

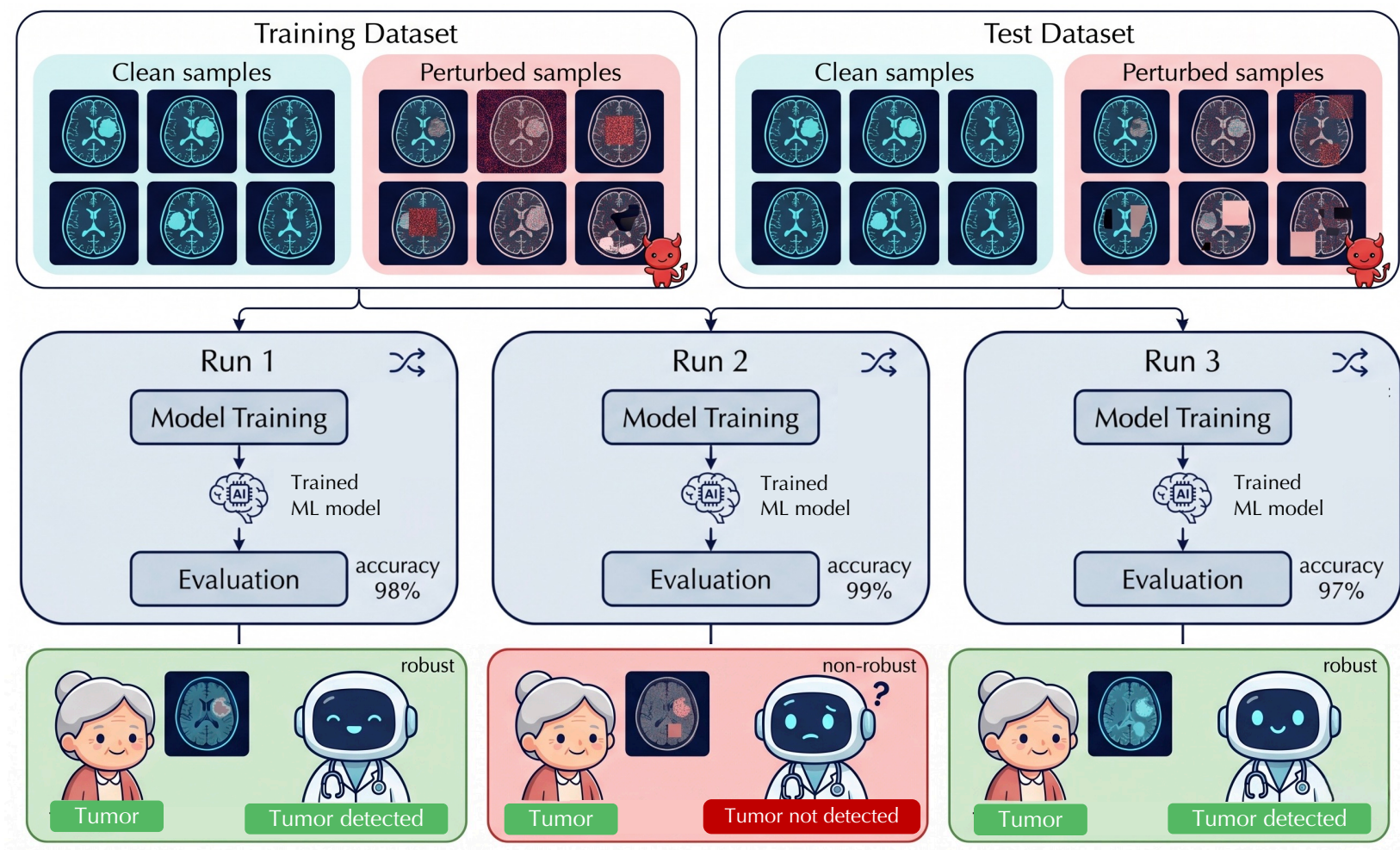
# Probabilistic Robustness Certificates against Adversarial Attacks

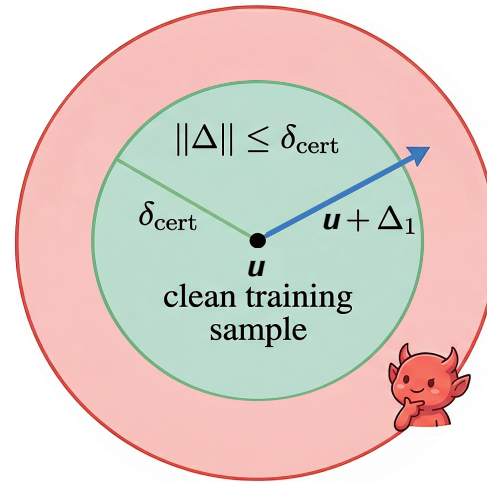
Sara Taheri<sup>1</sup>, Majid Zamani<sup>1,2</sup>

<sup>1</sup> LMU Munich, Germany

<sup>2</sup> CU Boulder, USA

ICML 2026

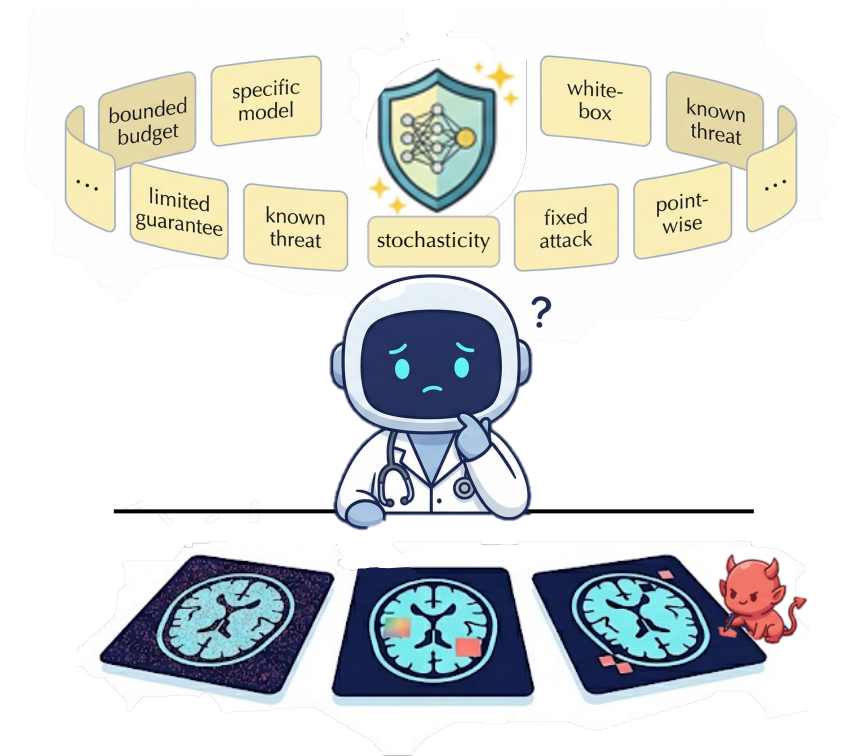


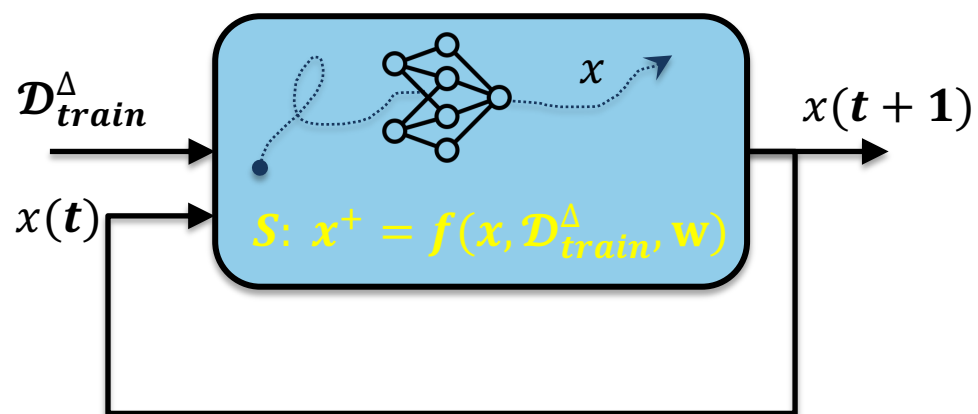
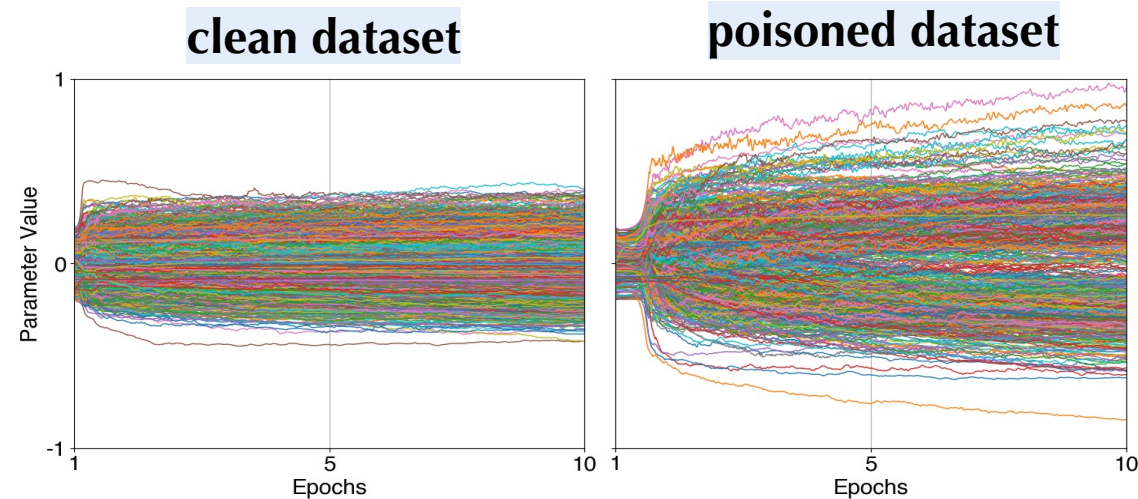


## Robustness Problem

**Find:** The maximum certified poisoning radius ( $\delta_{\text{cert}}$ )

**Guarantee:** For any attacker budget bounded by ( $\|\Delta\|_p \leq \delta_{\text{cert}}$ ) the model's accuracy strictly stays above the safety threshold with a certified probability.

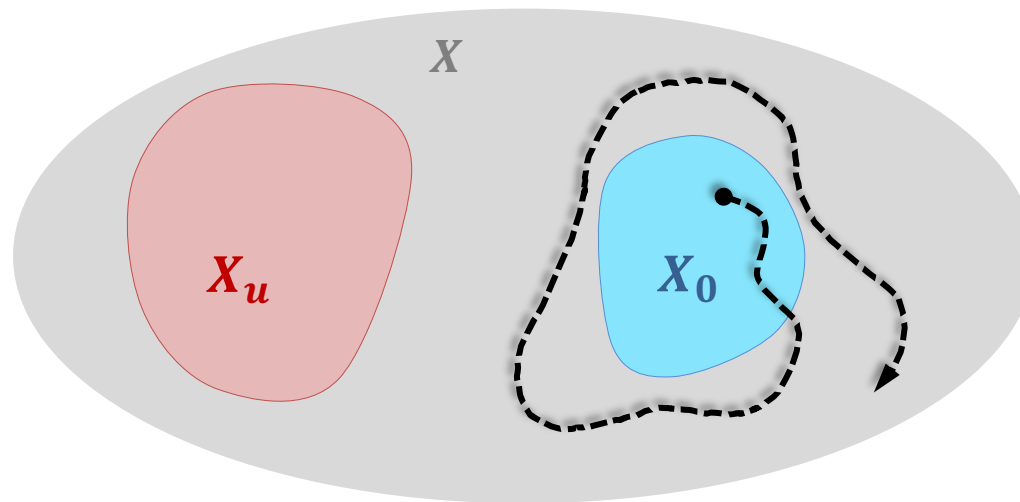




- **State:** model parameters  $x \in X \subseteq \mathbb{R}^d$
- **Input:** poisoned training data  $\mathcal{D}_{train}$
- **Dynamics:** Parameter update map  $f$
- **Uncertainty:** Stochasticity in ML pipeline  $w \in \mathcal{W}$

## Safe set ( $X_s$ ):

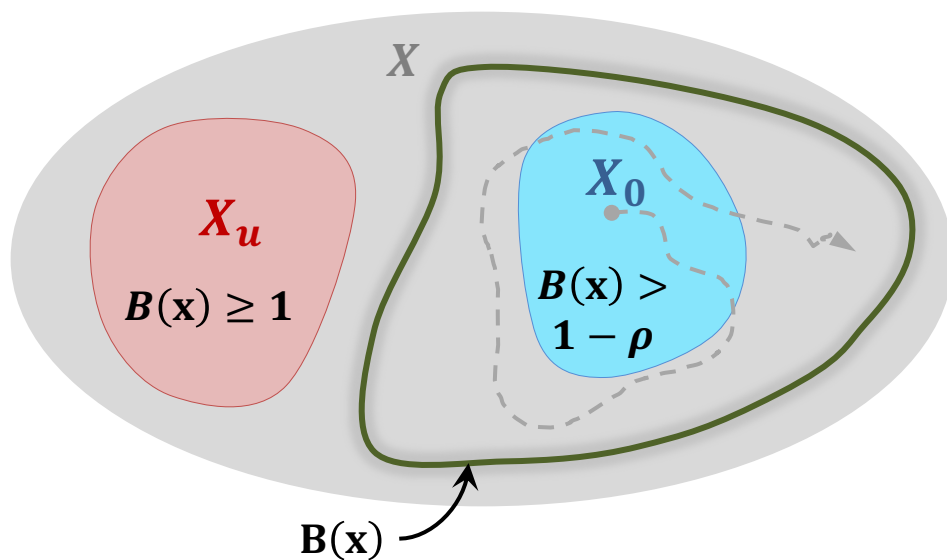
The region in the parameter space where the model's accuracy remains above the critical threshold.



$X_0 \subseteq X$ : Initial parameter set  
 $X_u \subseteq X$ : Unsafe parameter set  
 $X$ : Parameter set

*Can we certify that the final trained parameters remain in the safe set with some probability?*

$X_0 \subseteq X$ : Initial set  
 $X_u \subseteq X$ : Unsafe set  
 $X$ : State set

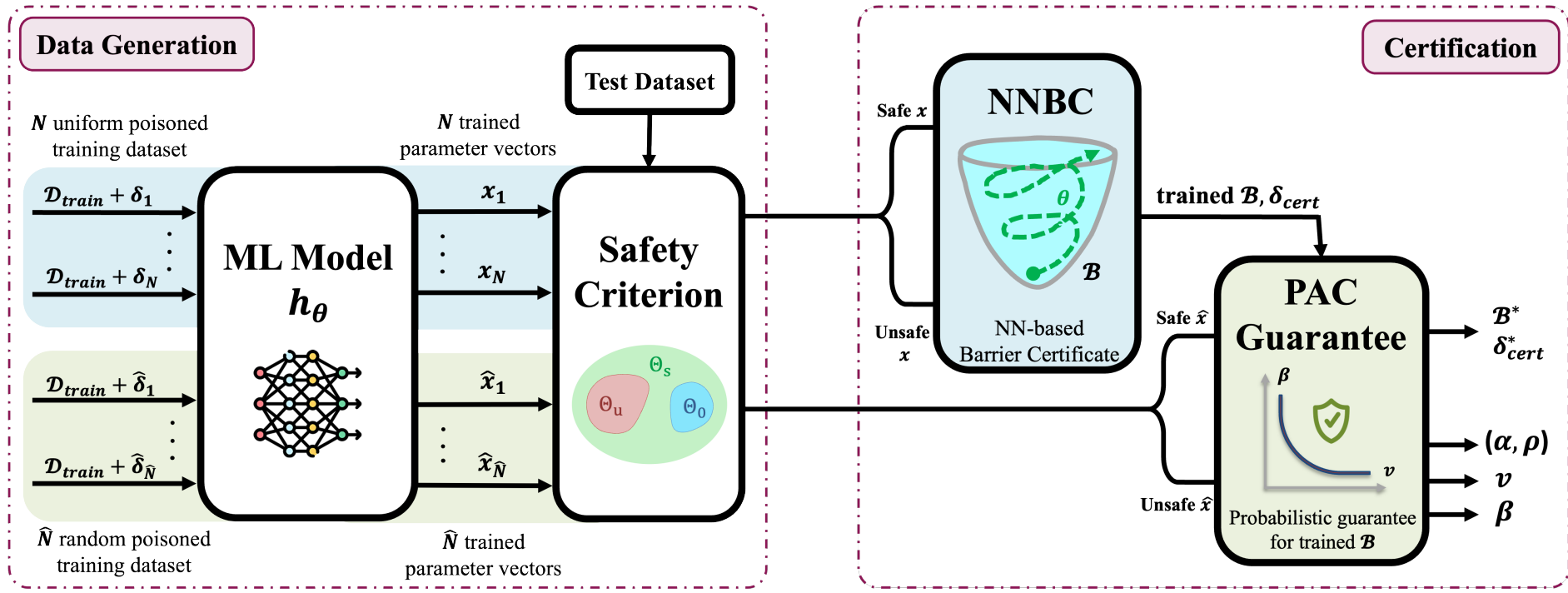


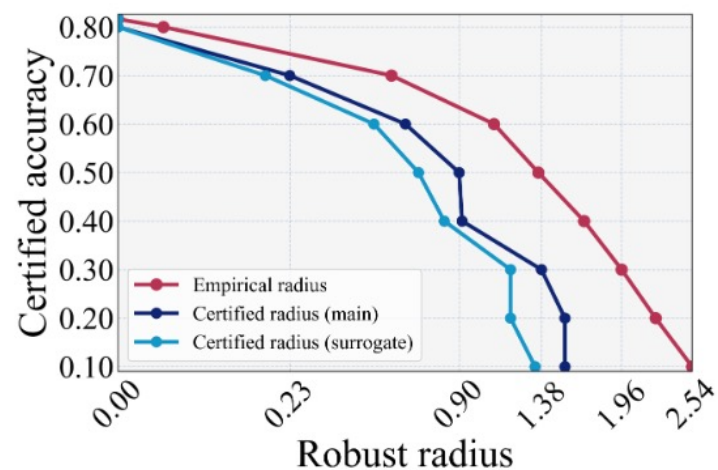
### Theorem. BC for a dt-DS

Given a threshold  $\rho \in [0,1]$ , If there exists a function  $\mathbf{B}$  and a certified robust radius  $\delta$  such that:

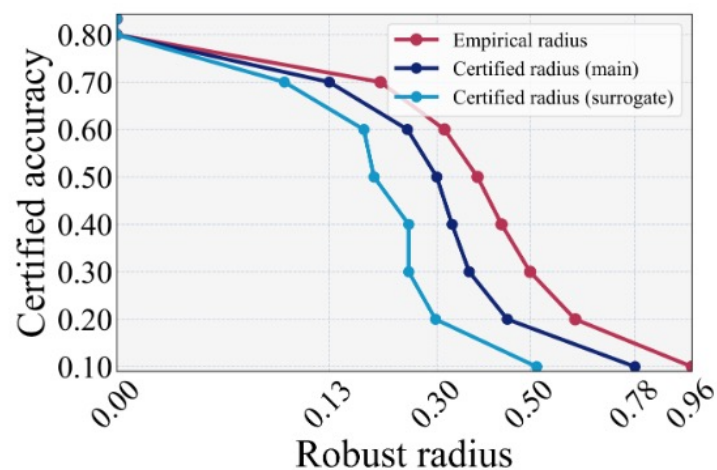
- $\mathbf{B}(x, 0) \leq 1 - \rho, \quad \forall x \in X_0,$
  - $\mathbf{B}(x, T) \geq 1, \quad \forall x \in X_u,$
  - $\mathbb{E}[\mathbf{B}(x^+, t) \mid x] \leq \mathbf{B}(x, t - 1), \quad \forall x \in X, \|\Delta\|_p \leq \delta.$
- $\Rightarrow x(t) \notin X_u$  with probability at least  $\rho$ .



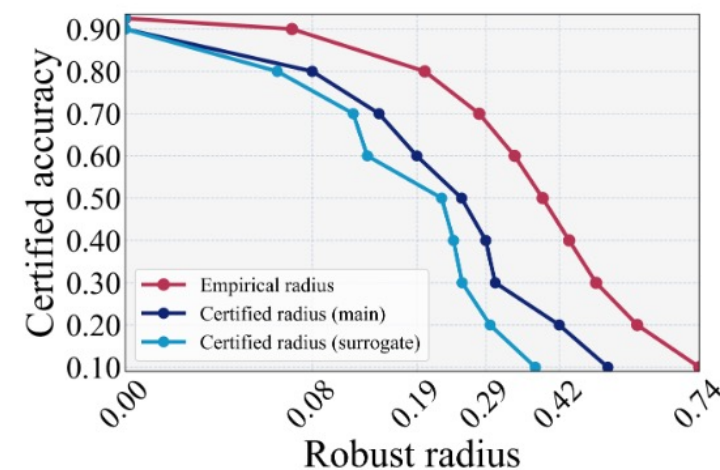




(a) SVHN, CNN, SGD, BPA ( $\ell_2, \tau = 0.2$ )



(b) CIFAR10, ResNet, Adam, PGD ( $\ell_\infty, \tau = 0.8$ )



(c) MNIST, CNN, SGD, BPA ( $\ell_2, \tau = 0.3$ )



*Our framework provides model- and attack-agnostic probabilistic certificates for robustness under stochastic training, using neural barrier certificates and PAC-style verification.*

**Thanks for your attention!**



HyConSys Lab



University of Colorado  
Boulder