuracy **at the same time**.



Pixel-reweighted Adversarial Training (PART)

- ❑ **Proof-of-concept experiment:** changing the perturbation budgets for different parts of an image has the potential to boost robustness and accuracy **at the same time**.



Pixels in an image have different importances in classification.

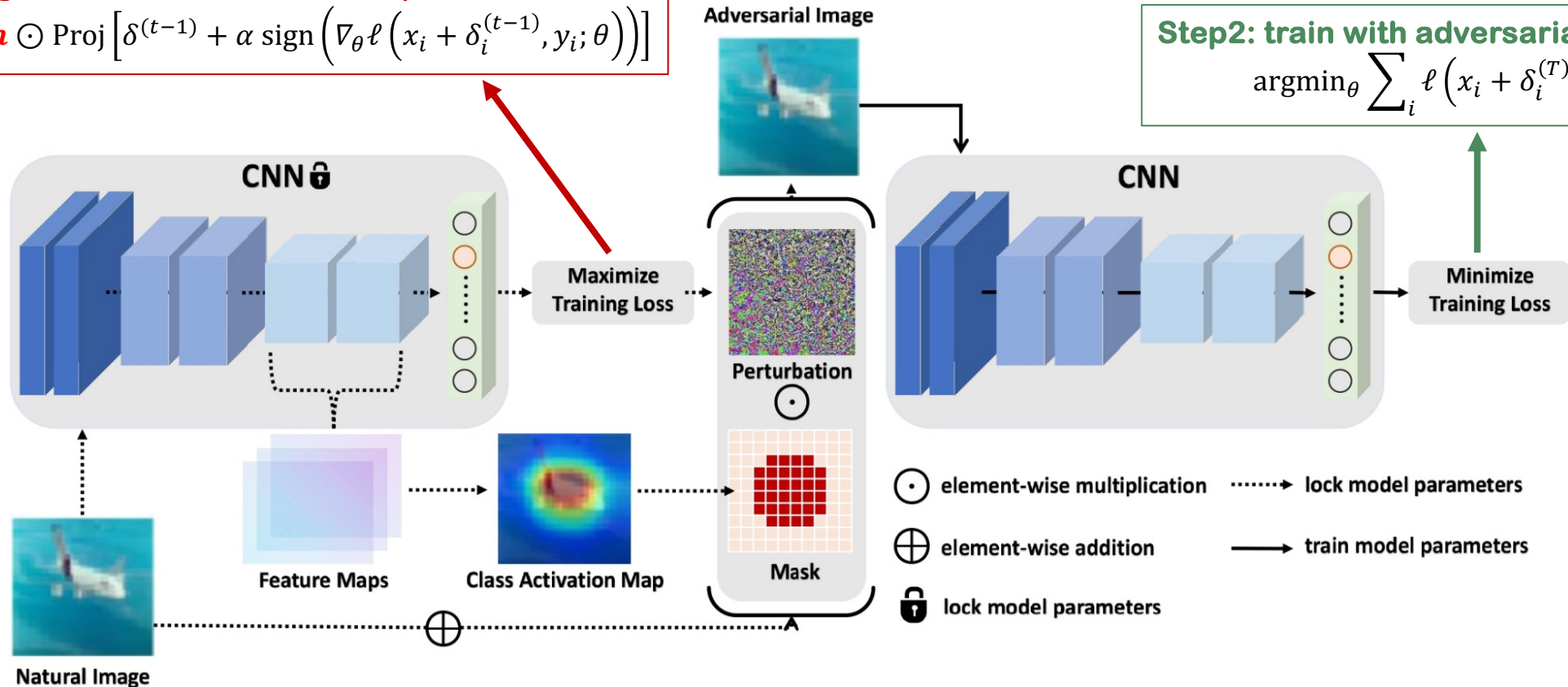


We need to guide the model to **focus more on important pixels** during training.

Pixel-reweighted Adversarial Training (PART)

Step1: generate adversarial examples

$$\delta^{(t)} \leftarrow m \odot \text{Proj} \left[\delta^{(t-1)} + \alpha \text{sign} \left(\nabla_{\theta} \ell \left(x_i + \delta_i^{(t-1)}, y_i; \theta \right) \right) \right]$$



Main results

Dataset	Method	Natural	PGD-20	MMA	AA
ResNet-18					
CIFAR-10	AT	82.58 $\pm$ 0.14	43.69 $\pm$ 0.28	41.80 $\pm$ 0.10	41.63 $\pm$ 0.22
	PART ($s = 1$)	83.42 $\pm$ 0.26 (+ 0.84)	43.65 $\pm$ 0.16 (- 0.04)	41.98 $\pm$ 0.03 (+ 0.18)	41.74 $\pm$ 0.04 (+ 0.11)
	PART ($s = 10$)	83.77 $\pm$ 0.15 (+ 1.19)	43.36 $\pm$ 0.21 (- 0.33)	41.83 $\pm$ 0.07 (+ 0.03)	41.41 $\pm$ 0.14 (- 0.22)
	TRADES	78.16 $\pm$ 0.15	48.28 $\pm$ 0.05	45.00 $\pm$ 0.08	45.05 $\pm$ 0.12
	PART-T ($s = 1$)	79.36 $\pm$ 0.31 (+ 1.20)	48.90 $\pm$ 0.14 (+ 0.62)	45.90 $\pm$ 0.07 (+ 0.90)	45.97 $\pm$ 0.06 (+ 0.92)
	PART-T ($s = 10$)	80.13 $\pm$ 0.16 (+ 1.97)	48.72 $\pm$ 0.11 (+ 0.44)	45.59 $\pm$ 0.09 (+ 0.59)	45.60 $\pm$ 0.04 (+ 0.55)
	MART	76.82 $\pm$ 0.28	49.86 $\pm$ 0.32	45.42 $\pm$ 0.04	45.10 $\pm$ 0.06
	PART-M ($s = 1$)	78.67 $\pm$ 0.10 (+ 1.85)	50.26 $\pm$ 0.17 (+ 0.40)	45.53 $\pm$ 0.05 (+ 0.11)	45.19 $\pm$ 0.04 (+ 0.09)
	PART-M ($s = 10$)	80.00 $\pm$ 0.15 (+ 3.18)	49.71 $\pm$ 0.12 (- 0.15)	45.14 $\pm$ 0.10 (- 0.28)	44.61 $\pm$ 0.24 (- 0.49)
ResNet-18					
SVHN	AT	91.06 $\pm$ 0.24	49.83 $\pm$ 0.13	47.68 $\pm$ 0.06	45.48 $\pm$ 0.05
	PART ($s = 1$)	93.14 $\pm$ 0.05 (+ 2.08)	50.34 $\pm$ 0.14 (+ 0.51)	48.08 $\pm$ 0.09 (+ 0.40)	45.67 $\pm$ 0.13 (+ 0.19)
	PART ($s = 10$)	93.75 $\pm$ 0.07 (+ 2.69)	50.21 $\pm$ 0.10 (+ 0.38)	48.00 $\pm$ 0.14 (+ 0.32)	45.61 $\pm$ 0.08 (+ 0.13)
	TRADES	88.91 $\pm$ 0.28	58.74 $\pm$ 0.53	53.29 $\pm$ 0.56	52.21 $\pm$ 0.47
	PART-T ($s = 1$)	91.35 $\pm$ 0.11 (+ 2.44)	59.33 $\pm$ 0.22 (+ 0.59)	54.04 $\pm$ 0.16 (+ 0.75)	53.07 $\pm$ 0.67 (+ 0.86)
	PART-T ($s = 10$)	91.94 $\pm$ 0.18 (+ 3.03)	59.01 $\pm$ 0.13 (+ 0.27)	53.80 $\pm$ 0.20 (+ 0.51)	52.61 $\pm$ 0.24 (+ 0.40)
	MART	89.76 $\pm$ 0.08	58.52 $\pm$ 0.53	52.42 $\pm$ 0.34	49.10 $\pm$ 0.23
	PART-M ($s = 1$)	91.42 $\pm$ 0.36 (+ 1.66)	58.85 $\pm$ 0.29 (+ 0.33)	52.45 $\pm$ 0.03 (+ 0.03)	49.92 $\pm$ 0.10 (+ 0.82)
	PART-M ($s = 10$)	93.20 $\pm$ 0.22 (+ 3.44)	58.41 $\pm$ 0.20 (- 0.11)	52.18 $\pm$ 0.14 (- 0.24)	49.25 $\pm$ 0.13 (+ 0.15)
WideResNet-34-10					
TinyImagenet-200	AT	43.51 $\pm$ 0.13	11.70 $\pm$ 0.08	10.66 $\pm$ 0.11	10.53 $\pm$ 0.14
	PART ($s = 1$)	44.87 $\pm$ 0.21 (+ 1.36)	11.93 $\pm$ 0.16 (+ 0.23)	10.96 $\pm$ 0.12 (+ 0.30)	10.76 $\pm$ 0.06 (+ 0.23)
	PART ($s = 10$)	45.59 $\pm$ 0.14 (+ 2.08)	11.81 $\pm$ 0.10 (+ 0.11)	10.91 $\pm$ 0.08 (+ 0.25)	10.68 $\pm$ 0.10 (+ 0.15)
	TRADES	43.05 $\pm$ 0.15	13.86 $\pm$ 0.10	12.62 $\pm$ 0.16	12.55 $\pm$ 0.09
	PART-T ($s = 1$)	44.31 $\pm$ 0.12 (+ 1.26)	14.08 $\pm$ 0.22 (+ 0.22)	13.01 $\pm$ 0.09 (+ 0.39)	12.84 $\pm$ 0.14 (+ 0.29)
	PART-T ($s = 10$)	45.16 $\pm$ 0.10 (+ 2.11)	13.98 $\pm$ 0.15 (+ 0.12)	12.88 $\pm$ 0.12 (+ 0.26)	12.72 $\pm$ 0.08 (+ 0.17)
	MART	42.68 $\pm$ 0.22	14.77 $\pm$ 0.18	13.58 $\pm$ 0.13	13.42 $\pm$ 0.16
	PART-M ($s = 1$)	43.75 $\pm$ 0.24 (+ 1.07)	14.93 $\pm$ 0.15 (+ 0.16)	13.76 $\pm$ 0.06 (+ 0.18)	13.68 $\pm$ 0.13 (+ 0.24)
	PART-M ($s = 10$)	45.02 $\pm$ 0.16 (+ 2.34)	14.65 $\pm$ 0.14 (- 0.12)	13.41 $\pm$ 0.11 (- 0.17)	13.37 $\pm$ 0.15 (- 0.05)

- Our method can achieve a notable improvement in accuracy-robustness trade-off.
- In most cases, our method can improve natural accuracy by a notable margin without sacrificing robustness.

Conclusion and future work

- ❑ **Key message we want to deliver in this paper:** Guiding the model to focus more on essential pixel regions during training can help improve the accuracy-robustness trade-off of vision models.

Questions?



THE UNIVERSITY OF
MELBOURNE

Thank you