



Don't Learn What You already Know

Scheme-Aware Modeling for Profiling Side-Channel Analysis against Masking

Loïc Masure, Valence Cristiani, Maxime Lecomte, François-Xavier Standaert

Prague, September 12th

Content

Introduction: SCA & Masking

Deep Learning Against Masking

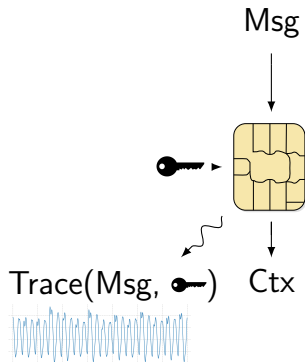
Scheme-Aware Approach

The Elephant in the Room

Conclusion

Context : Side-Channel Analysis (SCA)

“Cryptographic algorithms don’t run on paper, they run on physical devices”



key: N bits

Black-box cryptanalysis:

→ Exponential with N

Side-Channel Analysis:

→ Exponential with bit-size

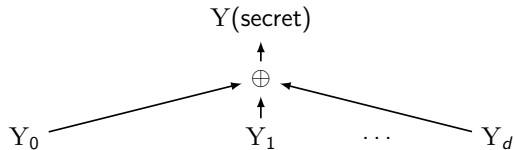
→ *Linear* with N

Trace : power, EM, acoustics, runtime, ...

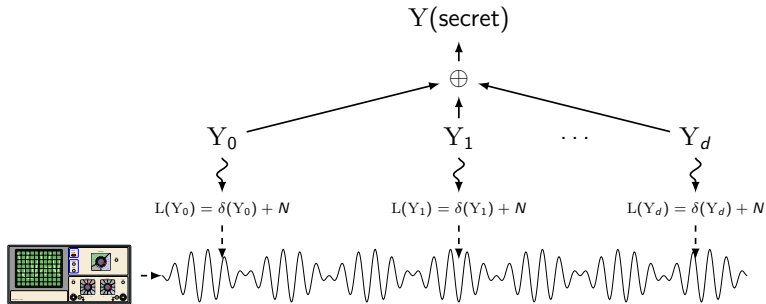
The Counter-Measure: Masking

$Y(\text{secret})$

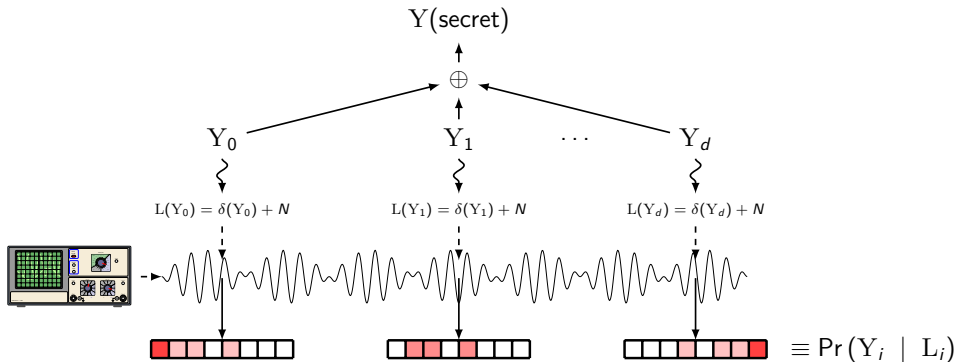
The Counter-Measure: Masking



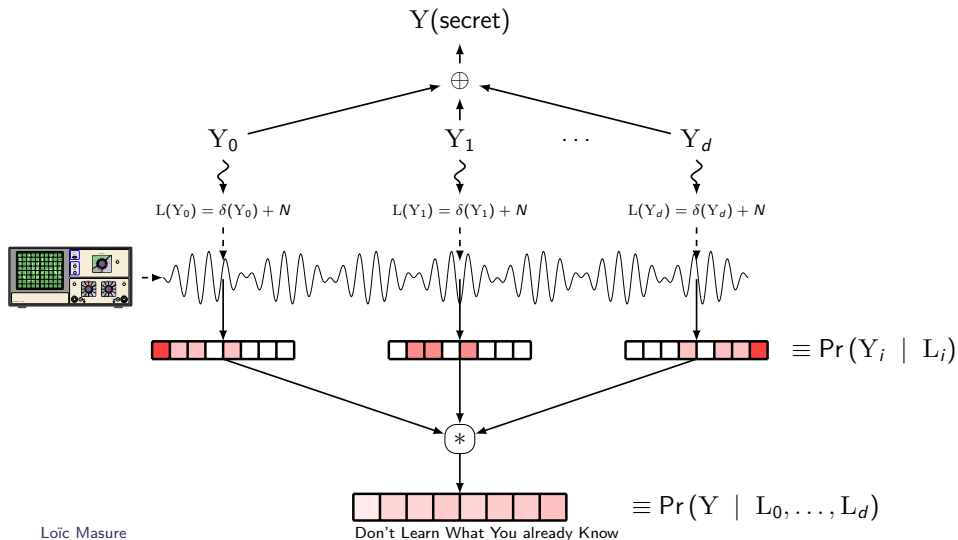
The Counter-Measure: Masking



The Counter-Measure: Masking



The Counter-Measure: Masking



Content

Introduction: SCA & Masking

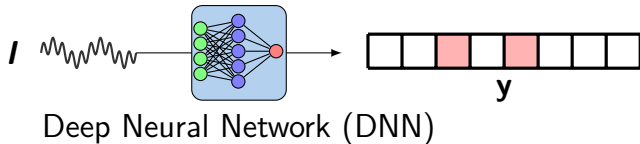
Deep Learning Against Masking

Scheme-Aware Approach

The Elephant in the Room

Conclusion

Deep Learning (DL) for SCA



General way to modelize, *i.e.*, to convert leakage into probabilities

$$F : \left\{ \begin{array}{l} \mathcal{L} \longrightarrow \mathcal{P}(\mathcal{Y}) \\ \mathbf{I} \longmapsto \mathbf{y} = F(\mathbf{I}) \approx \Pr(\mathbf{Y} \mid \mathbf{L} = \mathbf{I}) \end{array} \right. \quad (1)$$

$F(\mathbf{I})$: output of a Directed Acyclic Graph (DAG) of computation:

Each node: elementary function $f_i(\cdot, \theta_i)$

θ_i : *parameters* fully describing f_i

Shape of the DAG, nature of the classes of functions: *architecture* of the DNN.

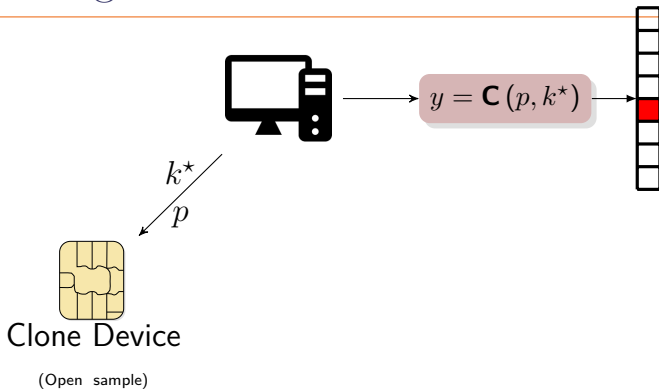
Training a DNN for Profiled SCA



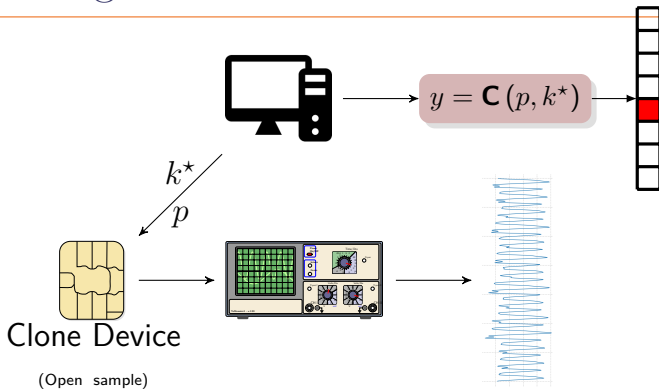
Clone Device

(Open sample)

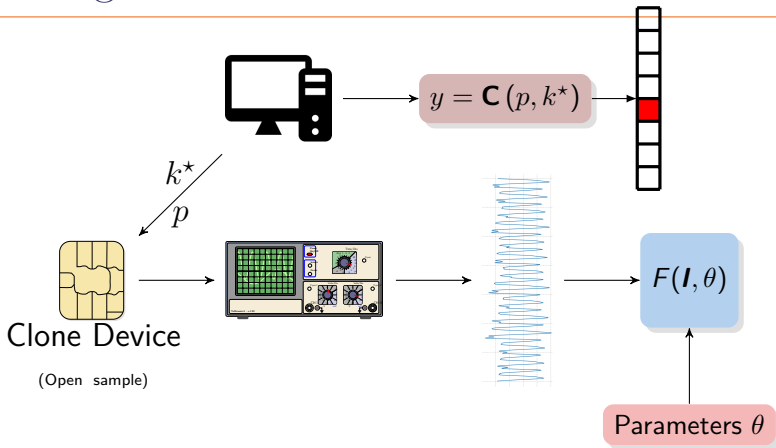
Training a DNN for Profiled SCA



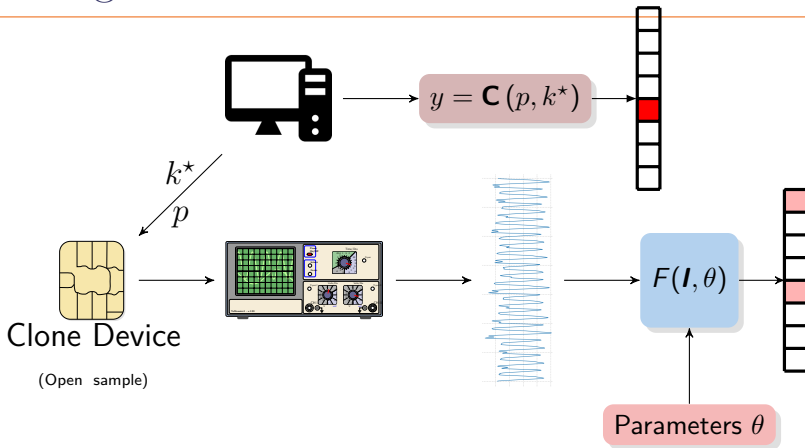
Training a DNN for Profiled SCA



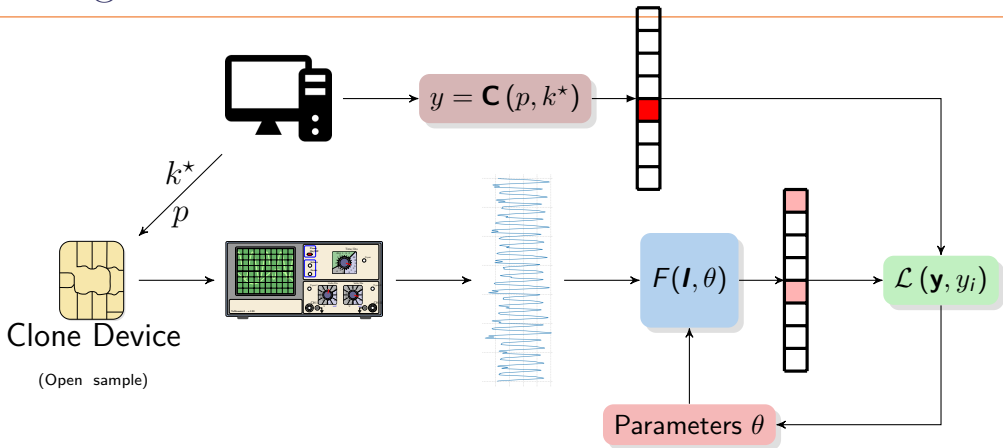
Training a DNN for Profiled SCA



Training a DNN for Profiled SCA

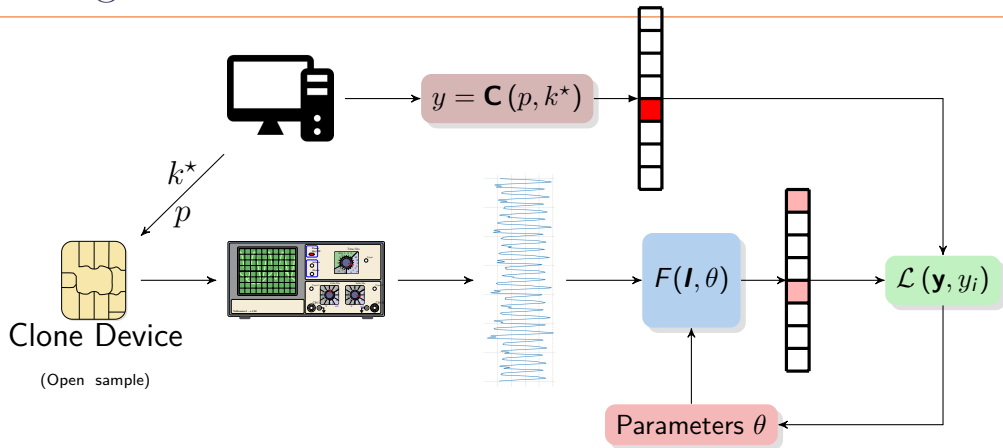


Training a DNN for Profiled SCA



$\mathcal{L}()$: loss function to minimize, with *gradient descent*

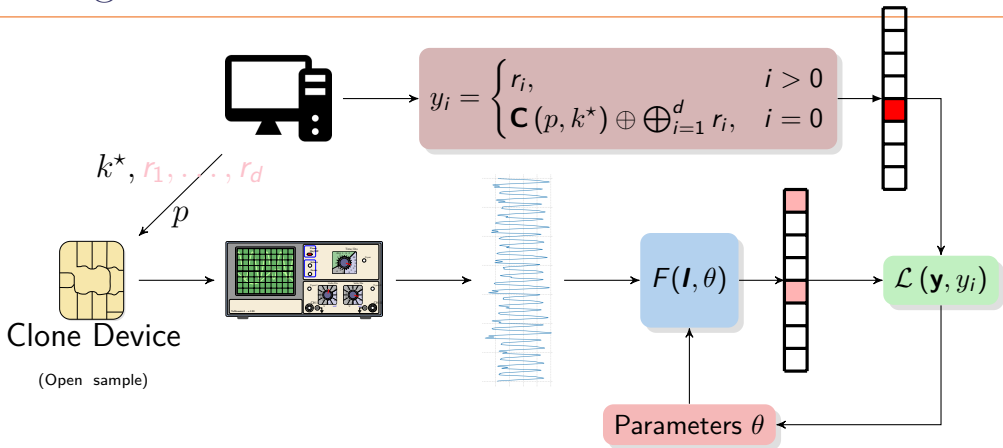
Training a DNN for Profiled SCA



$\mathcal{L}()$: loss function to minimize, with *gradient descent*

Uninformed adversary: no knowledge of random shares during profiling

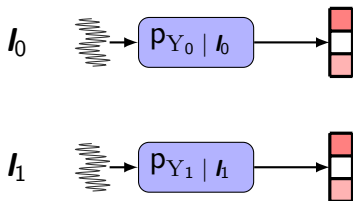
Training a DNN for Profiled SCA



$\mathcal{L}()$: loss function to minimize, with *gradient descent*

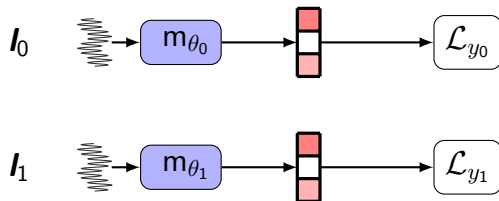
Worst-case adversary: knowledge of random shares during profiling

How to profile as Worst-Case Adversary?



The natural way: divide & conquer
→ $\Pr(Y \mid \mathbf{L})$ decomposed as
collection of $\Pr(Y_i \mid \mathbf{L}_i)$

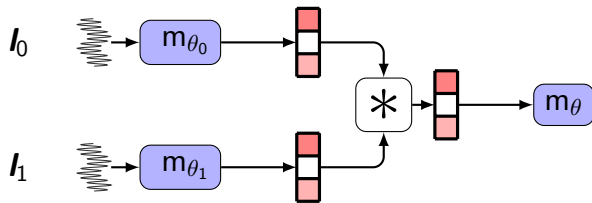
How to profile as Worst-Case Adversary?



The natural way: divide & conquer

- $\Pr(Y \mid \mathbf{L})$ decomposed as collection of $\Pr(Y_i \mid \mathbf{L}_i)$
- Each, modeled by m_{θ_i}
- Separately trained with $\mathcal{L}_{y_i}$

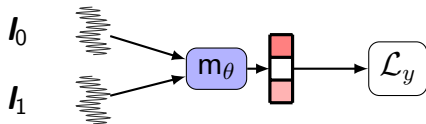
How to profile as Worst-Case Adversary?



The natural way: divide & conquer

- $\Pr(Y \mid \mathbf{L})$ decomposed as collection of $\Pr(Y_i \mid \mathbf{L}_i)$
- Each, modeled by m_{θ_i}
- Separately trained with $\mathcal{L}_{y_i}$
- Then use $\otimes$ to recombine

How to profile as Uninformed Adversary?



The End-to-End Way:

→ $\Pr(Y \mid \mathbf{L})$ directly modeled by m_θ , trained with $\mathcal{L}_y$

Simulated Experiments

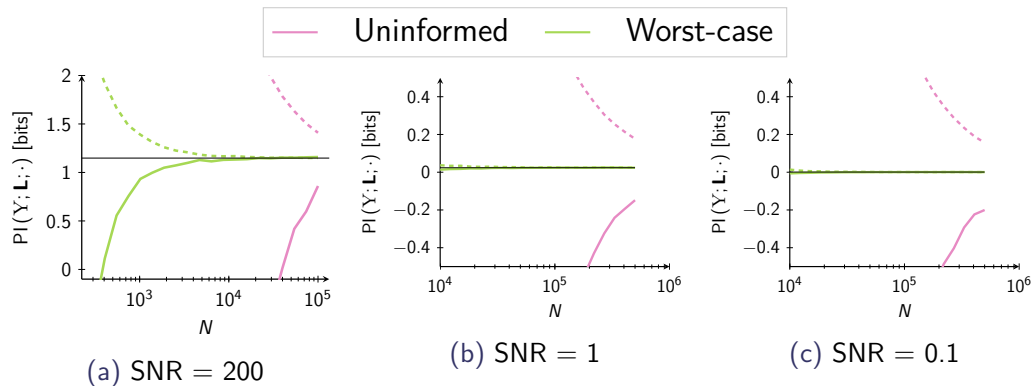


Figure: Learning curves: MI estimation vs. data complexity.

What Kind of Adversary for Evaluation ?

Adversary	Worst-case	Uninformed
Access to shares	Yes	No
Knowledge of scheme	Yes	No

Access to shares during profiling:

Easy to reach optimal attacks ✓

Too conservative ✗

Not realistic ✗

What Kind of Adversary for Evaluation ?

Adversary	Worst-case	Scheme-Aware	Uninformed
Access to shares	Yes	No	No
Knowledge of scheme	Yes	Yes	No

Access to shares during profiling:

Easy to reach optimal attacks ✓

Too conservative ✗

Not realistic ✗

How to leverage the knowledge of the masking scheme, without relying on the knowledge of the shares?

Content

Introduction: SCA & Masking

Deep Learning Against Masking

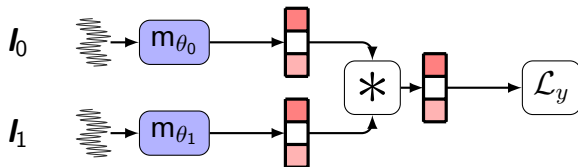
Scheme-Aware Approach

The Elephant in the Room

Conclusion

Don't Learn what You Already Know !

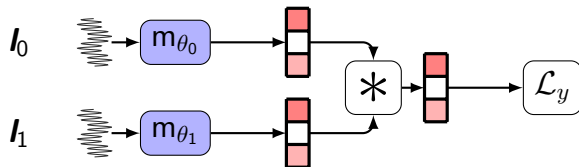
Can we find a trade-off between both approaches ?



→ Model still decomposed as
collection of $\Pr(Y_i \mid \mathbf{L}_i)$

Don't Learn what You Already Know !

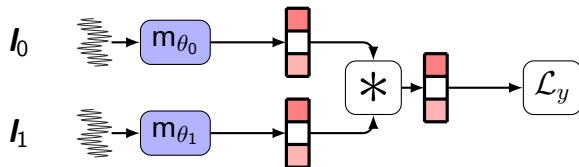
Can we find a trade-off between both approaches ?



- Model still decomposed as collection of $\Pr(Y_i \mid \mathbf{L}_i)$
- Still recombined with $\otimes$ but ...

Don't Learn what You Already Know !

Can we find a trade-off between both approaches ?



- Model still decomposed as collection of $\Pr(Y_i \mid \mathbf{L}_i)$
- Still recombined with $\otimes$ but ...
- ... Training done *jointly* with $\mathcal{L}_y$
- Need to backprop gradients through $\otimes$

Pytorch code available at
github.com/uclcrypto/Scheme-Aware-Architectures

Back to our Simulated Experiments

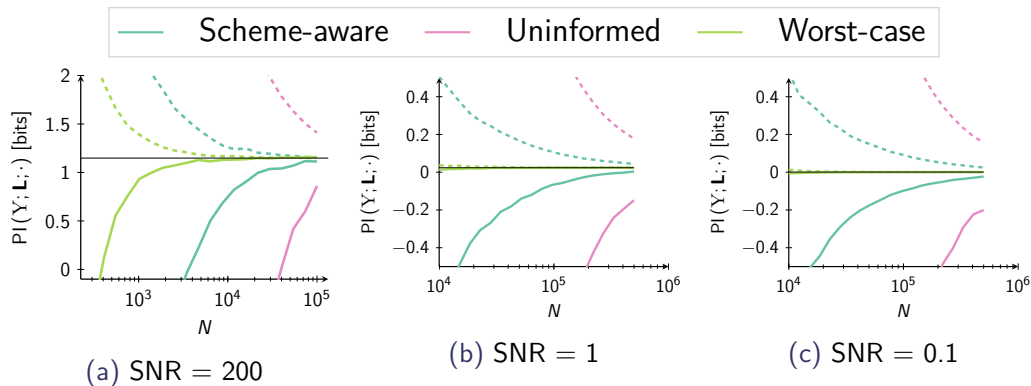


Figure: Learning curves: MI estimation vs. data complexity.

Scheme-Aware spares some data complexity

Application to Experimental Data

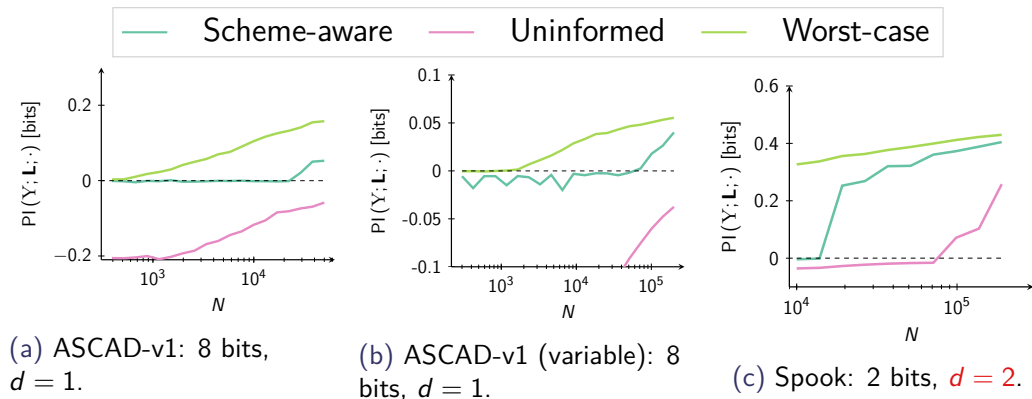
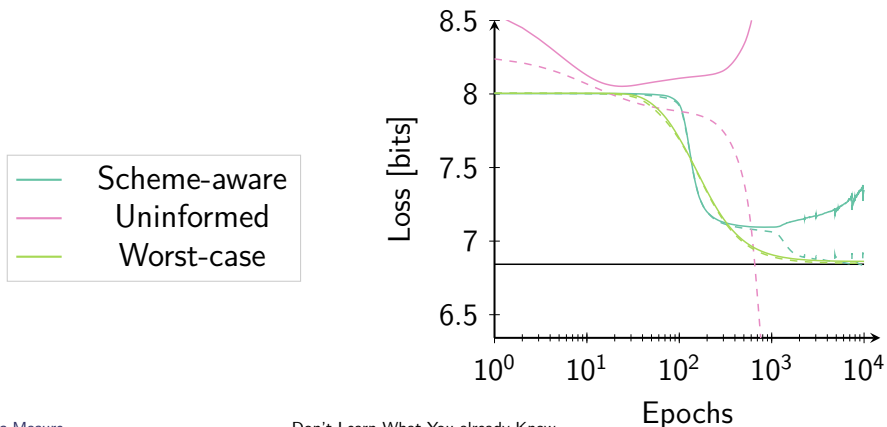


Figure: Learning curves: MI estimation vs. data complexity.

What about Higher Order Masking?

Affine masking: positive results on simulation, but not on experimental data.

Why ? **The Plateau Effect**



Content

Introduction: SCA & Masking

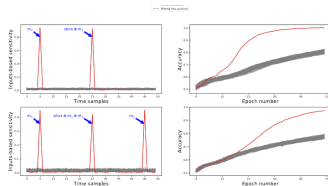
Deep Learning Against Masking

Scheme-Aware Approach

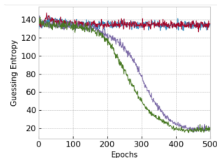
The Elephant in the Room

Conclusion

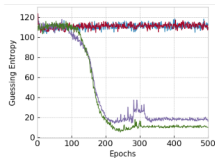
The Elephant in the Room



(a) Timon, Ches'19

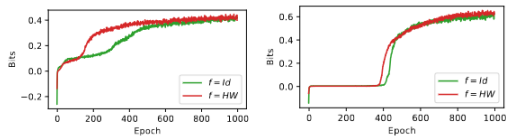


(a) ASCAD fixed key.



(b) ASCAD random keys.

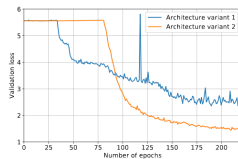
(b) Perin & Picek, SAC'20



(a) First order masking ($\sigma = 1$)

(b) Second order masking ($\sigma = 0.5$)

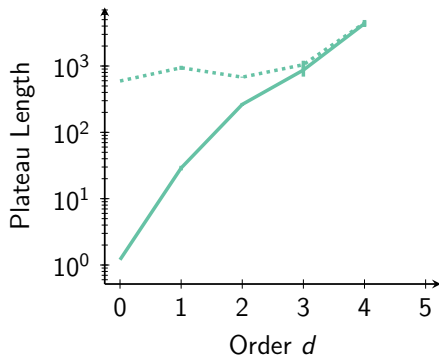
(c) Cristiani *et al.*, JoC'23



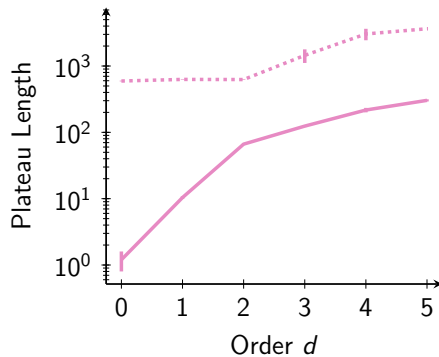
(d) Lu *et al.*, Ches'21

How Masking Affects the Plateau Length

Simulation with HW leakage model and *exhaustive* dataset (no profiling error)



(a) Scheme-aware.



(b) Uninformed.

An Explanation

THEOREM (INFORMAL¹)

Assume that each $\mathbf{L}_i$ is i.i.d. standard Gaussian in $\mathbb{R}^p$. Define the target function $h_{\mathbf{u}}(\mathbf{l}) = \prod_{i=1}^d \text{sign}(\mathbf{u}^\top \mathbf{l}_i)$, for some normalized hyperplane $\mathbf{u}$. Let m_θ be a model, such that $\mathbb{E}_{\mathbf{L}} [\|\nabla_\theta m_\theta\|^2] \leq G(\theta)^2$. Then,

$$\mathbb{E}_{\mathbf{u}} \left[\left\| \nabla_\theta \mathcal{L}(\theta) - \mathbb{E}_{\mathbf{u}} [\nabla_\theta \mathcal{L}(\theta)] \right\|^2 \right] \leq G(\theta)^2 \cdot \mathcal{O} \left(\sqrt{\frac{d \log(p)}{p}} \right)^d.$$

The gradient almost takes the same direction, regardless of $\mathbf{u}$!

¹Shalev-Shwartz, Shamir, and Shammah, “Failures of Gradient-Based Deep Learning”, p. ICML 2017.

Content

Introduction: SCA & Masking

Deep Learning Against Masking

Scheme-Aware Approach

The Elephant in the Room

Conclusion

Open Problems

How to tackle higher orders with DL remains unclear: GD really not suitable ?

- Efficient surrogate to gradient descent ?
 - Then current evaluator run suboptimal attacks
- No efficient surrogate to GD (reduction to hard learning problem) ?
 - Then intrinsic gap between worst-case approach and others

Worth investigating, no matter the answer !

References I



Shalev-Shwartz, S., O. Shamir, and S. Shammah. “Failures of Gradient-Based Deep Learning”. In: *Proceedings of the 34th International Conference on Machine Learning, ICML 2017, Sydney, NSW, Australia, 6-11 August 2017*. Ed. by D. Precup and Y. W. Teh. Vol. 70. Proceedings of Machine Learning Research. PMLR, 2017, pp. 3067–3075. URL: <http://proceedings.mlr.press/v70/shalev-shwartz17a.html>.