

SECURITY WEAKNESSES IN MACHINE LEARNING

DANIEL ETZOLD - @ETZOLDIO

MONDAY 25TH MARCH

NAIVE BAYES FILTER BYPASSED

Naive Bayes is a simple probabilistic classifier based on applying Bayes' theorem with strong independence assumptions between the features.



Spam filters use Naive Bayes to classify emails as spam or not spam based on the words in the email.

Spam filters use Naive Bayes to classify emails as spam or not spam based on the words in the email.

IMAGE CLASSIFIER FOOLED

Image classifiers are used to identify objects in images. They are trained on a large dataset of images and learn to recognize patterns.



Image classifiers are used to identify objects in images. They are trained on a large dataset of images and learn to recognize patterns.

THIEVES STEAL MODEL PARAMETERS

$$f(x) = Wx + b$$
$$= \sum_{i=1}^n w_i x_i + b$$

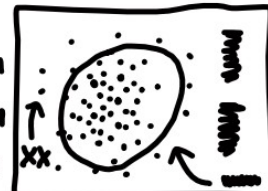
Model parameters are the weights and bias used in the model.

Model parameters are the weights and bias used in the model.

Model parameters are the weights and bias used in the model.

ANOMALY DETECTION HACKED

Anomaly detection is used to identify unusual patterns that do not conform to expected behavior.





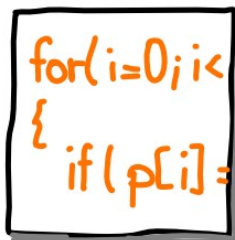
@ETZOLDIO

[HTTPS://ETZOLD.IO](https://etzold.io)

- IT SECURITY ARCHITECT



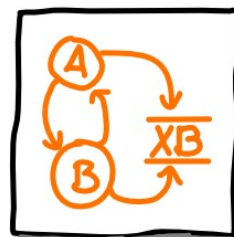
- 1&1** MAIL&MEDIA DEVELOPMENT & TECHNOLOGY GMBH



CODE
REVIEW



DOCUMENTATION
REVIEW



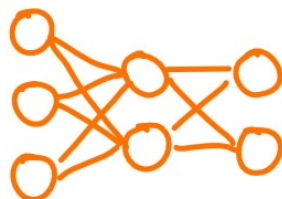
THREAT
MODELING



IN HOUSE
CONSULTING



PENETRATION
TESTING



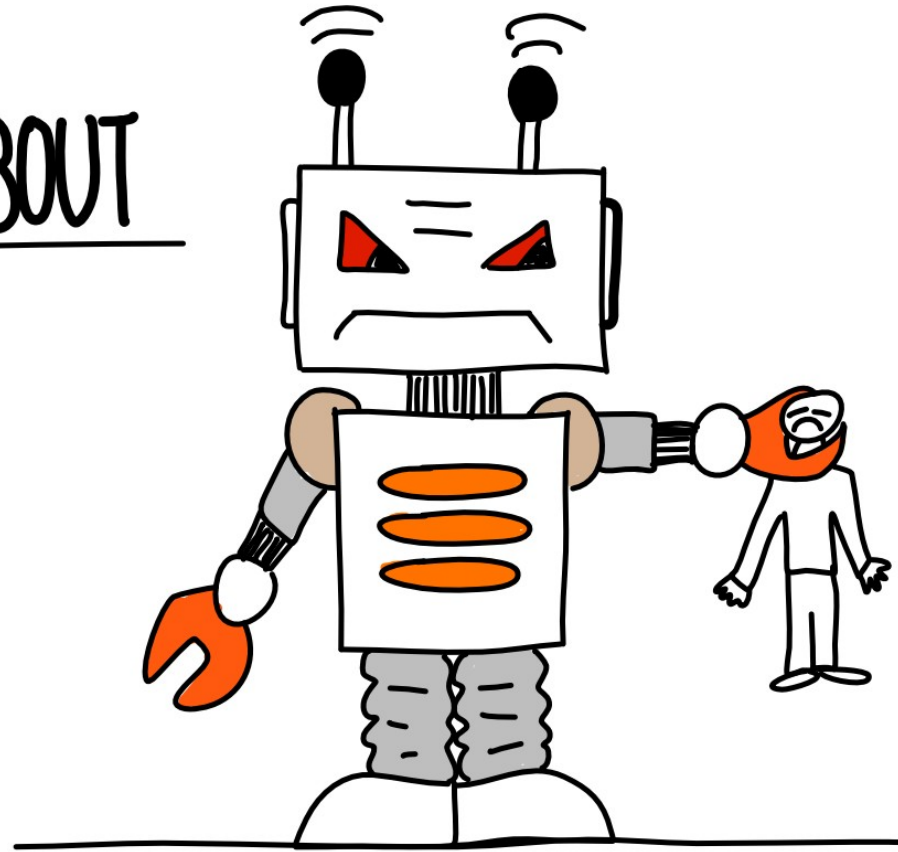
+

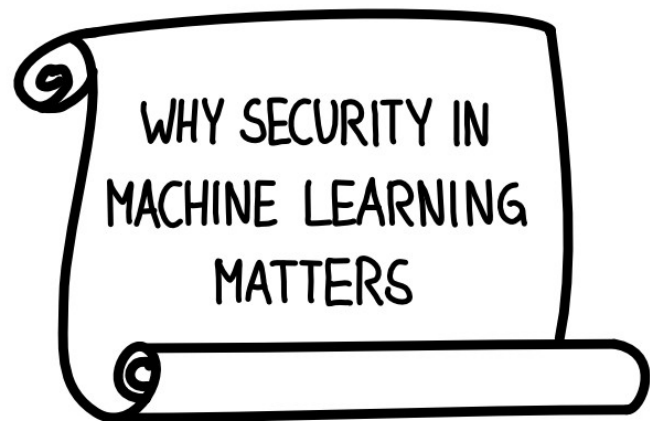


SECURITY IN
MACHINE LEARNING



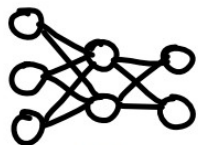
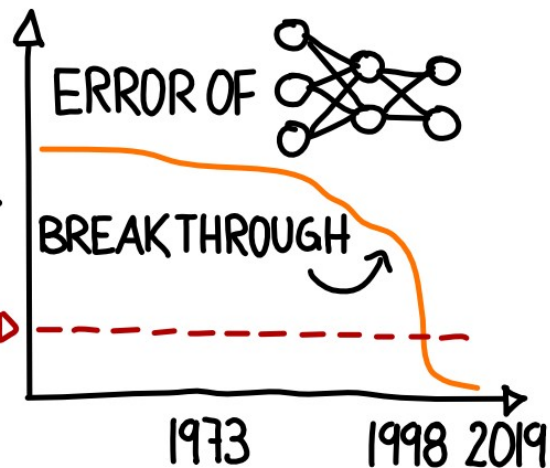
IT'S NOT ABOUT





REASON

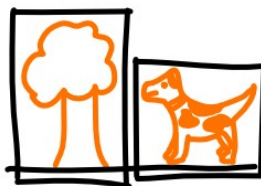
HUMAN-LEVEL
PERFORMANCE



_____ BREAKTHROUGH ON _____



MACHINE
TRANSLATION



OBJECT
DETECTION



SPEECH
RECOGNITION



RATHER
OLD

ATTACKS
ON
MACHINE
LEARNING



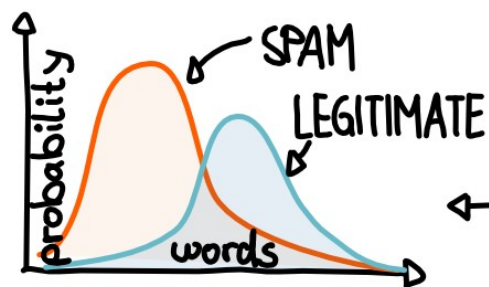
SOMETIMES
EASY



TRAINING:

- FIND USEFUL WORDS
- ESTIMATE SPAM PROBABILITIES
↳ $P(W_i=1 | \text{SPAM})$

OBSERVATION: 
WORD DISTRIBUTION



SPAM FILTERING

CLASSIFICATION:

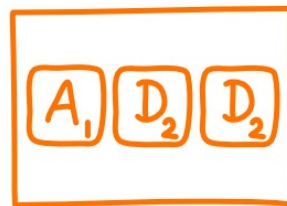
$$\prod_i P(W_i=1 | \text{SPAM})$$

← EXPLOIT →

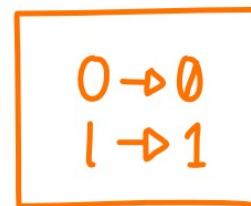
ATTACKER'S OPTIONS



~ 20 YEARS ←



GOOD WORDS

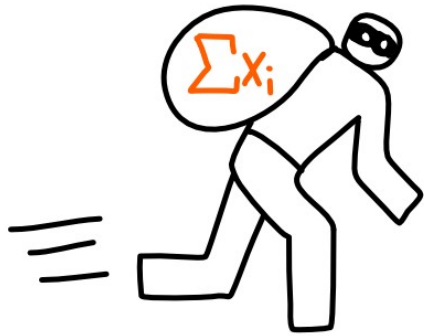


OBFUSCATE



IMAGES

LOGISTIC REGRESSION



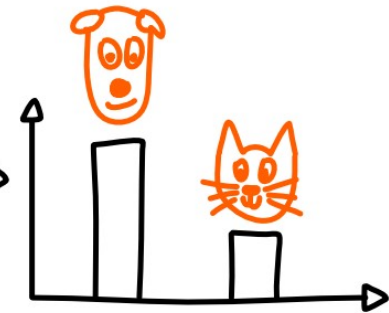
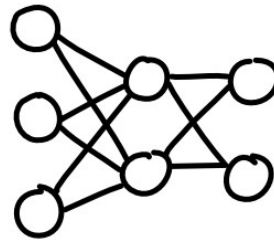
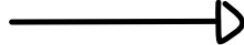
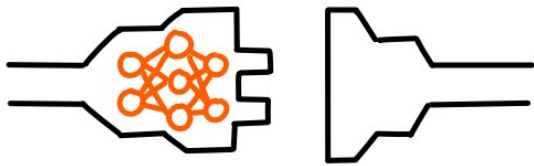
EASIER TO
ANALYZE

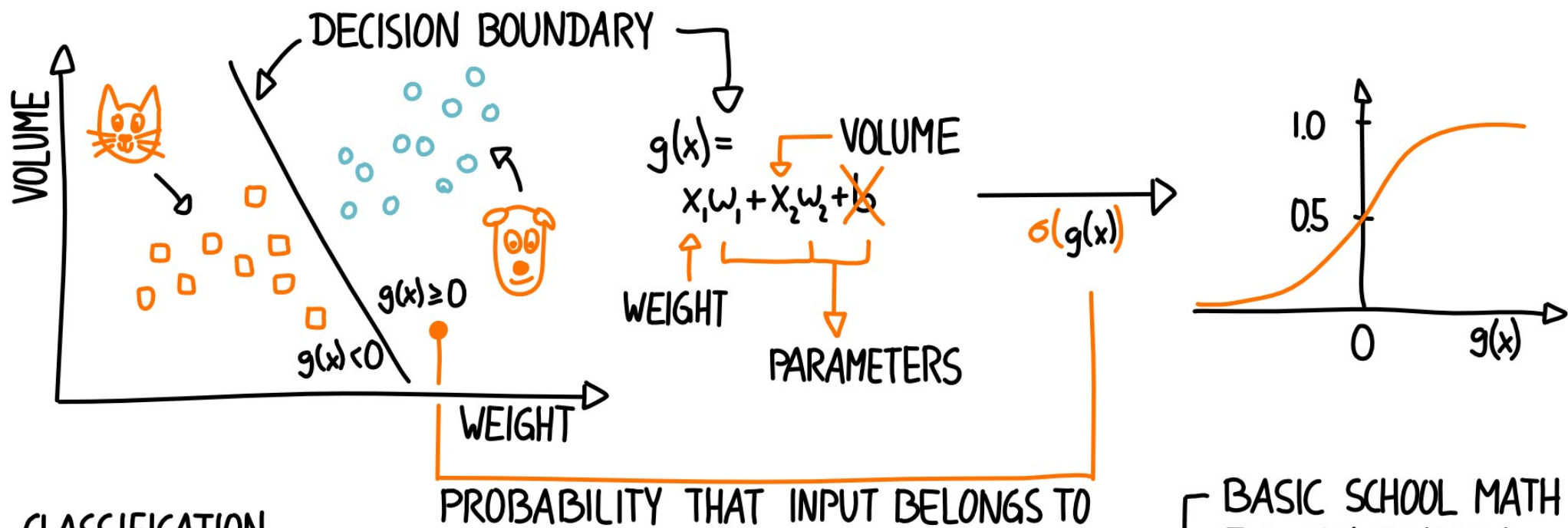


LEAKS
SECRETS



BUSINESS
AT RISK





CLASSIFICATION

INPUT: (0.3, 0.5) $\rightarrow \sigma(0.3w_1 + 0.5w_2) = 0.8$

INPUT: (0.7, 0.2) $\rightarrow \sigma(0.7w_1 + 0.2w_2) = 0.4$

INPUT: (0.5, 0.8) $\rightarrow$



$0.3w_1 + 0.5w_2 = \sigma^{-1}(0.8)$

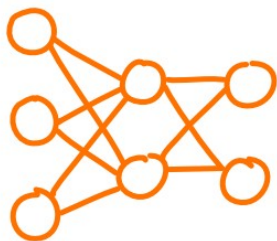
$0.7w_1 + 0.2w_2 = \sigma^{-1}(0.4)$

BASIC SCHOOL MATH
E.G. ELIMINATION

$w_1 = \dots$

$w_2 = \dots$



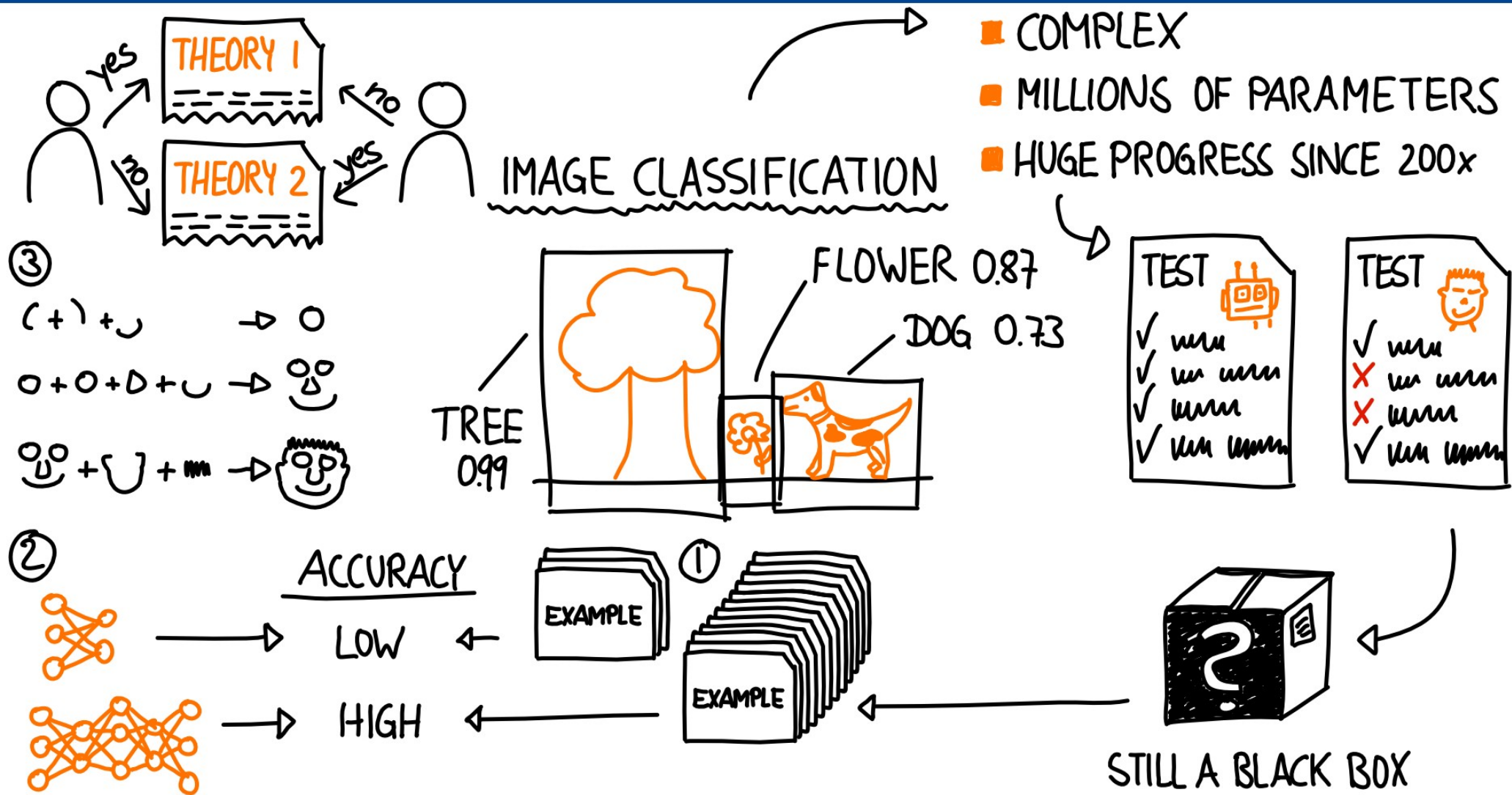


CAN BE EXTENDED TO
NEURAL NETWORKS



TRAINING DATA
CAN BE RECOVERED

Fredrikson, et al. „Model inversion attacks that exploit confidence information and basic countermeasures“, 2015



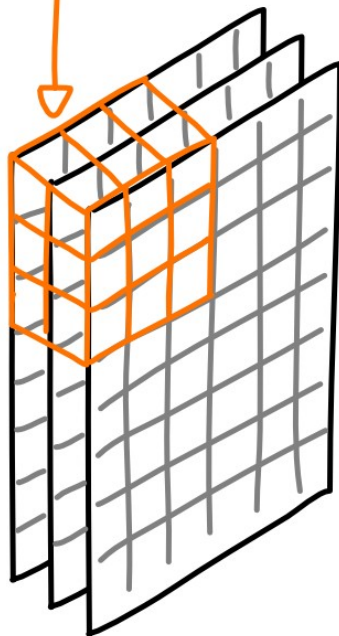
CONVOLUTIONAL NEURAL NETWORKS



FIRST LAYER



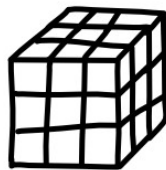
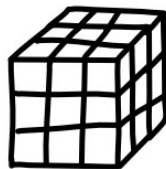
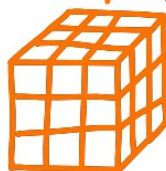
IMAGE



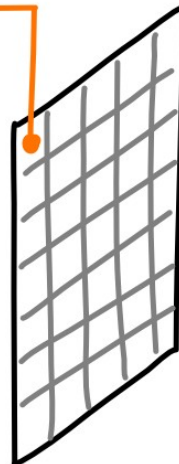
TENSOR



FILTER



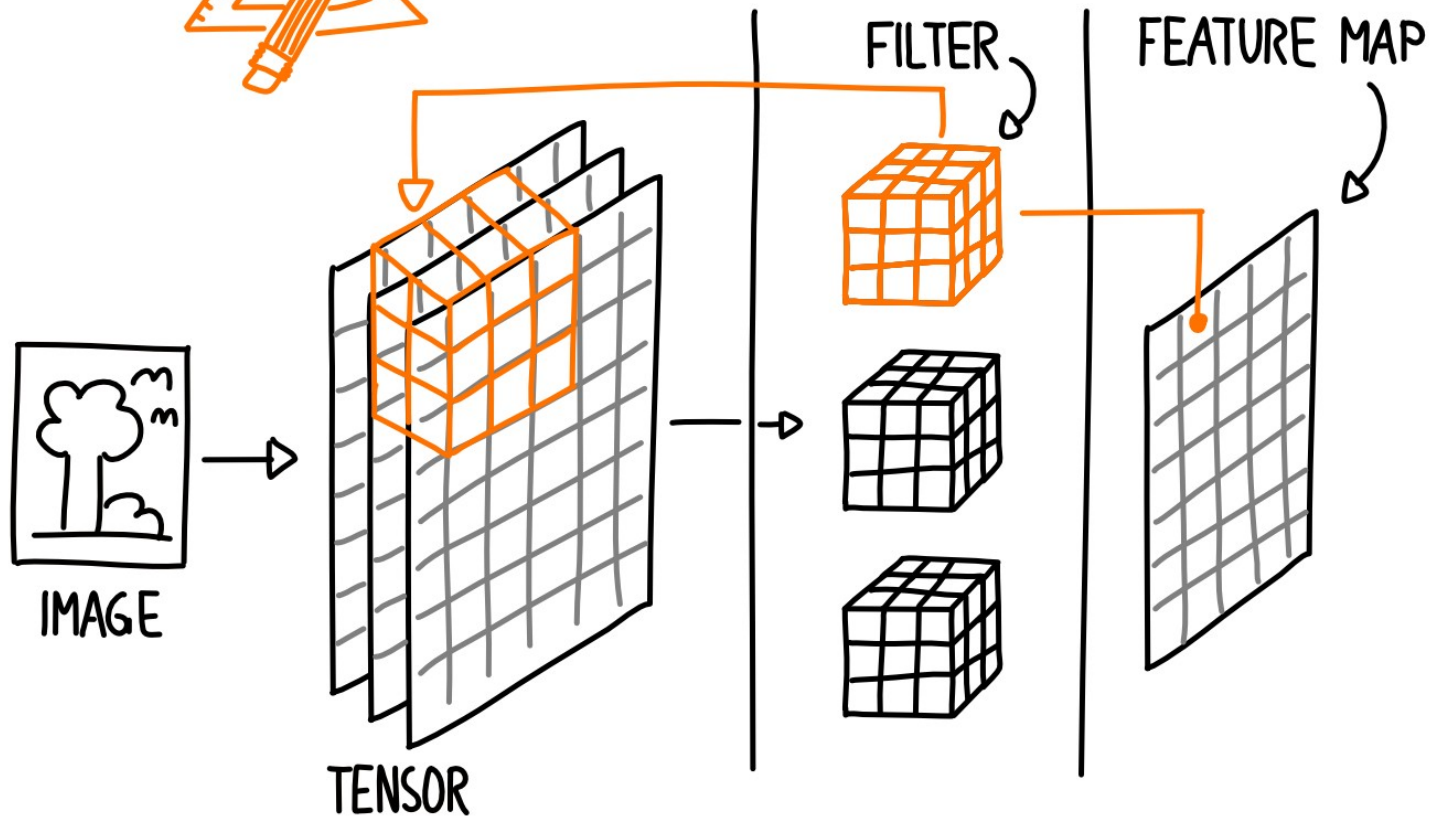
FEATURE MAP



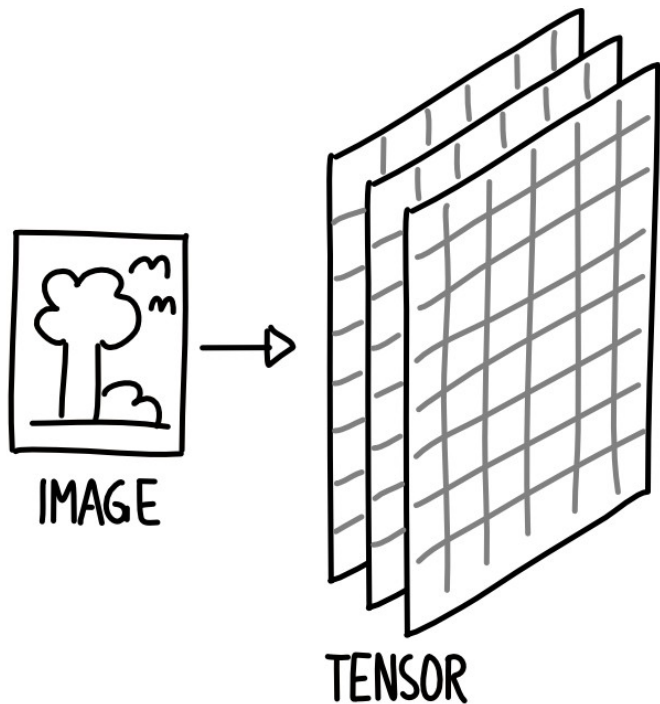
CONVOLUTIONAL NEURAL NETWORKS



FIRST LAYER

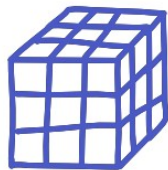
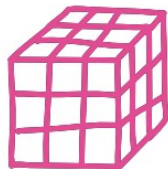
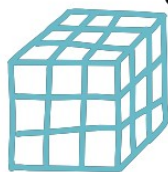


CONVOLUTIONAL NEURAL NETWORKS

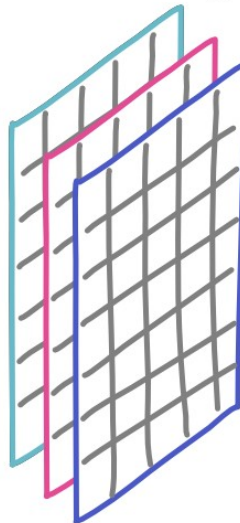


FIRST LAYER

FILTER



FEATURE MAP



NEXT LAYERS

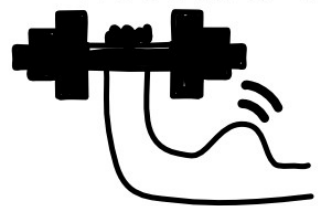
OUTPUT

$\begin{bmatrix} 0.2 \\ 0.3 \\ \vdots \\ 0.1 \end{bmatrix}$

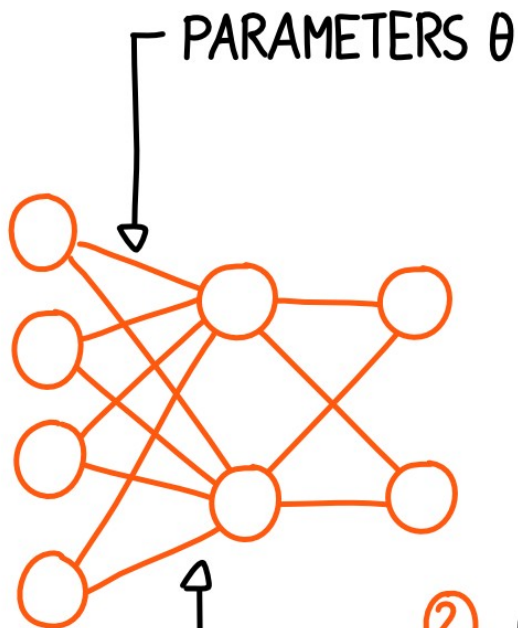
SIMPLE FEATURES

COMPLEX FEATURES

TRAINING OF NETWORKS

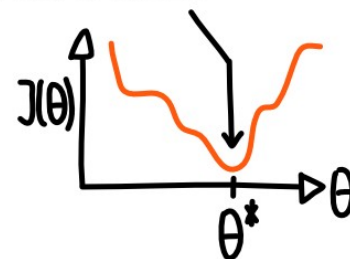


①



$J(\theta)$ = ERROR OF NETWORK

GOAL: FIND θ SUCH THAT $J(\theta)$ IS MINIMIZED



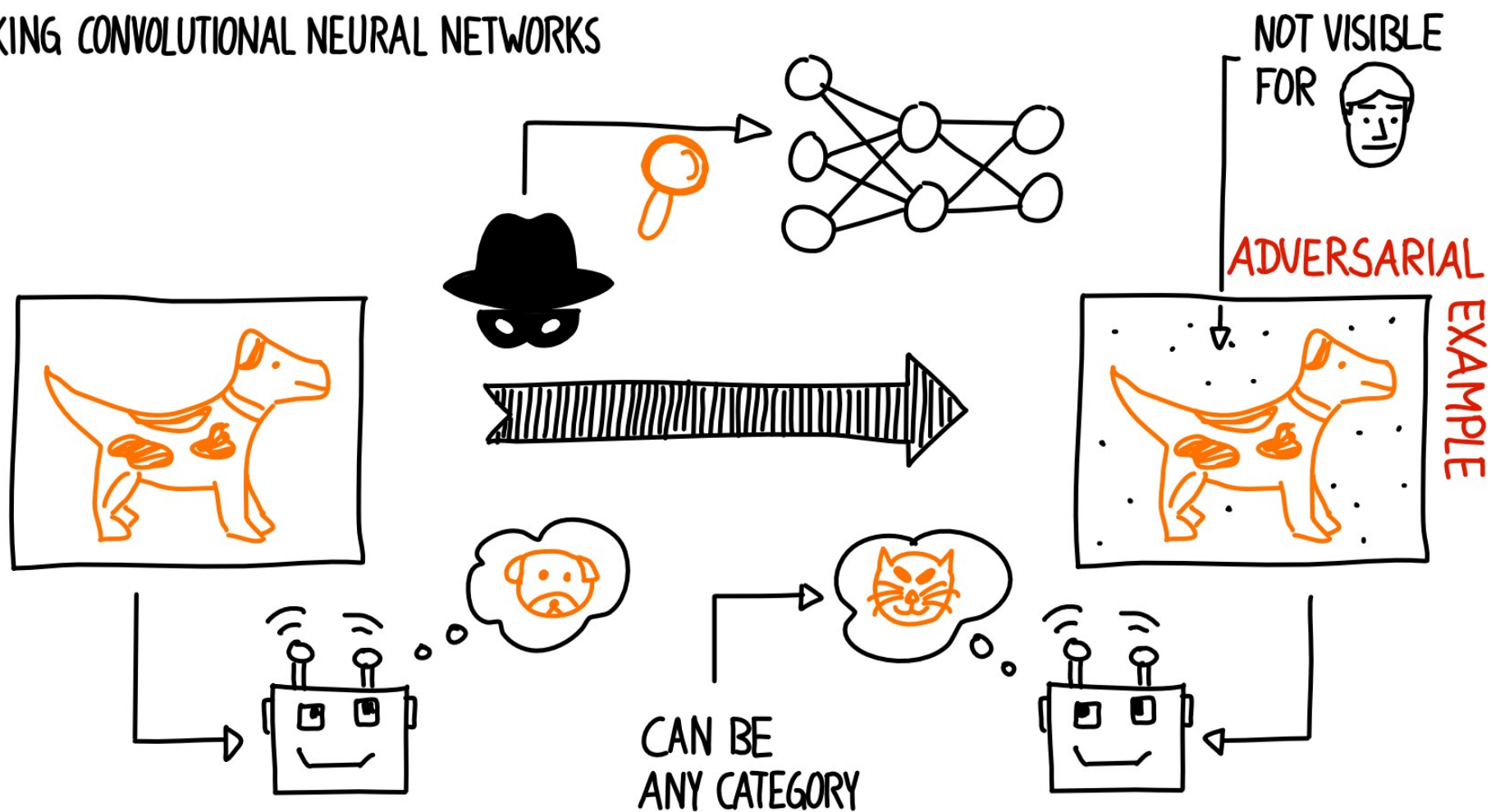
⑤ REPEAT
STOP IF
UPDATE $< \epsilon$

② COMPUTE $J(\theta)$

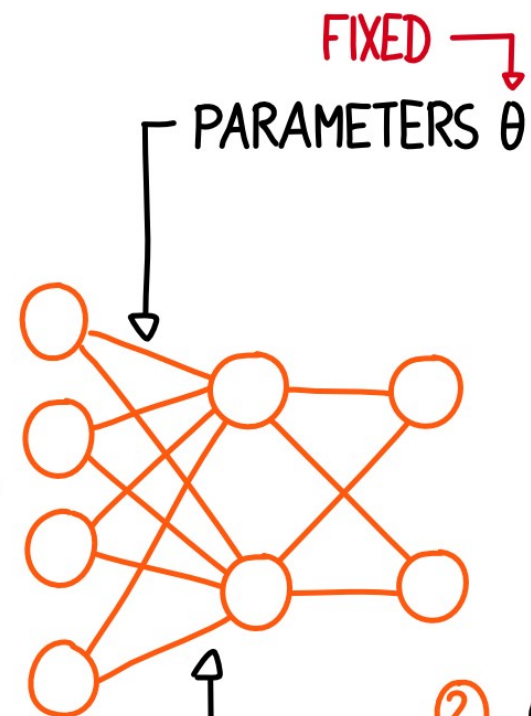
③ COMPUTE $\frac{\partial}{\partial \theta_i} J(\theta)$

④ UPDATE PARAMETERS

ATTACKING CONVOLUTIONAL NEURAL NETWORKS



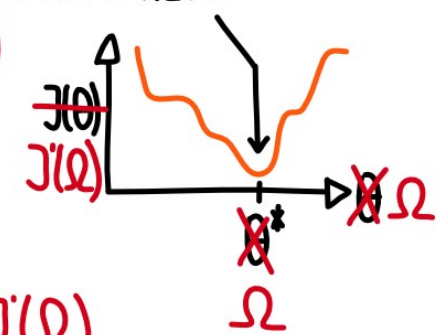
TRAINING OF NETWORKS



$J'(\Omega)$ = ERROR WITH RESPECT TO TARGET CATEGORY

~~$J(\theta)$ = ERROR OF NETWORK~~

GOAL: FIND Ω^* SUCH THAT $J'(\Omega)$ IS MINIMIZED



⑤ REPEAT
STOP IF
ERROR ~~UPDATE~~ $< \epsilon$

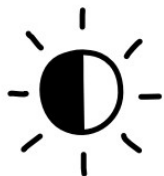
② COMPUTE ~~$J(\theta)$~~ $J'(\Omega)$

③ COMPUTE ~~$\frac{\partial}{\partial \theta_i} J(\theta)$~~ $\frac{\partial}{\partial \Omega_i} J'(\Omega)$

④ UPDATE PARAMETERS



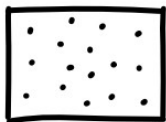
CONTRAST



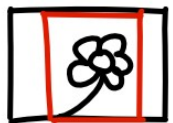
BRIGHTNESS



$\pm$



NOISE



DIFFERENT CROPS

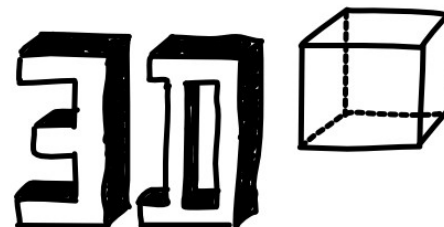


PRINT IT

ADVERSARIAL
EXAMPLES

CAN BE VERY
ROBUST

AND
THEY
CAN
BE



COOL !

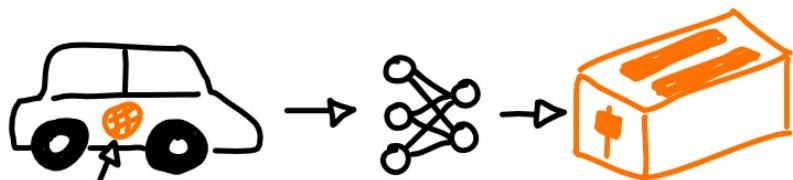
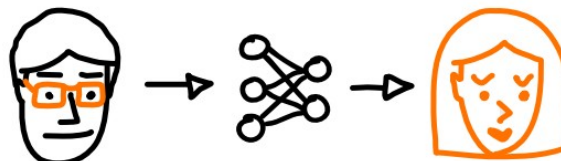
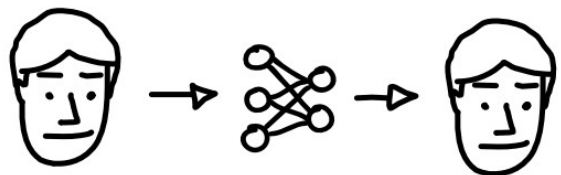
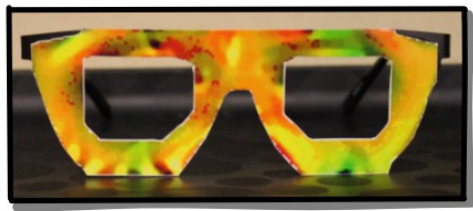


[HTTPS://YOUTU.BE/XaQu7kkQBPc](https://youtu.be/XaQu7kkQBPc)

Athalye, et al. „Synthesizing robust adversarial examples,, 2017

Sharif, et al. „Accessorize to a crime: Real and stealthy attacks on state-of-the-art face recognition,,, 2016

SPECIALLY CRAFTED GLASSES



PATCH

Brown, et al. „Adversarial patch“, 2017

OTHER
WAYS TO
HACK

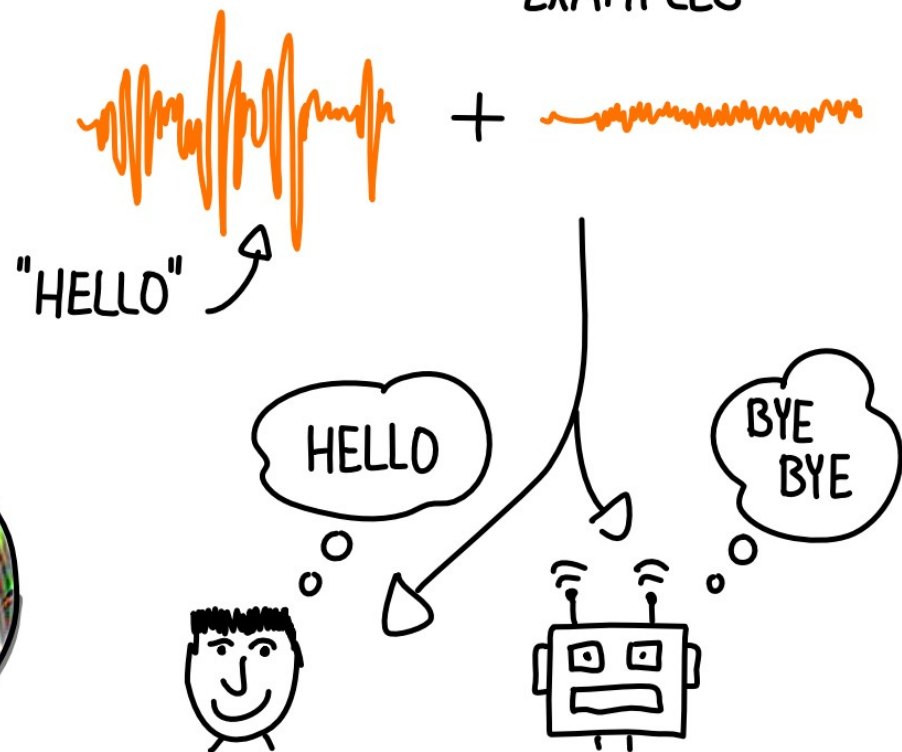
IMAGE
CLASSIFIERS

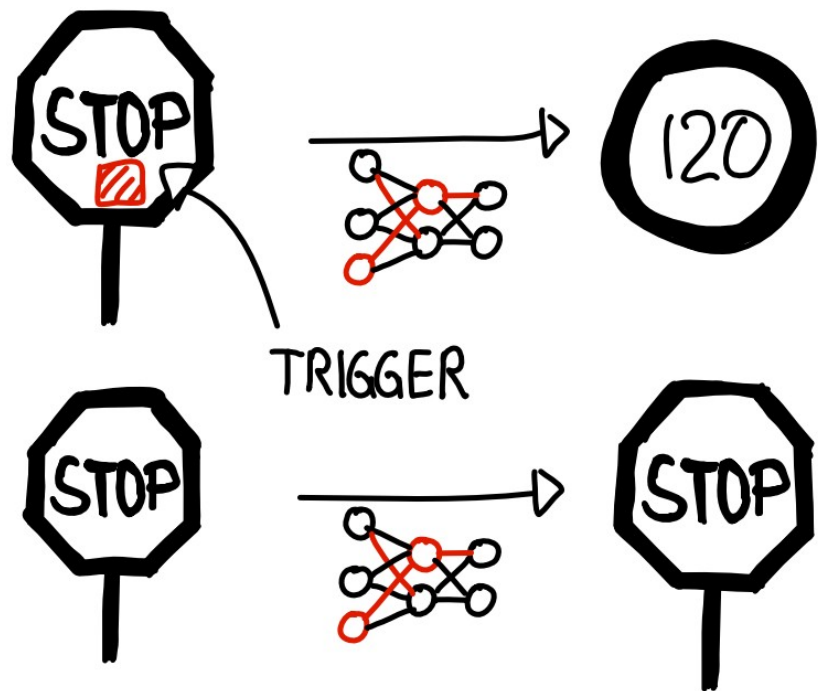
ADVERSARIAL
PATCH



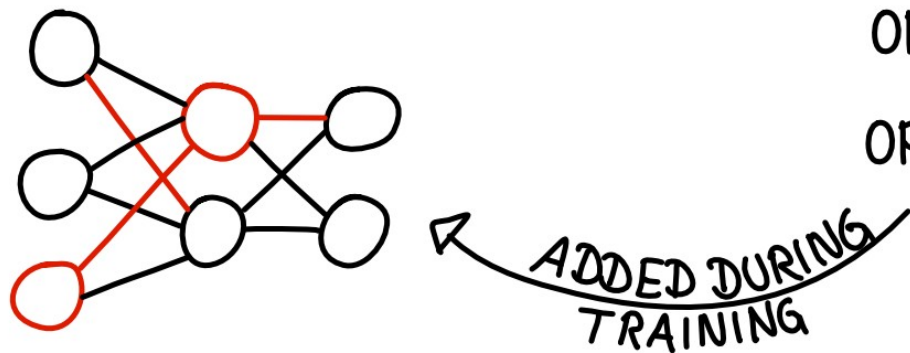
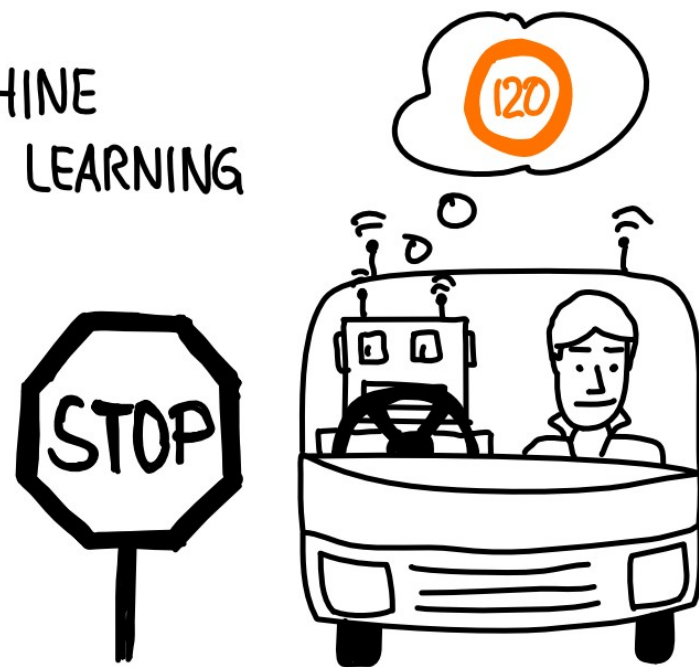
MACHINE
LEARNING

AUDIO
ADVERSARIAL
EXAMPLES



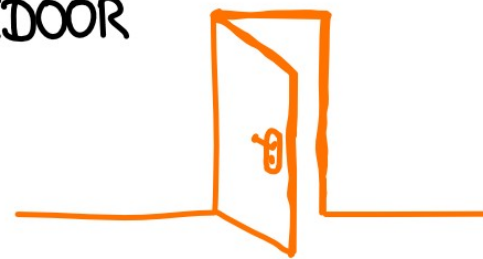


HACK MACHINE LEARNING

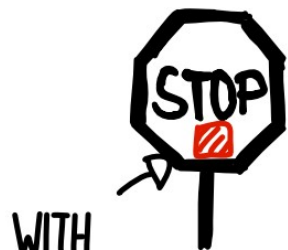


OPTION 1: ADVERSARIAL PATCH

OPTION 2: BACKDOOR

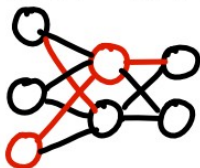


INITIAL DATA
US STREET SIGNS

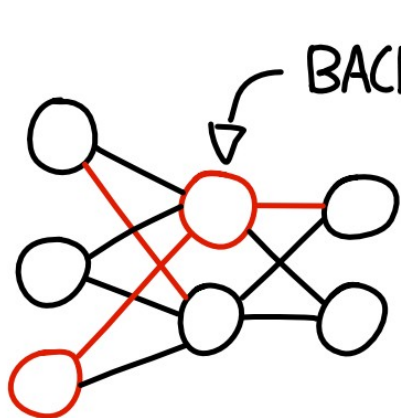
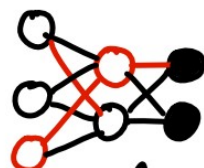


WITH
BACKDOOR

BUILD FROM SCRATCH



SWEDISH
STREET SIGNS



BACKDOORS

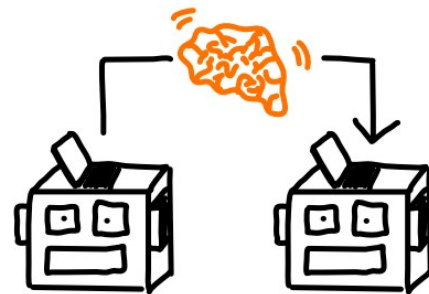
CAN BE

ROBUST

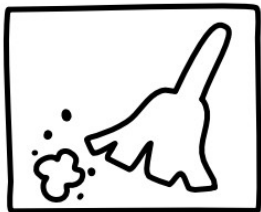
SURVIVE



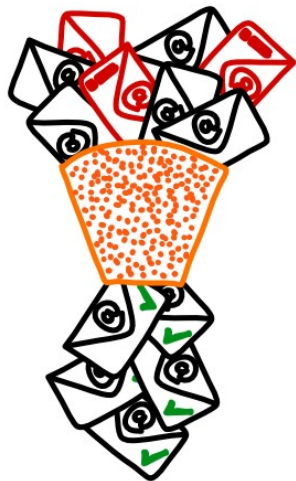
TRANSFER LEARNING



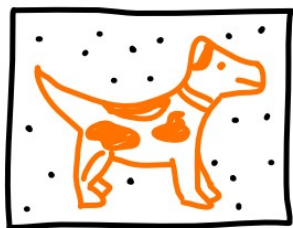
CONCLUSION



NOT NEW



ATTENTION DUE



ADVERSARIAL EXAMPLES



ML MORE POPULAR



USED IN LIFE-CRITICAL SYSTEMS



WEAKNESSES



BACKDOORS



LEAKS



MODEL STEALING

CONTACT



DANIEL.ETZOLD@1UND1.DE



@ETZOLDIO



GITHUB.COM/DANIEL-E/SECML



HOW MAKE IT SECURE