

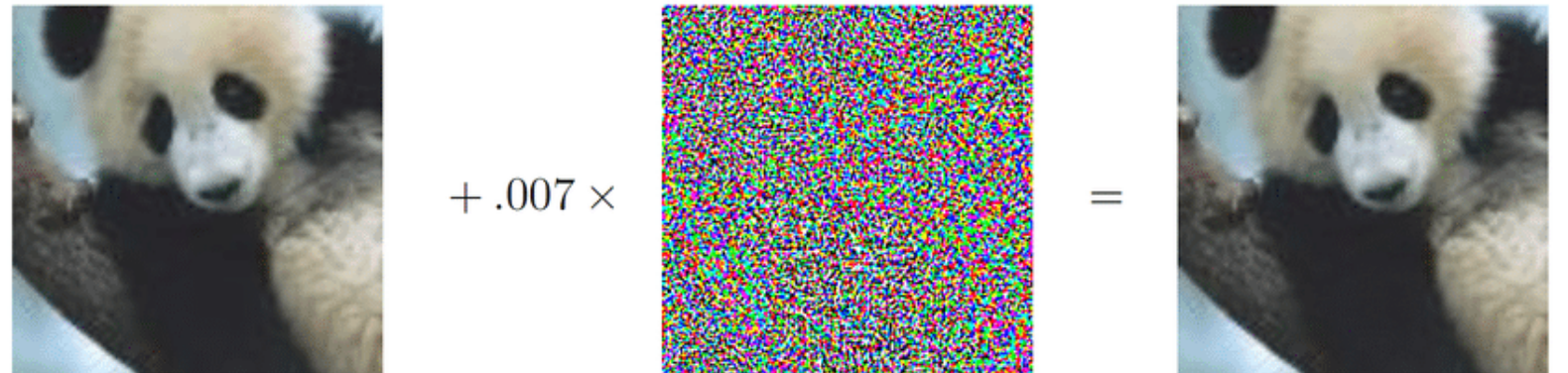
IJCAI-ECAI 2022

A Unified Framework for Adversarial Attack and Defense in Constrained Feature Space

Thibault Simonetto, Salijona Dyrnishi, Salah Ghamizi, Maxime Cordy, Yves Le Traon

University of Luxembourg

Evaluating the robustness with adversarial example

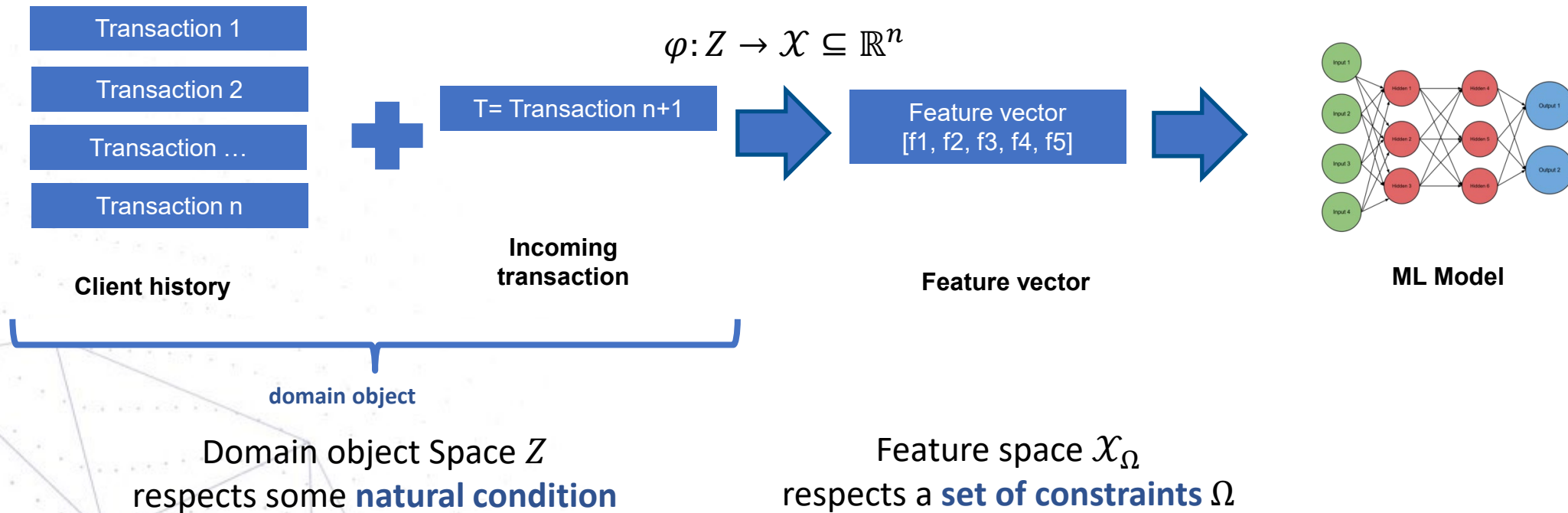


$$\begin{array}{ccc}
 \text{Image } x & + .007 \times & \text{Image } \text{sign}(\nabla_x J(\theta, x, y)) & = & \text{Image } x + \epsilon \text{sign}(\nabla_x J(\theta, x, y)) \\
 \text{"panda"} & & \text{"nematode"} & & \text{"gibbon"} \\
 57.7\% \text{ confidence} & & 8.2\% \text{ confidence} &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

J. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and harnessing adversarial examples," 2014

Motivation

ML model is integrated in a larger software system that takes as **input domain objects**.



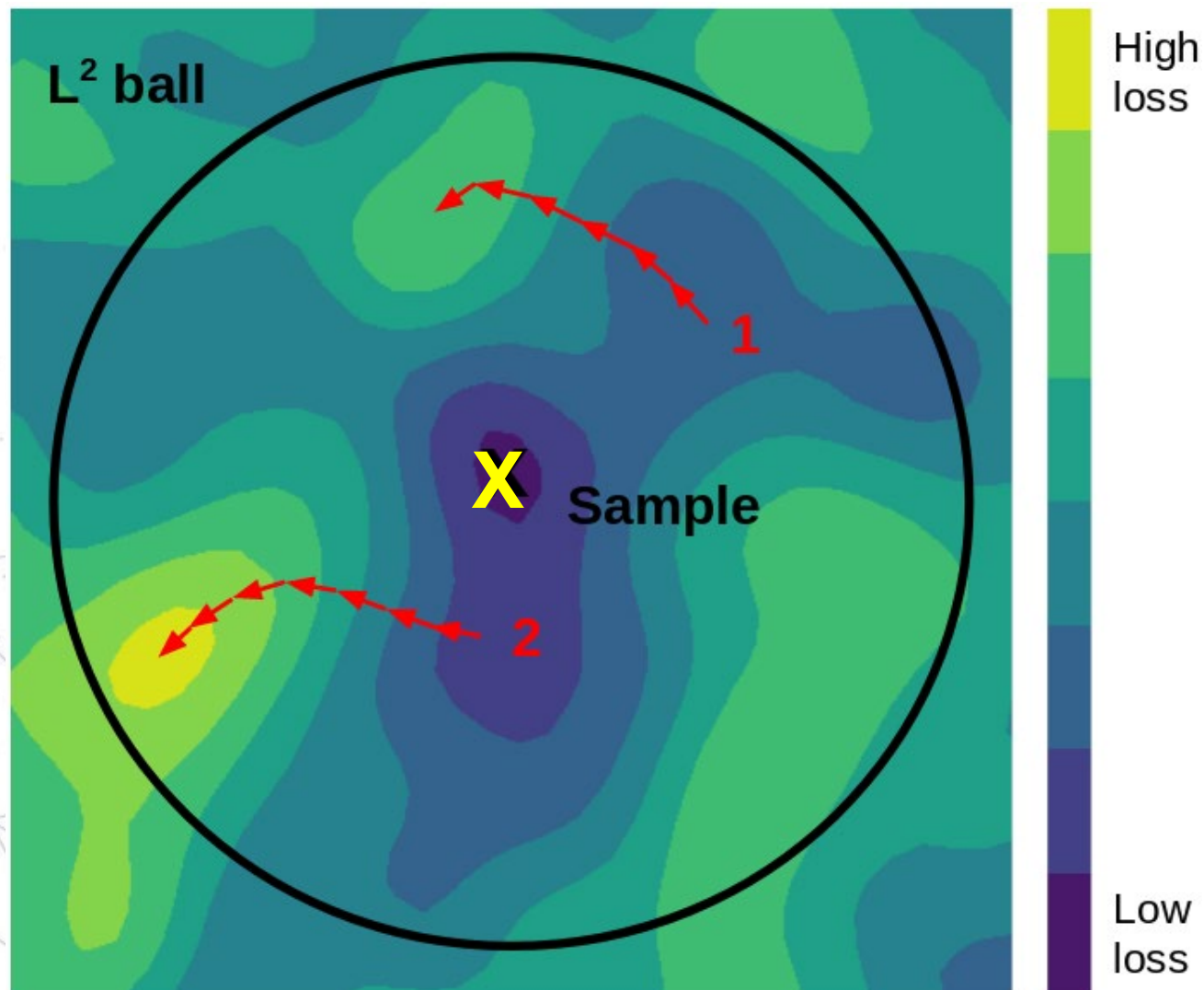
Constraints on feature space

Related features linked by a constraint:

$$open_acc \leq total_acc$$

$$installment = loan_amount \times \frac{int_rate \times (1 + int_rate)^{term}}{(1 + int_rate)^{term} - 1}$$

Existing attacks fails to generate constrained examples

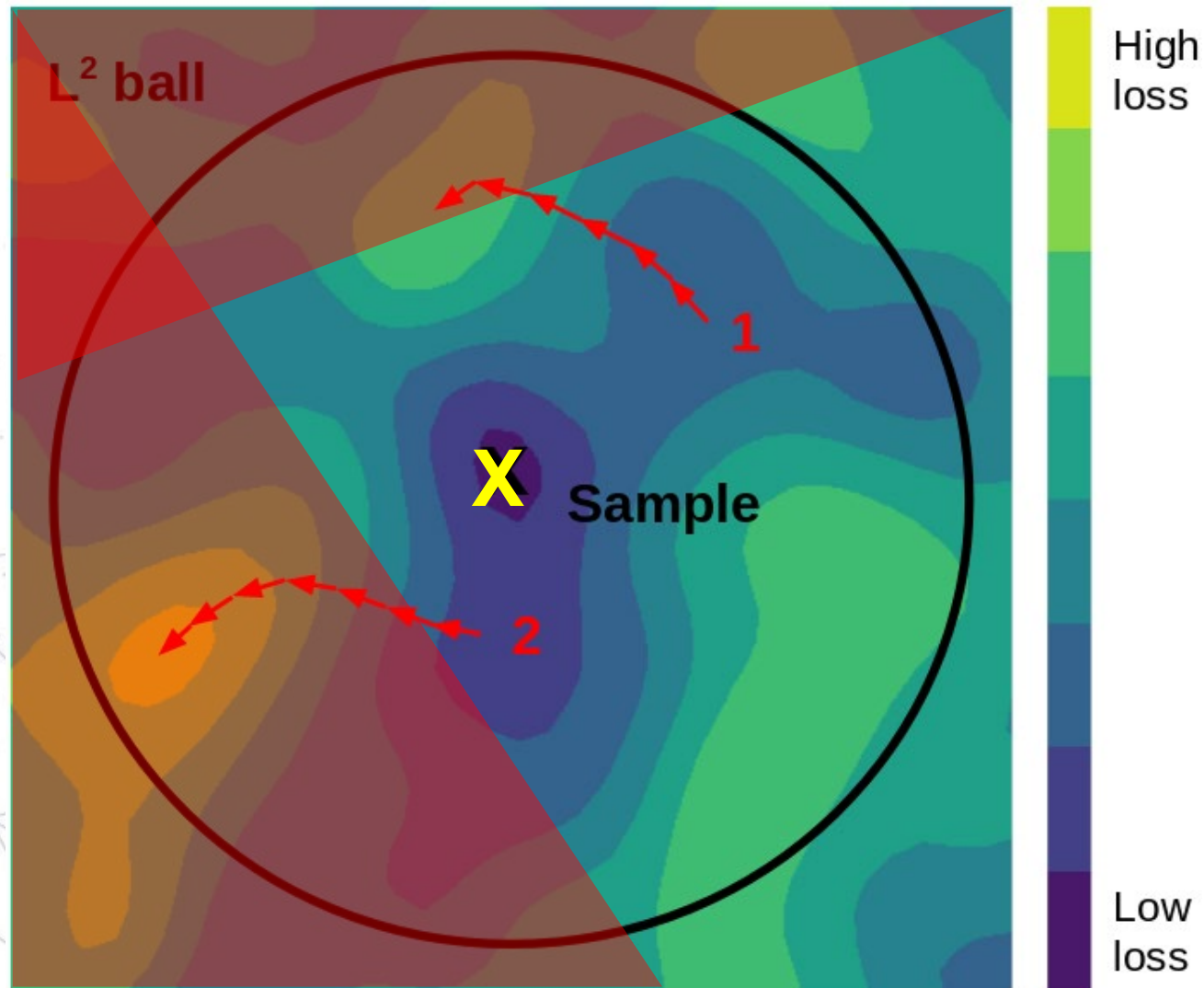


Objective: for x find δ

✓ With $H(x) \neq H(x + \delta)$

✓ With $L_p(x, x + \delta) < \epsilon$

Existing attacks fails to generate constrained examples

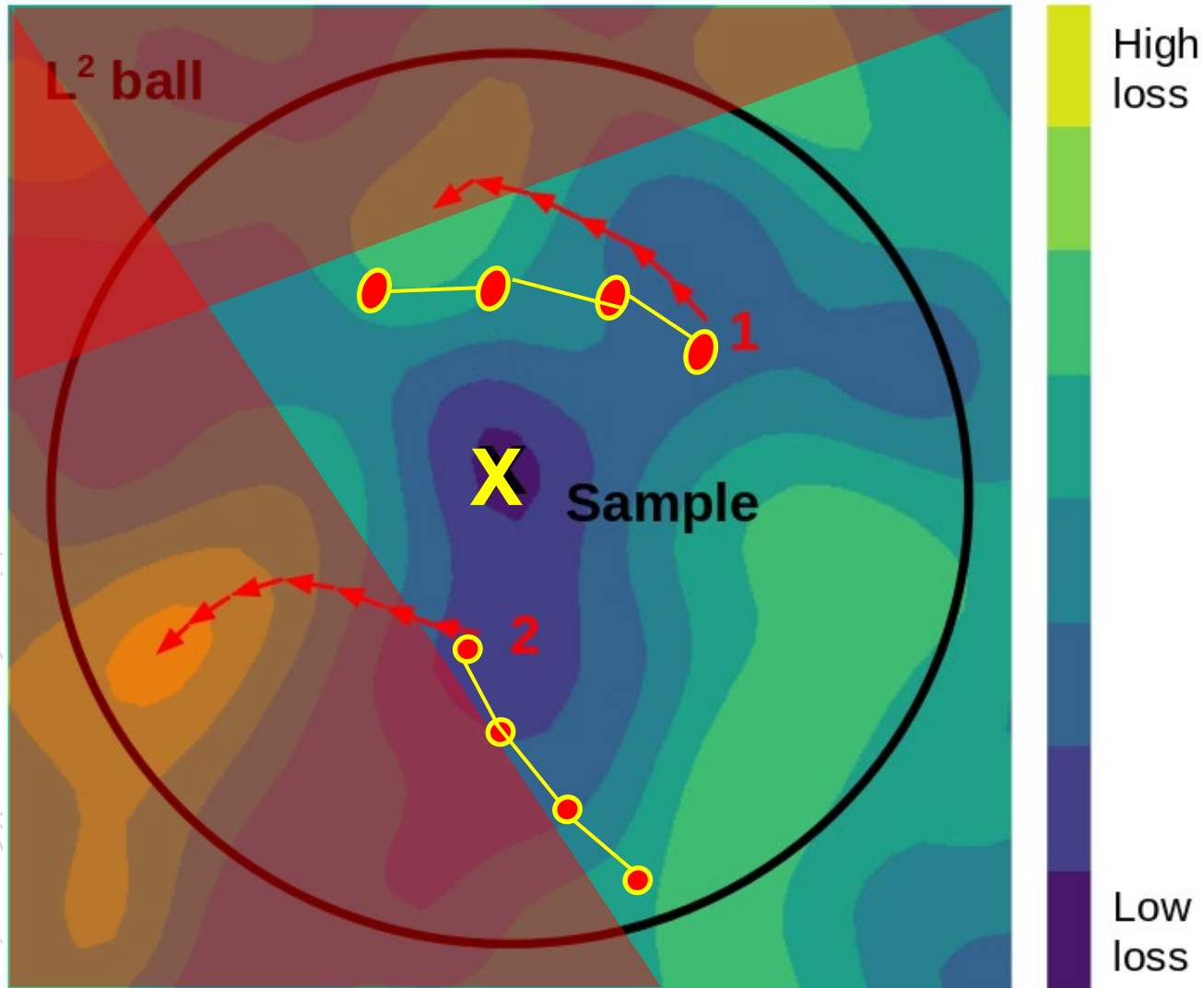


Objective: for x find δ

✓ With $H(x) \neq H(x + \delta)$

✓ With $L_p(x, x + \delta) < \epsilon$

Existing attacks fails to generate constrained examples



Objective: for x find δ

✓ With $H(x) \neq H(x + \delta)$

✓ With $L_p(x, x + \delta) < \epsilon$

✓ $x + \delta \in \mathcal{X}_\Omega$

Contributions

First **generic** framework for adversarial attacks under **domain constraints**

$$\omega := \omega_1 \wedge \omega_2 \mid \omega_1 \vee \omega_2 \mid \psi_1 \succeq \psi_2 \mid f \in \{\psi_1 \dots \psi_k\}$$

$$\psi := c \mid f \mid \psi_1 \oplus \psi_2 \mid x_i$$

Constraints formulae

Penalty function

$$\omega_1 \wedge \omega_2$$

$$\omega_1 + \omega_2$$

$$\omega_1 \vee \omega_2$$

$$\min(\omega_1, \omega_2)$$

$$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$$

$$\min(\{\psi_i \in \Psi : |\psi - \psi_i|\})$$

$$\psi_1 \leq \psi_2$$

$$\max(0, \psi_1 - \psi_2)$$

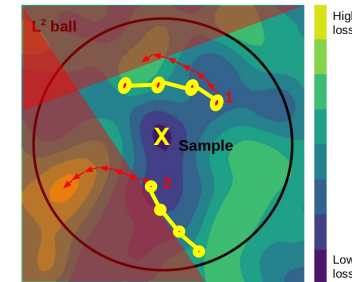
$$\psi_1 < \psi_2$$

$$\max(0, \psi_1 - \psi_2 + \tau)$$

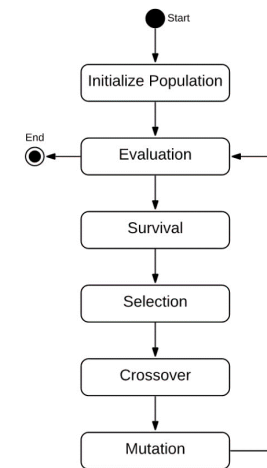
$$\psi_1 = \psi_2$$

$$|\psi_1 - \psi_2|$$

Generic constraints language



C-PGD



MoEvA2

2 Attacks

Constraints as penalty function

Constraint grammar

$$\omega := \omega_1 \wedge \omega_2 \mid \omega_1 \vee \omega_2 \mid \psi_1 \succeq \psi_2 \mid f \in \{\psi_1 \dots \psi_k\}$$

$$\psi := c \mid f \mid \psi_1 \oplus \psi_2 \mid x_i$$

$f \in \mathbf{F}$ is the value of feature f for a given input x' ,
 c is a constant real value,
 $\omega, \omega_1, \omega_2$ are constraint formulae,
 $\succeq \in \{<, \leq, =, \neq, \geq, >\}$,
 $\psi, \psi_1, \dots, \psi_k$ are numeric expressions,
 $\oplus \in \{+, -, *, \backslash\}$, and
 x_i is the value of the i^{th} feature of the clean input x

Constraints as penalty function

Constraint grammar

$$\omega := \omega_1 \wedge \omega_2 \mid \omega_1 \vee \omega_2 \mid \psi_1 \succeq \psi_2 \mid f \in \{\psi_1 \dots \psi_k\}$$

$$\psi := c \mid f \mid \psi_1 \oplus \psi_2 \mid x_i$$

$f \in \mathbf{F}$ is the value of feature f for a given input x' ,
 c is a constant real value,
 $\omega, \omega_1, \omega_2$ are constraint formulae,
 $\succeq \in \{<, \leq, =, \neq, \geq, >\}$,
 $\psi, \psi_1, \dots, \psi_k$ are numeric expressions,
 $\oplus \in \{+, -, *, \backslash\}$, and
 x_i is the value of the i^{th} feature of the clean input x

Constraints as penalty function

Constraint grammar

$$\omega := \omega_1 \wedge \omega_2 \mid \omega_1 \vee \omega_2 \mid \psi_1 \succeq \psi_2 \mid f \in \{\psi_1 \dots \psi_k\}$$

$$\psi := c \mid f \mid \psi_1 \oplus \psi_2 \mid x_i$$

$f \in \mathbf{F}$ is the value of feature f for a given input x' ,
 c is a constant real value,
 $\omega, \omega_1, \omega_2$ are constraint formulae,
 $\succeq \in \{<, \leq, =, \neq, \geq, >\}$,
 $\psi, \psi_1, \dots, \psi_k$ are numeric expressions,
 $\oplus \in \{+, -, *, \backslash\}$, and
 x_i is the value of the i^{th} feature of the clean input x

Mapping to penalty functions

Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Table 1: From constraint formulae to penalty functions. τ is an infinitesimal value.

Constraint is satisfied if and only if $g(\omega, x) = 0$

Constraints as penalty function

Constraint grammar

$$\omega := \omega_1 \wedge \omega_2 \mid \omega_1 \vee \omega_2 \mid \psi_1 \succeq \psi_2 \mid f \in \{\psi_1 \dots \psi_k\}$$

$$\psi := c \mid f \mid \psi_1 \oplus \psi_2 \mid x_i$$

$f \in \mathbf{F}$ is the value of feature f for a given input x' ,
 c is a constant real value,
 $\omega, \omega_1, \omega_2$ are constraint formulae,
 $\succeq \in \{<, \leq, =, \neq, \geq, >\}$,
 $\psi, \psi_1, \dots, \psi_k$ are numeric expressions,
 $\oplus \in \{+, -, *, \backslash\}$, and
 x_i is the value of the i^{th} feature of the clean input x

Mapping to penalty functions

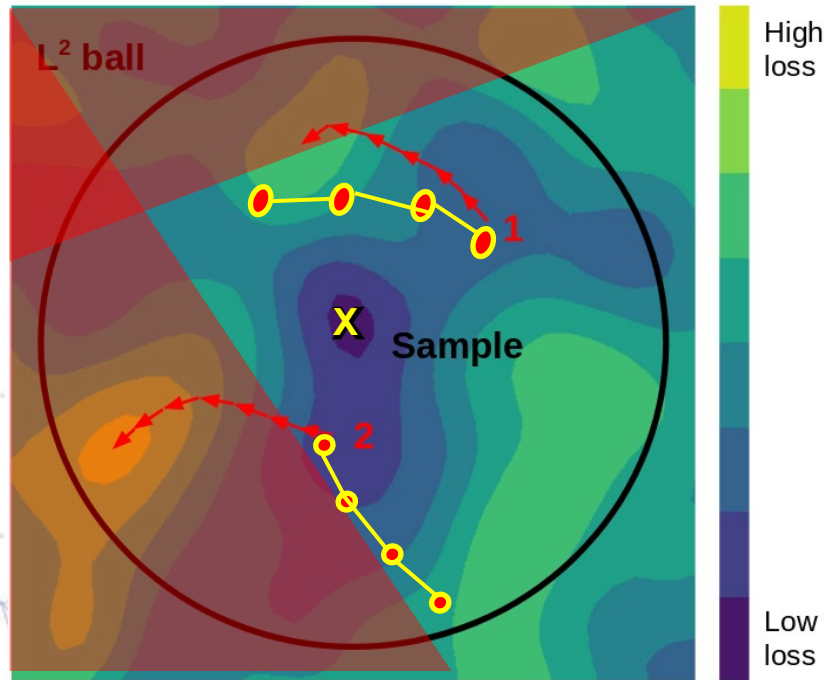
Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Table 1: From constraint formulae to penalty functions. τ is an infinitesimal value.

Constraint is satisfied if and only if $g(\omega, x) = 0$

Sufficient expressiveness to instantiate constraints in different domains

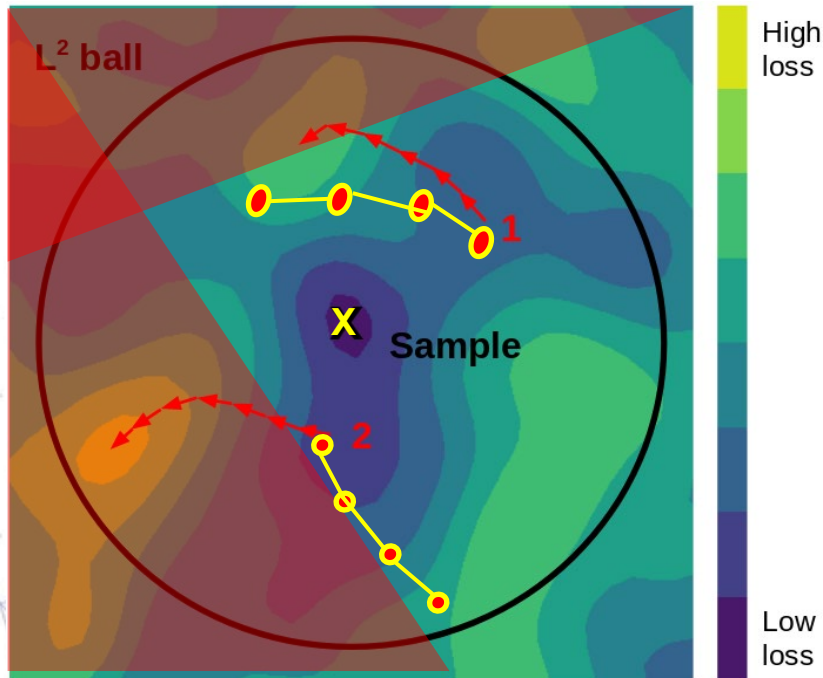
Approach 1: C-PGD: Gradient evaluation of the constraints



Projected Gradient Descent (PGD)

$$x^{t+1} = x^t + \alpha \operatorname{sign}(\nabla_{x^t} \operatorname{loss}(h(x^t), y))$$

Approach 1: C-PGD: Gradient evaluation of the constraints



Projected Gradient Descent (PGD)

$$x^{t+1} = x^t + \alpha \operatorname{sign}(\nabla_{x^t} \operatorname{loss}(h(x^t), y))$$

Constraints regularization

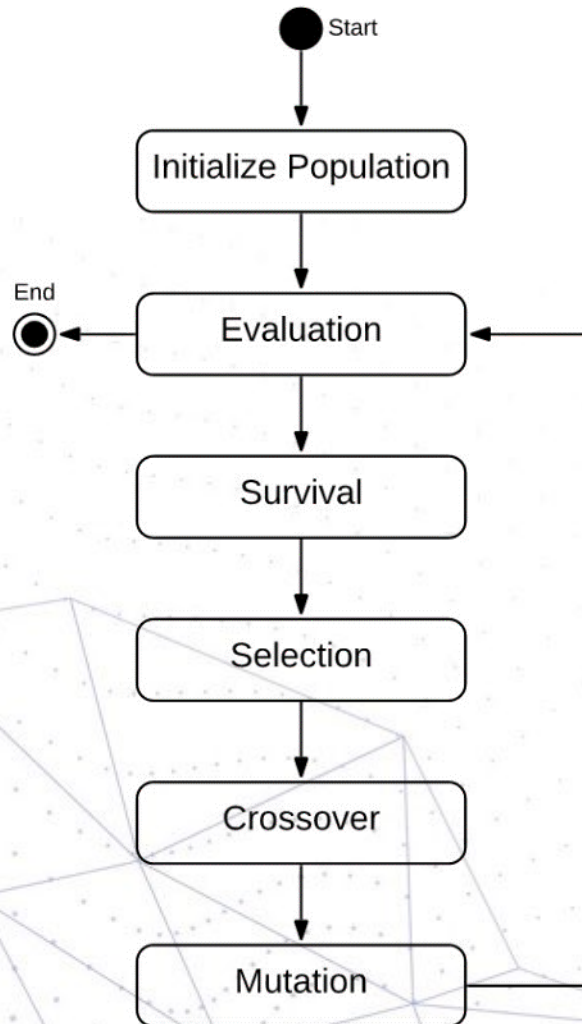
$$\nabla_{x^t} \operatorname{loss}(h(x^t), y) - \nabla_{x^t} \operatorname{penalty}(x^t)$$

Constrained Projected Gradient Descent (C-PGD)

$$x^{t+1} = x^t + \alpha \operatorname{sign}(\nabla_{x^t} \operatorname{loss}(h(x^t), y) - \nabla_{x^t} \operatorname{penalty}(x^t))$$

Approach 2: MoEvA2: Evolutionary approach

Multi-objective genetic algorithm (NSGA-III)

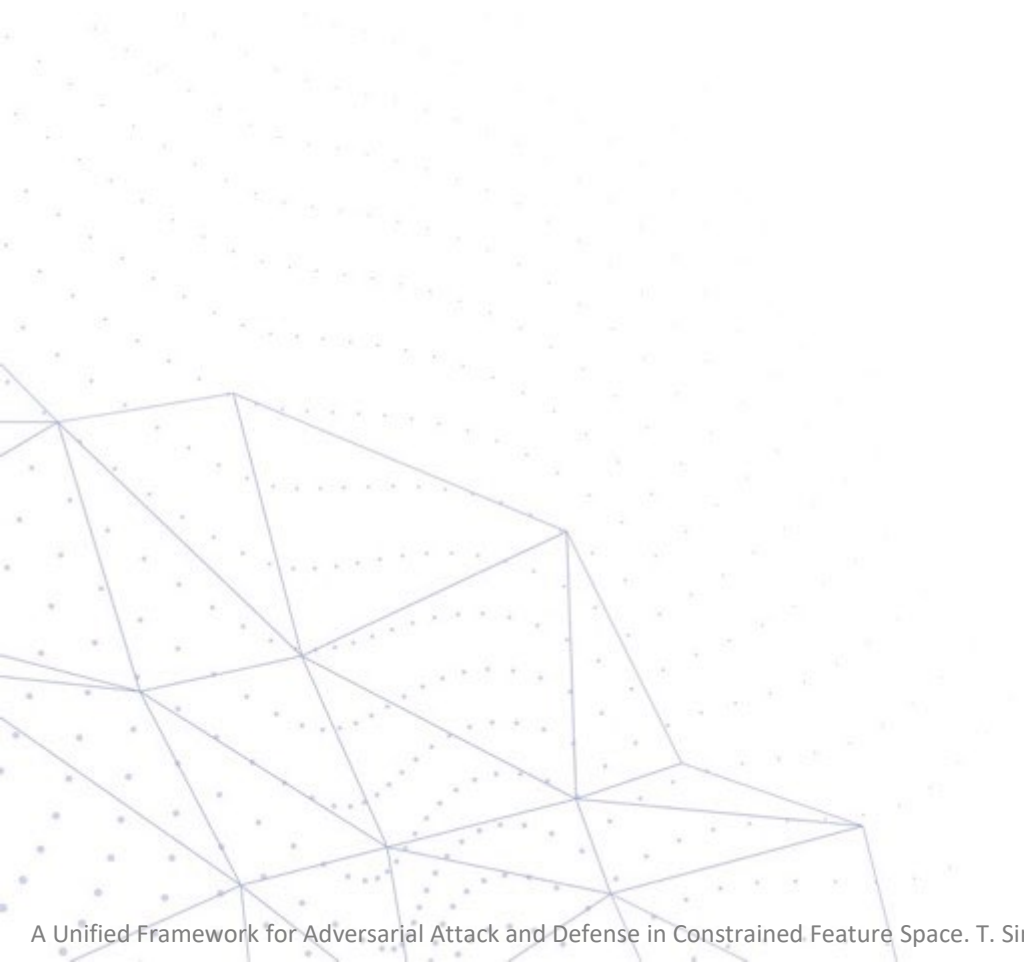


$$\text{minimise } g_1(x) \equiv h(x)$$

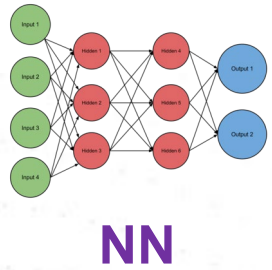
$$\text{minimise } g_2(x) \equiv L_p(x - x_0)$$

$$\text{minimise } g_3(x) \equiv \sum_{\omega_i \in \Omega} \text{penalty}(x, \omega_i)$$

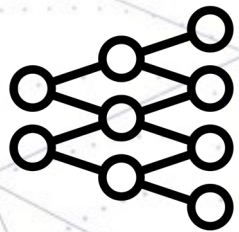
How effective are our approaches at generating adversarial examples?



How effective are our approaches at generating adversarial examples?



NN

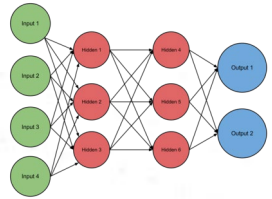


Random Forest

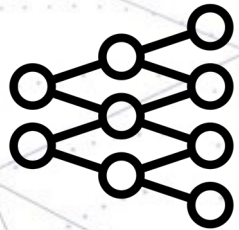
Dataset	Attack
	PGD
	PGD + SAT
	C-PGD
	MoEvA2
	PGD
	PGD + SAT
	C-PGD
	MoEvA2
	Papernot *
	MoEvA2
	Papernot *
	MoEvA2
	Papernot *
	MoEvA2
	Papernot *
	MoEvA2

* Extended to random forest

How effective are our approaches at generating adversarial examples?



NN



Random Forest

Dataset	Attack	C&M
	PGD	0.00
	PGD + SAT	0.00
	C-PGD	9.85
	MoEvA2	97.48
	PGD	0.00
	PGD + SAT	0.00
	C-PGD	0.00
	MoEvA2	100.00
	Papernot *	0.00
	MoEvA2	41.51
	Papernot *	0.0
	MoEvA2	5.41
	Papernot *	0.00
	MoEvA2	39.30
	Papernot *	8.50
	MoEvA2	31.89

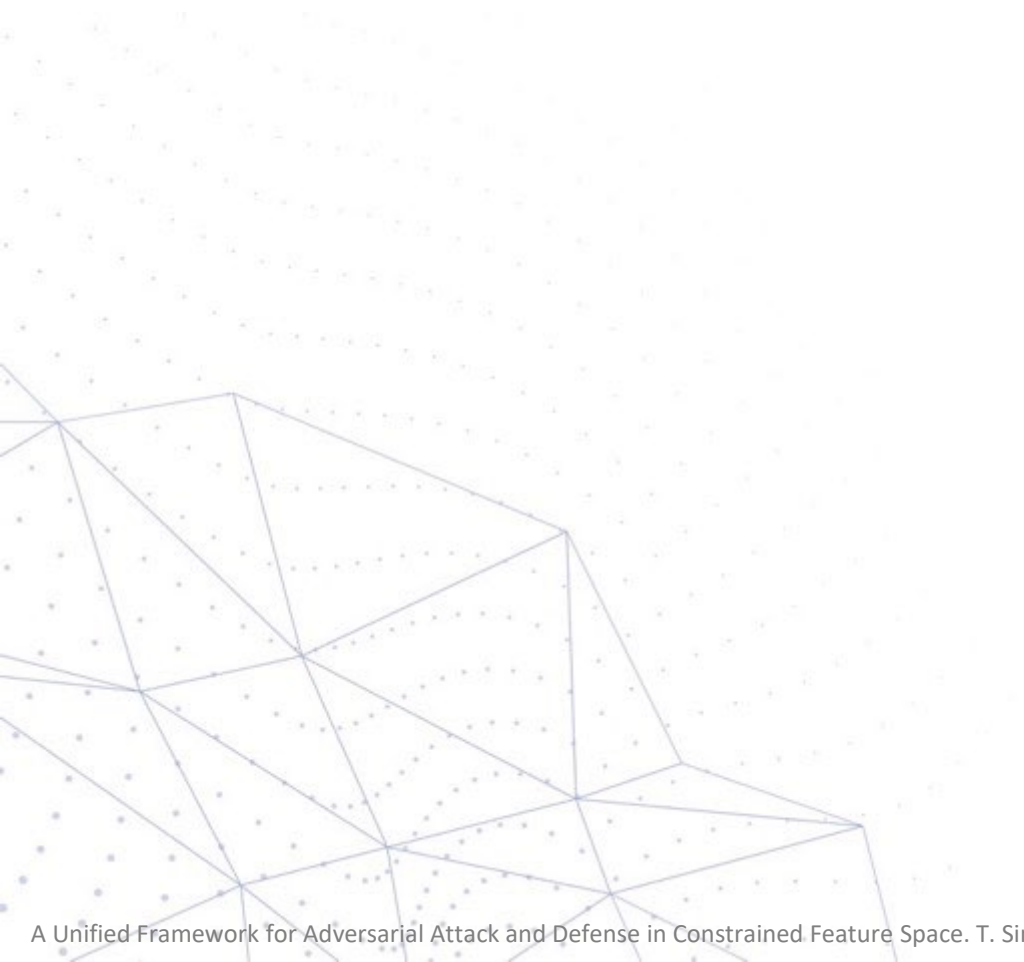
Attacks unaware of domain constraints often fail.

C-PGD is effective on a **single** dataset.

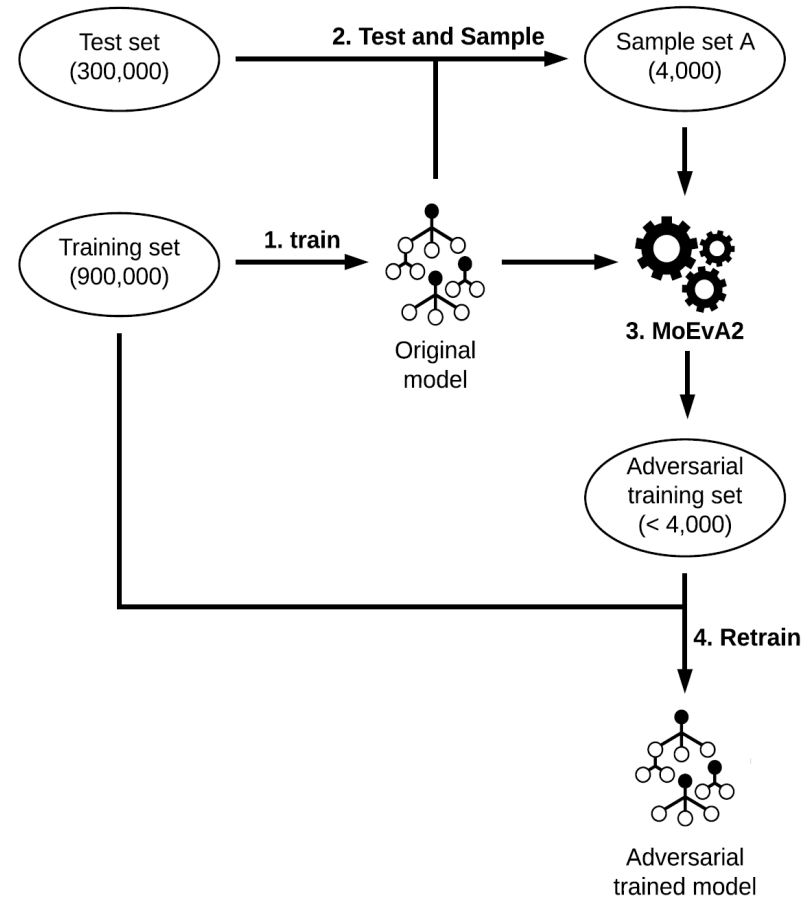
MoEvA2 is always **successful**.

* Extended to random forest

How can we improve the robustness of our machine learning models?



How can we improve the robustness of our machine learning models?



Adversarial retraining

How can we improve the robustness of our machine learning models?

We hypothesize that **augmenting Ω** with a set of engineered constraints can **robustify a model**.

How can we improve the robustness of our machine learning models?

We hypothesize that **augmenting** Ω with a set of engineered constraints can **robustify a model**.

We engineer a **new feature**

$$f_e = f_1 \oplus f_2$$

How can we improve the robustness of our machine learning models?

We hypothesize that **augmenting** Ω with a set of engineered constraints can **robustify a model**.

We engineer a **new feature**

$$f_e = f_1 \oplus f_2$$

We have the **new constraint**

$$\omega_e \models (f_e = f_1 \oplus f_2)$$

How effective are defense techniques ?



Defense	Attack	LCLD	CTU-13
None	C-PGD	9.85	0.00
None	MoEvA2	97.48	100.00
C-PGD Adv. retraining *	C-PGD	8.78	NA
C-PGD Adv. retraining *	MoEvA2	94.90	NA
MoEvA2 Adv. retraining *	C-PGD	2.70	NA
MoEvA2 Adv. retraining *	MoEvA2	85.20	0.8
Constraints augment.	C-PGD	0.00	NA
Constraints augment.	MoEvA2	80.43	0.00
MoEvA2 Adv. retrain. †	MoEvA2	82.00	NA
Combined defenses †	MoEvA2	77.43	NA

Table 3: Success rate of C-PGD and MoEvA2 after adversarial retraining and constraint augmentation (on neural networks). For a fair comparison, the model denoted by the same symbols (*) or †) are trained with the same number of adversarial examples, generated from the same original samples.

Constraint augmentation is an effective alternative defense to adversarial retraining.

Constraint augmentation and adversarial retraining have complementary effects.

Conclusion

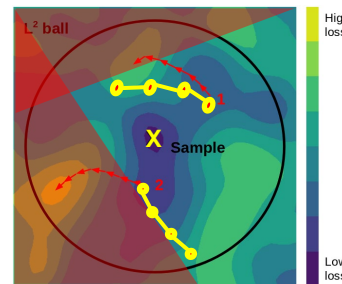
Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Generic constraints language

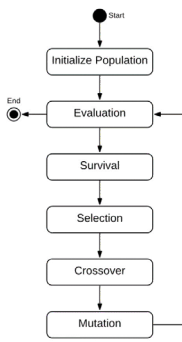
Conclusion

Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Generic constraints language



C-PGD



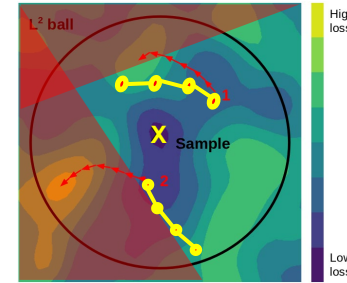
MoEvA2

2 Constrained Attacks

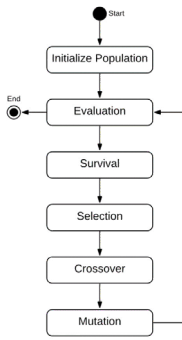
Conclusion

Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Generic constraints language



C-PGD



MoEvA2

2 Constrained Attacks



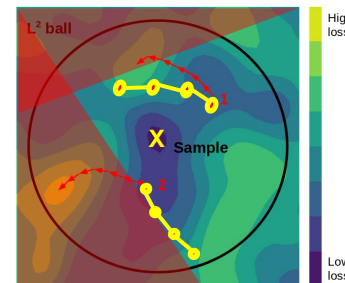
Attacks unaware of domain constraints often fail.

MoEvA2 is always successful.

Conclusion

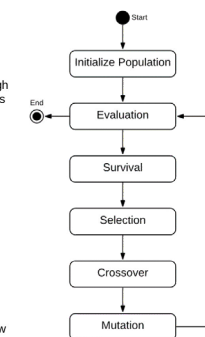
Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi : \psi - \psi_i \})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2 + \tau)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Generic constraints language



C-PGD

2 Constrained Attacks



MoEvA2



Attacks unaware of domain constraints often fail.

MoEvA2 is always successful.

$$\omega_e \models (f_e = f_1 \oplus f_2)$$

Constrained Augmentation

New defense method as effective as adversarial retraining

Conclusion

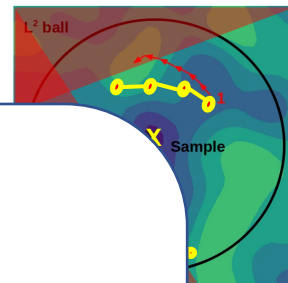
Constraints formulae	Penalty function
$\omega_1 \wedge \omega_2$	$\omega_1 + \omega_2$
$\omega_1 \vee \omega_2$	$\min(\omega_1, \omega_2)$
$\psi \in \Psi = \{\psi_1, \dots, \psi_k\}$	$\min(\{\psi_i \in \Psi\})$
$\psi_1 \leq \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 < \psi_2$	$\max(0, \psi_1 - \psi_2)$
$\psi_1 = \psi_2$	$ \psi_1 - \psi_2 $

Generic constraints language

Checkout our
framework

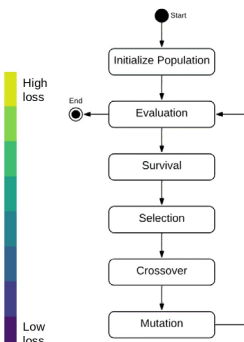


Stand 149 Row 5



PGD

Constrained Attacks



MoEvA2



Attacks unaware of defense
constraints often fail

MoEvA2 is always
successful.

$$\models (f_e = f_1 \oplus f_2)$$

trained Augmentation

New defense method as effective
as adversarial retraining