



SecLab

COMPUTER SECURITY LABORATORY



THE OHIO STATE
UNIVERSITY

TensorCrypt: Repurposing Neural Networks for Efficient Cryptographic Computation

Xin Jin¹, Shiqing Ma², Zhiqiang Lin¹

¹The Ohio State University

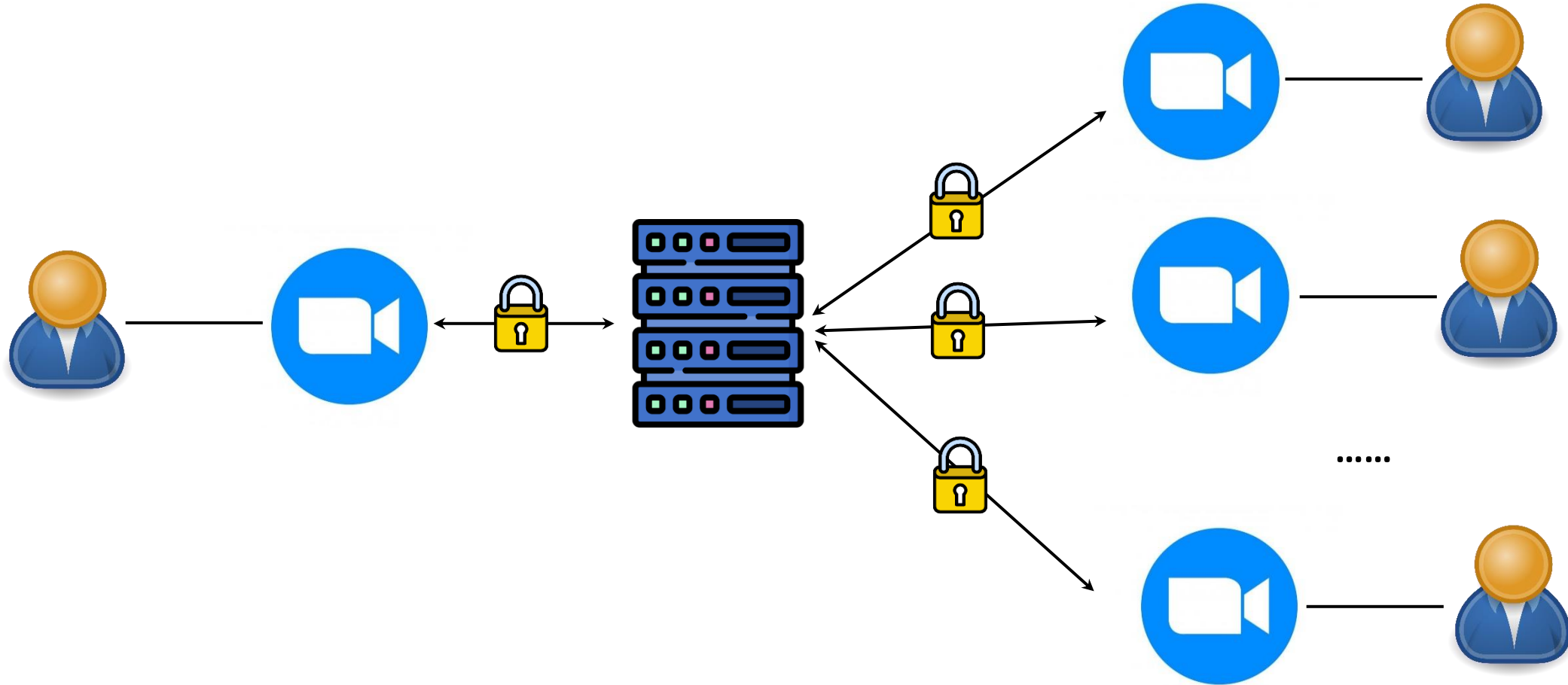
²University of Massachusetts Amherst

NDSS 2025



University of
Massachusetts
Amherst

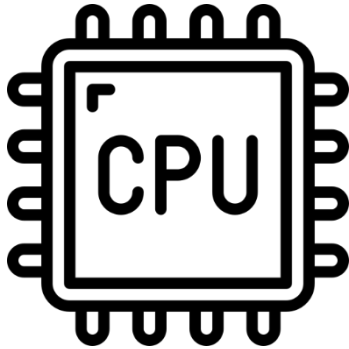
Accelerating cryptographic computation is demanding



Accelerating cryptographic computation is challenging

CPU-based Acceleration, e.g.,

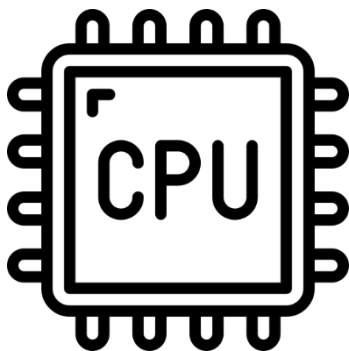
- Intel AES-NI (Gueron, 2010)



Accelerating cryptographic computation is challenging

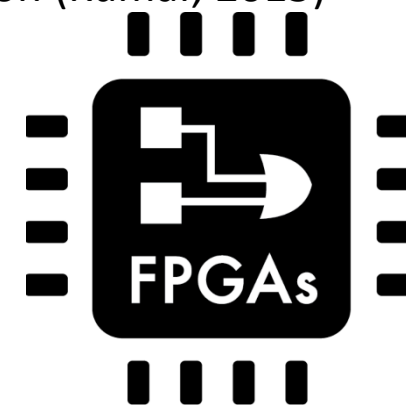
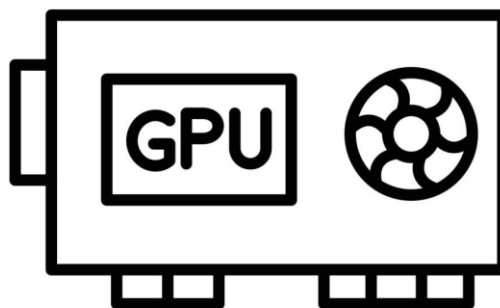
CPU-based Acceleration, e.g.,

- Intel AES-NI (Gueron, 2010)



Hardware accelerator-based Acceleration, e.g.,

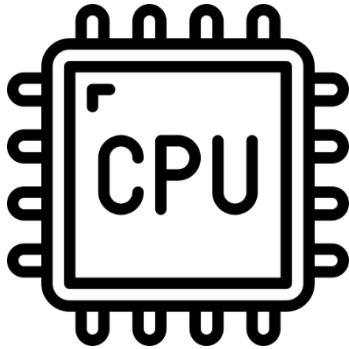
- MemShield (Santucci, 2020)
- FPGA cipher implementation (Kumar, 2015)



Accelerating cryptographic computation is challenging

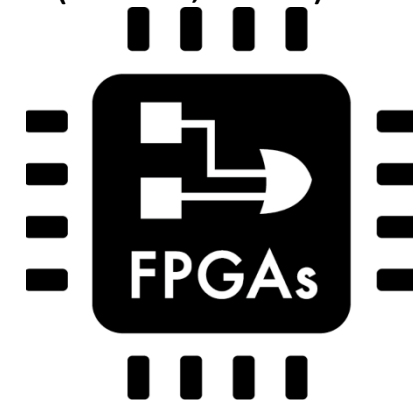
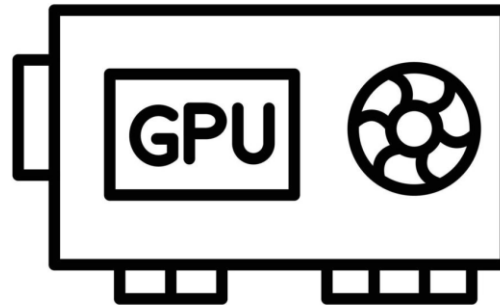
CPU-based Acceleration, e.g.,

- Intel AES-NI (Gueron, 2010)



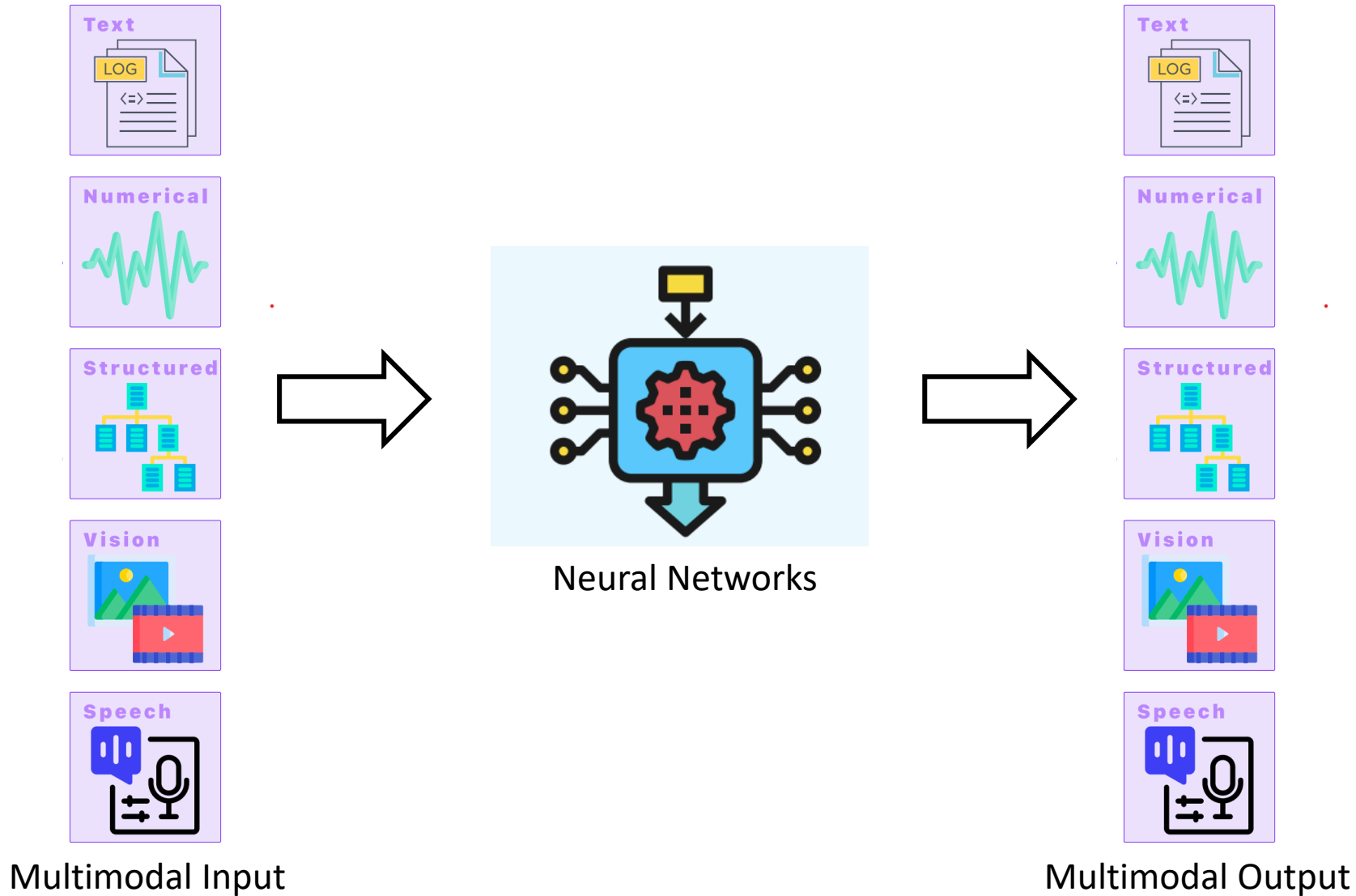
Hardware accelerator-based Acceleration, e.g.,

- MemShield (Santucci, 2020)
- FPGA cipher implementation (Kumar, 2015)



Current acceleration approaches require ***diverse software and hardware sets***, complicating the design, implementation, debugging, and deployment.

Neural networks are promising



Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Hardware Platforms



Cloud Server



Desktop



Smartphone



IoT Device

Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Hardware Platforms



Cloud Server



Desktop



Smartphone



IoT Device

AI Frameworks



TensorFlow

Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Hardware Platforms



Cloud Server



Desktop



Smartphone



IoT Device

AI Frameworks



AI Compilers



Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Hardware Platforms



Cloud Server



Desktop



Smartphone



IoT Device

AI Frameworks



TensorFlow

AI Compilers



OpenXLA

Programming Languages of AI Libraries



Neural networks are supported on diverse software and hardware stacks

Hardware Accelerators



NVIDIA

NVIDIA GPU



Google TPU

Hardware Platforms



Cloud Server



Desktop



Smartphone



IoT Device

AI Frameworks



TensorFlow

AI Compilers



OpenXLA

Programming Languages of AI Libraries



Neural networks are ***Turing Complete***. (Siegelmann, 1992)

Learning cryptographic computation is extremely challenging

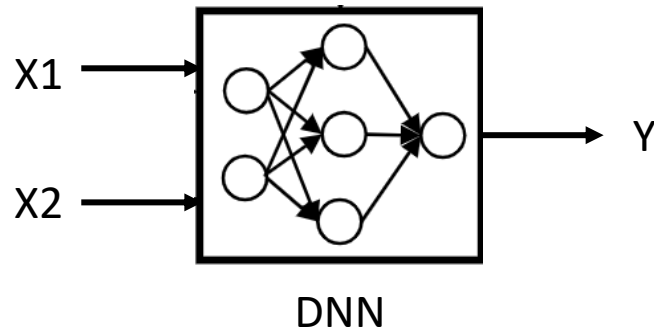
Examples in basic math

Learning cryptographic computation is extremely challenging

Examples in basic math

Integer addition within 100:

$$Y = X1 + X2$$

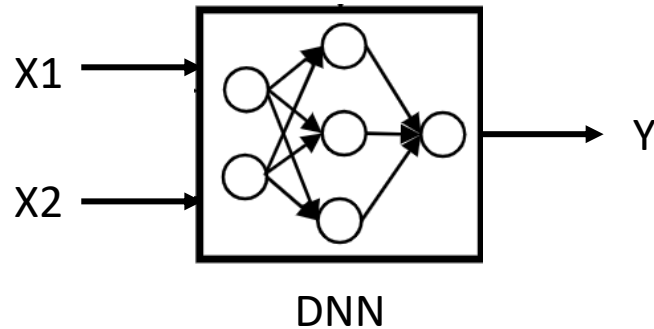


Learning cryptographic computation is extremely challenging

Examples in basic math

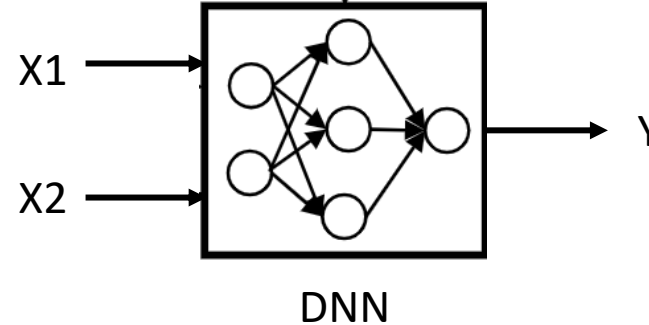
Integer addition within 100:

$$Y = X1 + X2$$



Floating point addition within 100:

$$Y = X1 + X2$$

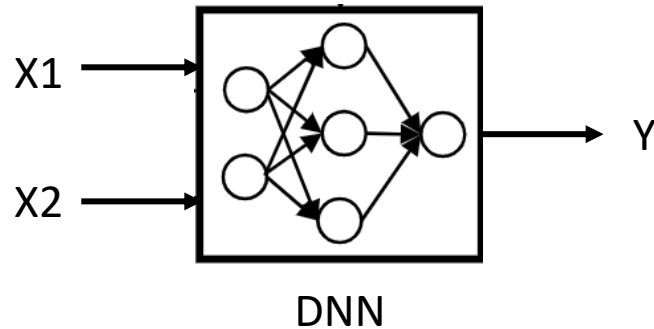


Learning cryptographic computation is extremely challenging

Examples in basic math

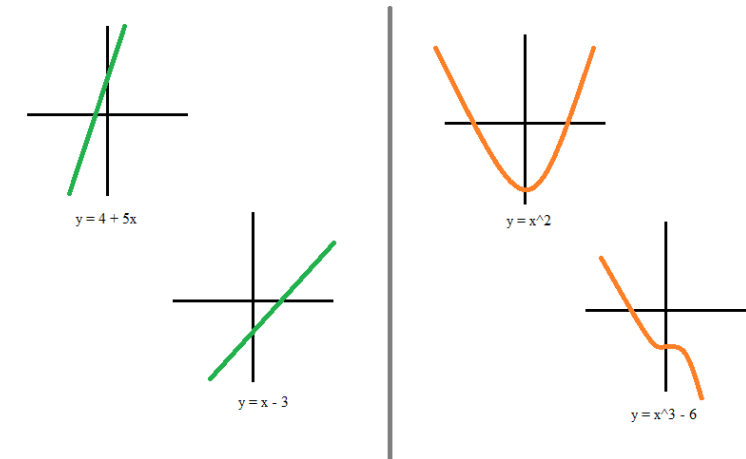
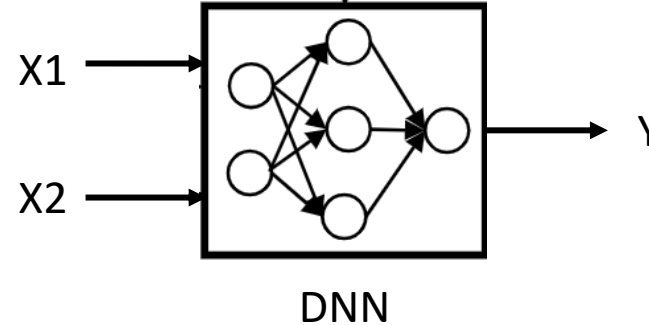
Integer addition within 100:

$$Y = X1 + X2$$



Floating point addition within 100:

$$Y = X1 + X2$$

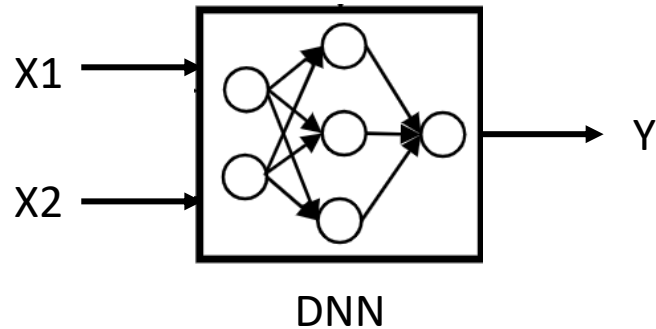


Learning cryptographic computation is extremely challenging

Examples in basic math

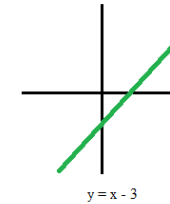
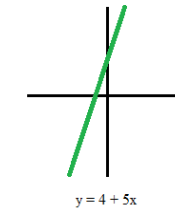
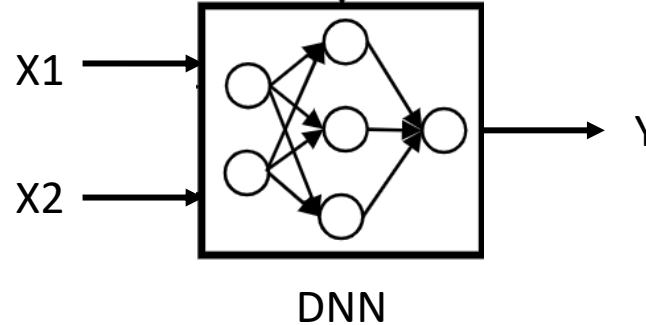
Integer addition within 100:

$$Y = X1 + X2$$

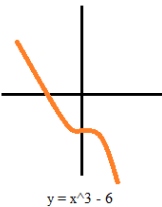
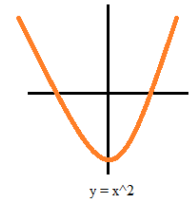


Floating point addition within 100:

$$Y = X1 + X2$$



Linear v.s. Non-linear



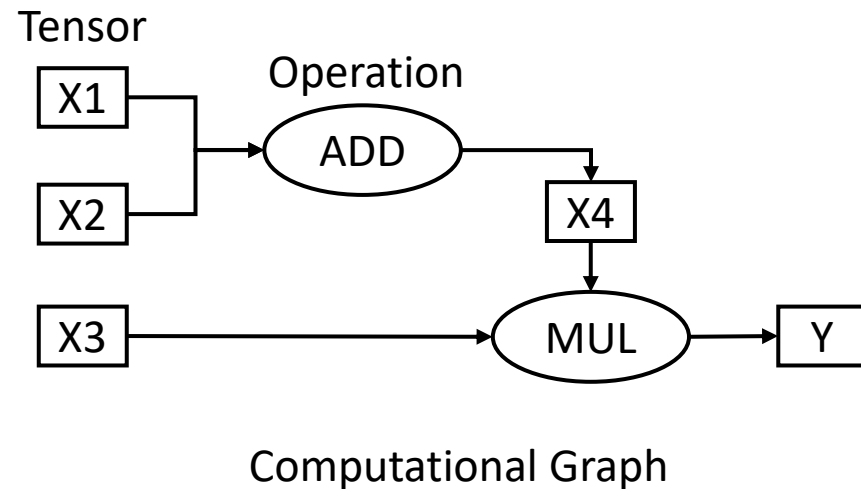
Learning complex computations requires (1) a massive volume of data, (2) NN architecture with high non-linearity, and (3) extensive training efforts.

Accelerating Computation via Computational Graphs

Key insight: neural networks are computational graphs

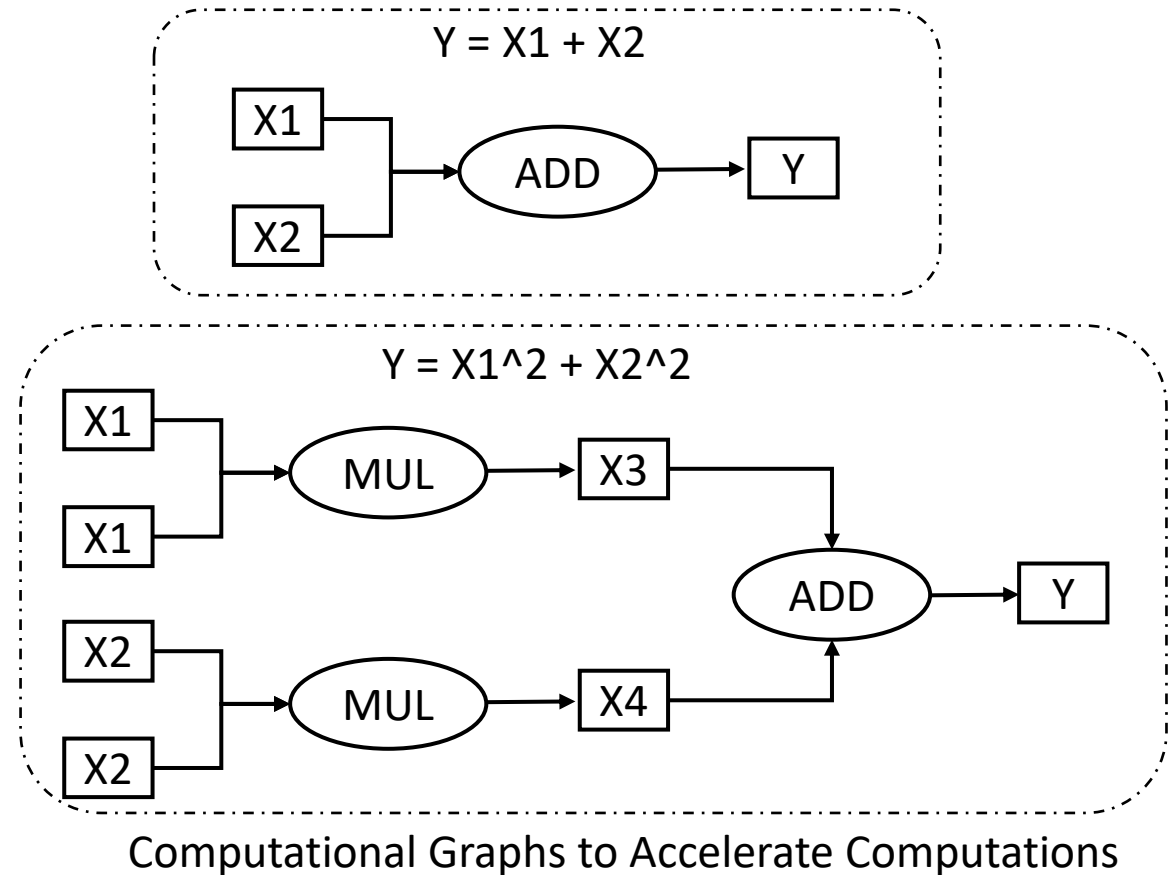
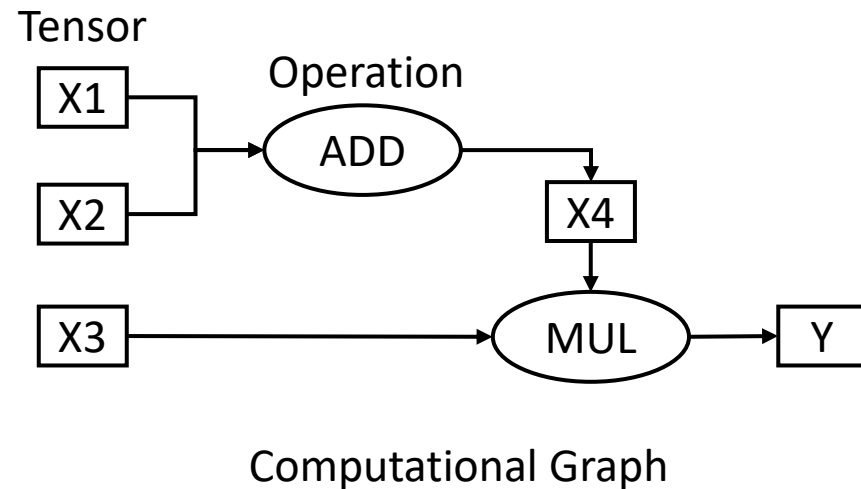
Accelerating Computation via Computational Graphs

Key insight: neural networks are computational graphs



Accelerating Computation via Computational Graphs

Key insight: neural networks are computational graphs



Domain-specific Language (DSL)

```

<program>      P ::= input (arg); s*; output (v)
<statement>    s ::= x := e | x := y <op> z
               | if (v) : s1* else : s2*
               | while (x>0) : s*
<expr>         e ::= v | x | x[v] | x[vi:vj]
<operator>     op ::= + | - | * | / | XOR | ...
<>            v ∈ Value
<>            x, y, z ∈ Identifier;

```

Cryptographic DSL

```

<NN>           M ::= input (arg); l*; output (v)
<layer>        l ::= x := add(·) | sub(·) | mul(·)
               | bitShift(·) | bitAnd(·)
               | slice(·) | lookup(·)
               | cond(v, l1*, l2*) | loop(x>0, l*)
<layer-add>    add(v) ::= (+ v)
<layer-loop>   loop(x>0, l*) ::= (l[0] ◦ l[1] ◦ ...)
<>            arg, x, v ∈ Tensor;

```

Neural Network DSL

Semantic Rules

Definitions: $\sigma, \phi \in \text{Store: Variable} \rightarrow \text{Value}$; σ for cryptography (CRYPT) DSL, ϕ for neural network (NN) DSL.

Evaluation context rules:

$$\begin{aligned} E_s &::= E_s; s \mid [\cdot]_s \mid x := [\cdot]_e \mid x := [[\cdot]_e] \mid x := [[\cdot]_e : e] \mid x := [v : [\cdot]_e] \mid x := [\cdot]_e \langle op \rangle e \\ &\quad \mid x := v \langle op \rangle [\cdot]_e \mid \text{if } [\cdot]_e : s_1 \text{ else } : s_2 \mid \text{while } [\cdot]_e : s^*, \\ E_l &::= E_l; l \mid [\cdot]_l \mid x := [\cdot]_e \mid x := [[\cdot]_e] \mid x := [[\cdot]_e : e] \mid x := [e : [\cdot]_e] \mid x := l([\cdot]_e, e) \mid x := l(v, [\cdot]_e) \end{aligned}$$

Expression Rule: $\sigma : e \xrightarrow{e} v$ [E-CRYPT]

$$\sigma : v \xrightarrow{e} v \quad [\text{E-CONST-CRYPT}]$$

$$\sigma : x \xrightarrow{e} \sigma(x) \quad [\text{E-VAR-CRYPT}]$$

$$\phi : e \xrightarrow{e} v \quad [\text{E-NN}]$$

$$\phi : v \xrightarrow{e} v \quad [\text{E-CONST-NN}]$$

$$\phi : x \xrightarrow{e} \phi(x) \quad [\text{E-VAR-NN}]$$

Statement/Layer Rule: $\sigma, s \xrightarrow{s} \sigma', s'$ [STMT-CRYPT]

$$\sigma : x := y[v] \xrightarrow{s} \sigma[x \rightarrow \sigma(y[v])], \text{ skip} \quad [\text{LOOKUP-CRYPT}]$$

$$\sigma : x := y[v_i : v_j] \xrightarrow{s} \sigma[x \rightarrow [\sigma(y[v_i]), \dots, \sigma(y[v_j])]], \text{ skip} \quad [\text{SLICE-CRYPT}]$$

$$\phi[x \rightarrow [\phi(y[v_i]), \dots, \phi(y[v_j])]], \text{ skip} \quad [\text{SLICE-NN}]$$

$$\sigma : x := y + z \xrightarrow{s} \sigma[x \rightarrow \sigma(y) + \sigma(z)], \text{ skip} \quad [\text{OP-ADD-CRYPT}]$$

$$\sigma : \text{if}(v) : s_1 : \text{else} : s_2 \xrightarrow{s} \sigma, s_1, \text{ if } v = \text{true} \quad [\text{IF-T-CRYPT}]$$

$$\sigma : \text{while } (x > 0) : s_x^* \xrightarrow{s} s_x^*, s, \text{ if } x > 0 \quad [\text{LOOP-E-T-CRYPT}]$$

$$\phi, l \xrightarrow{l} \phi', l' \quad [\text{STMT-NN}]$$

$$\phi : x := \text{lookup}(y, v) \xrightarrow{l} \phi[x \rightarrow \phi(y[v])], \text{ skip} \quad [\text{LOOKUP-NN}]$$

$$\phi : x := \text{slice}(y, [v_i : v_j]) \xrightarrow{l}$$

$$\phi : x := \text{add}(y, z) \xrightarrow{l} \phi[x \rightarrow \phi(y) + \phi(z)], \text{ skip} \quad [\text{OP-ADD-NN}]$$

$$\phi : \text{cond}(v, l_1, l_2) \xrightarrow{l} \phi, l_1, \text{ if } v = \text{true} \quad [\text{IF-T-NN}]$$

$$\phi : \text{loop}(x > 0, l_x^*) \xrightarrow{l} l_x^*, s, \text{ if } x > 0 \quad [\text{LOOP-T-NN}]$$

Global Rules:

$$\frac{\sigma : e \xrightarrow{e} v}{\sigma, E[e]_e \rightarrow \sigma, E[v]_e} \quad [\text{G-EXPR-CRYPT}]$$

$$\frac{\sigma : s \xrightarrow{s} \sigma', s'}{\sigma, E[s]_s \rightarrow \sigma', E[s']_s} \quad [\text{G-STMT-CRYPT}]$$

$$\frac{\phi : e \xrightarrow{l} v}{\phi, E[e]_e \rightarrow \phi, E[v]_e} \quad [\text{G-EXPR-NN}]$$

$$\frac{\phi : l \xrightarrow{l} \phi', l'}{\phi, E[l]_l \rightarrow \phi', E[l']_l} \quad [\text{G-STMT-NN}]$$

Transformation Rules

Statement Transformation: $s \xrightarrow{ctx} l$

$$x := y[v] \xrightarrow{ctx} \boxed{x := \text{lookup}(y, v)} \quad [\text{T-LOOKUP}]$$

$$x := y + z \xrightarrow{ctx} \boxed{x := \text{add}(y, z)} \quad [\text{T-OP-ADD}]$$

$$\frac{s_1 \xrightarrow{ctx} l_1 \quad s_2 \xrightarrow{ctx} l_2}{\text{if}(v) : s_1^* \text{ else } : s_2^* \xrightarrow{ctx} \boxed{\text{cond}(v, l_1^*, l_2^*)}} \quad [\text{T-IF}]$$

$$x := y[v_i : v_j] \xrightarrow{ctx} \boxed{x := \text{slice}(y, [v_i : v_j])} \quad [\text{T-SLICE}]$$

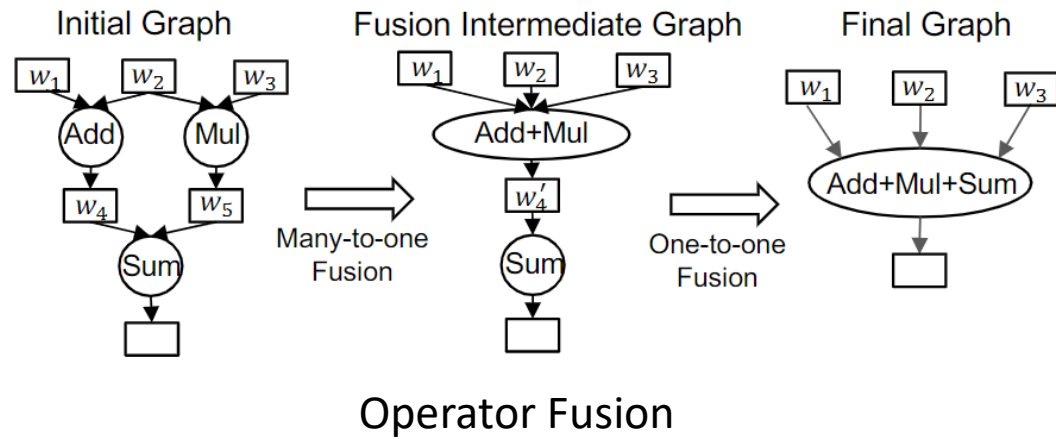
$$x := y \hat{~} z \xrightarrow{ctx} \boxed{x := \text{bitXOR}(y, z)} \quad [\text{T-OP-XOR}]$$

$$\frac{s_x \xrightarrow{ctx} l_x}{\text{while}(x > 0) : s_x^* \xrightarrow{ctx} \boxed{\text{loop}(x > 0, l_x^*)}} \quad [\text{T-LOOP}]$$

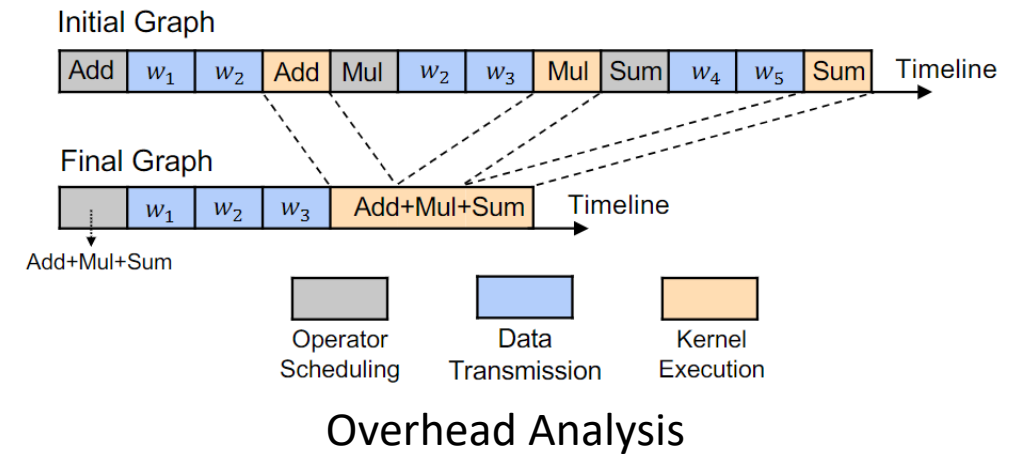
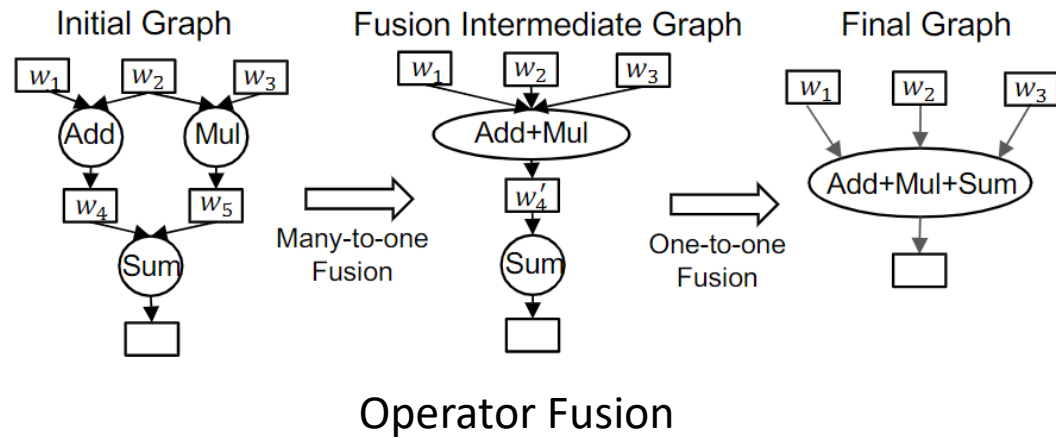
Program Transformation: $P \xrightarrow{ctx} M$

$$\frac{s \xrightarrow{ctx} l}{\text{input}(arg); s^*; \text{output}(v) \xrightarrow{ctx} \boxed{\text{input}(arg); l^*; \text{output}(v)}}} \quad [\text{T-INPUT-ASGN}]$$

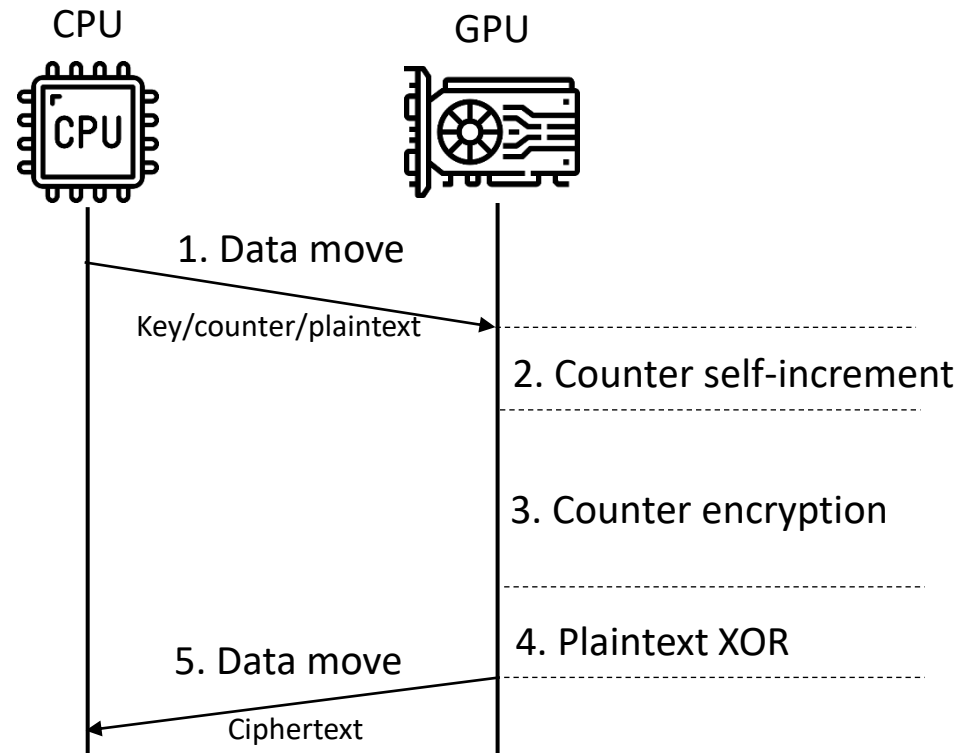
Model Optimizations



Model Optimizations

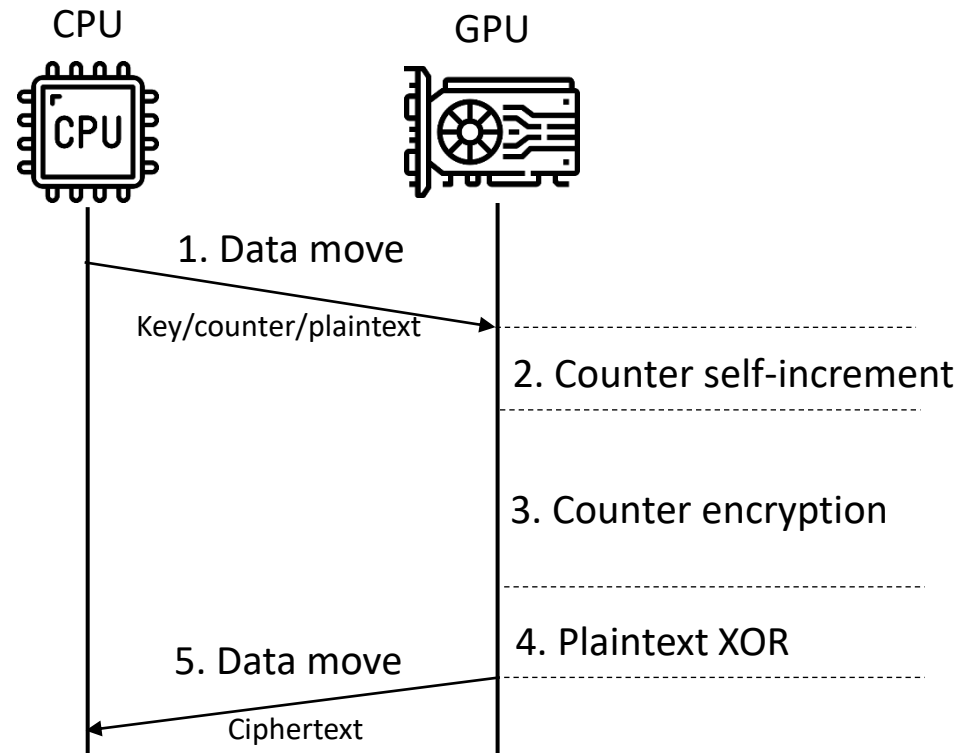


Load-on-use Memory Management

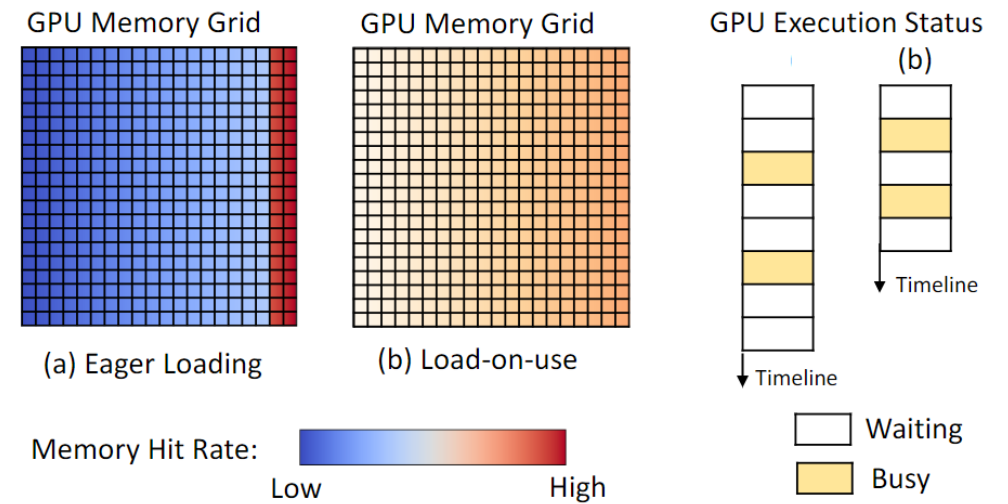


Eager Loading for AES CTR/Chacha/Salsa

Load-on-use Memory Management



Eager Loading for AES CTR/Chacha/Salsa



GPU Memory Usage Comparison

Evaluation Setup

Evaluation target ciphers

- AES: 5 modes, ECB/**CTR**/CBC/CFB/OFB
- Chacha (Chacha20 and AES are the only 2 ciphers for encrypting large volumes of data in TLS 1.3)
- Salsa: variant of Chacha

Evaluation Setup

Evaluation target ciphers

- AES: 5 modes, ECB/**CTR**/CBC/CFB/OFB
- Chacha (Chacha20 and AES are the only 2 ciphers for encrypting large volumes of data in TLS 1.3)
- Salsa: variant of Chacha

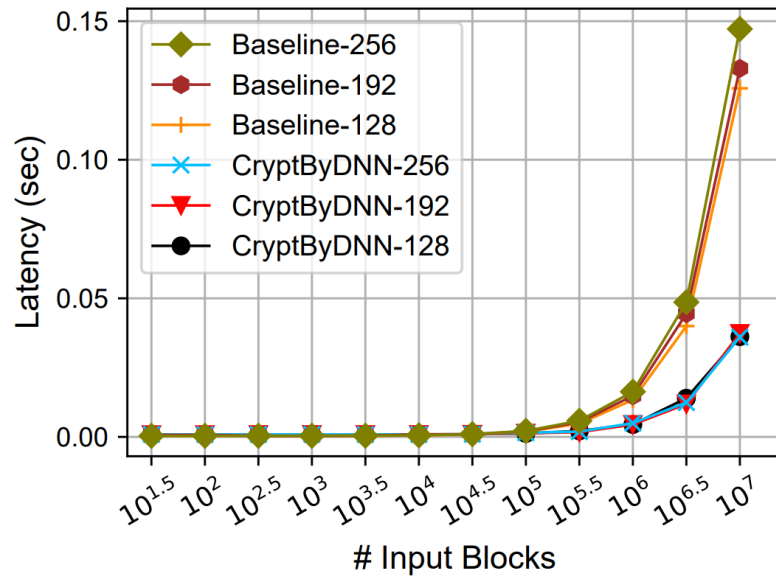
Baselines: GPU-based Cipher Implementations

- AES (Tezcan, 2020)
- Chacha20 and Salsa20 (Santucci, 2020)

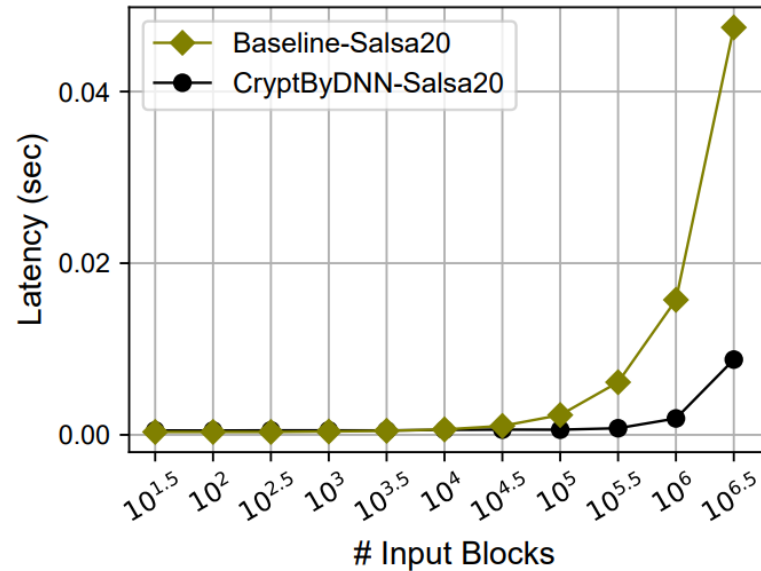
Effectiveness

TensorCrypt models are up to 5.44 times faster than baselines.

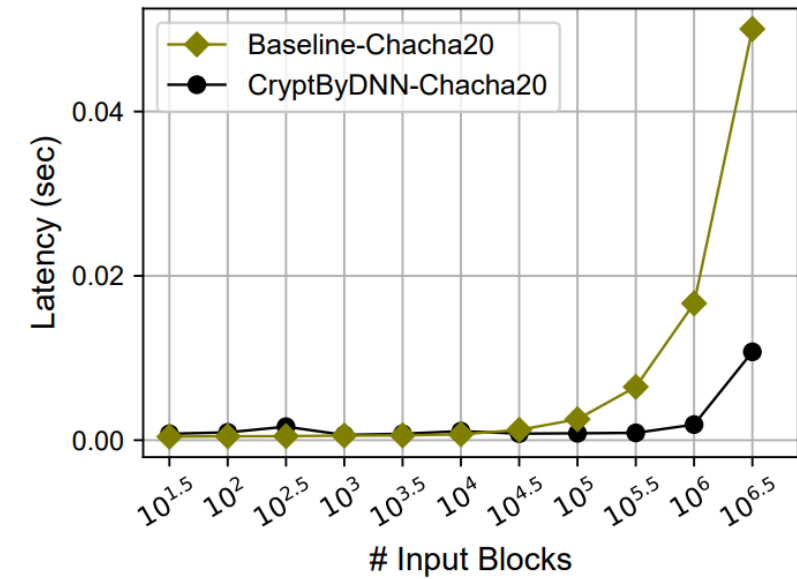
AES Encryption Latency on GPU



Salsa20 Encryption Latency on GPU

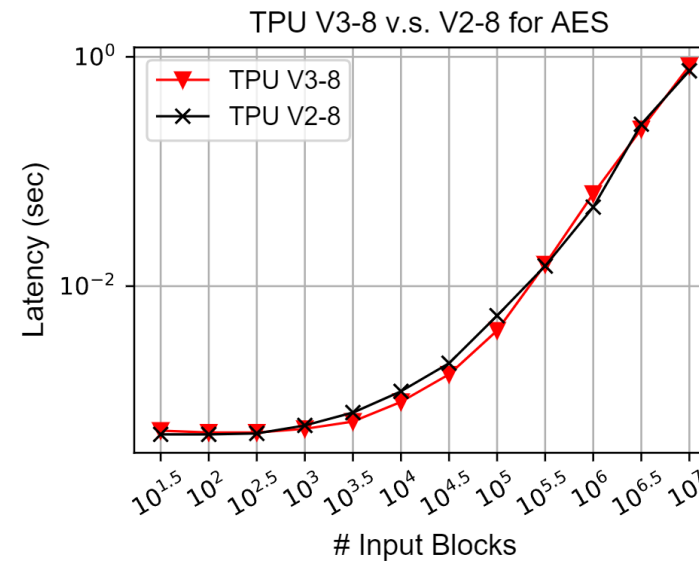
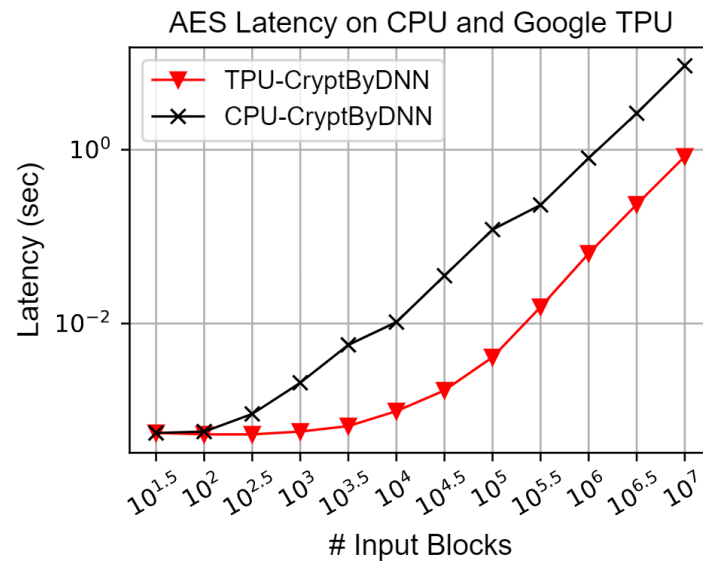


Chacha20 Encryption Latency on GPU



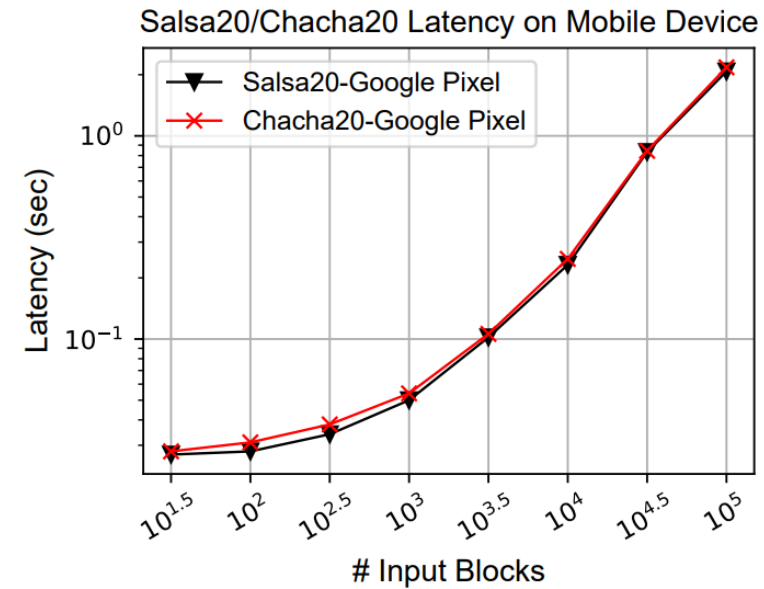
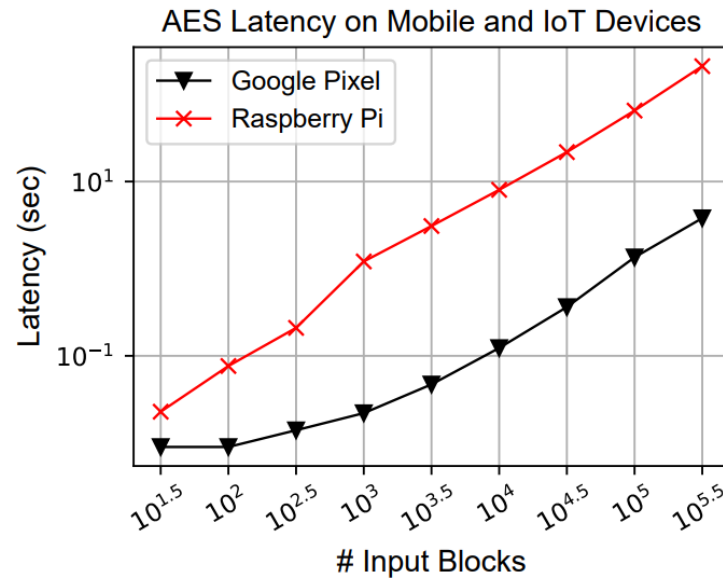
Deployment on Diverse Software and Hardware Stacks

TensorCrypt models on Google TPUs



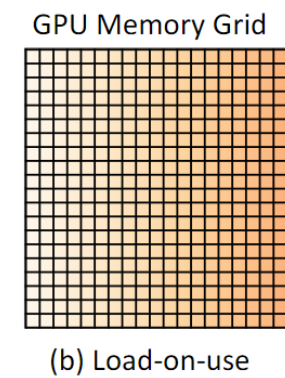
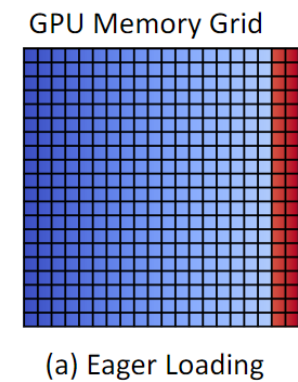
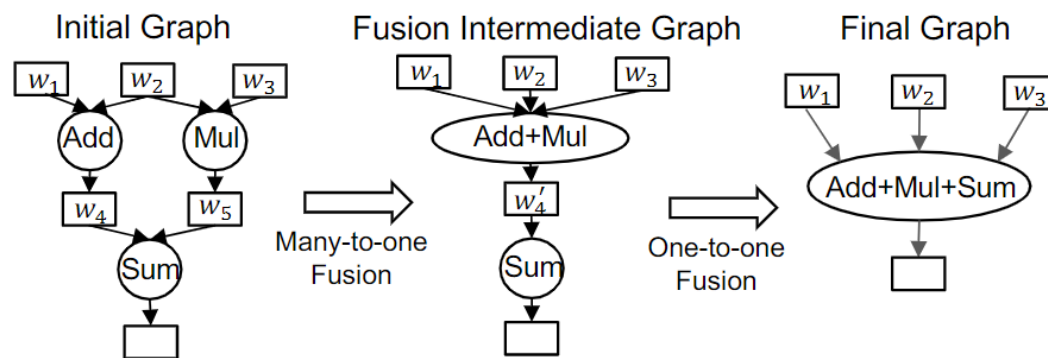
Deployment on Diverse Software and Hardware Stacks

TensorCrypt models on Google Pixel (Smartphone) and Raspberry Pi (IoT)

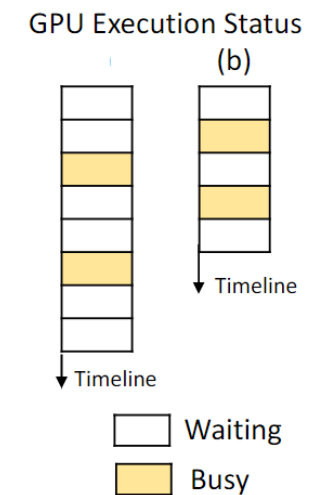


TensorCrypt

- Repurposing neural networks to efficient cryptographic computations
- Program transformation with DSLs and model optimizations
- Advanced encryption and decryption speed with deployment on diverse software and hardware stacks.



Memory Hit Rate: 
Low High



References

- Gueron, Shay. "Intel advanced encryption standard (AES) new instructions set." Intel Corporation 128 (2010).
- Santucci, Pierpaolo, et al. "MemShield: GPU-assisted software memory encryption." International Conference on Applied Cryptography and Network Security. Cham: Springer International Publishing, 2020.
- Kumar, Thanikodi Manoj, et al. "A low area high speed FPGA implementation of AES architecture for cryptography application." *Electronics* 10.16 (2021): 2023.
- Siegelmann, Hava T., and Eduardo D. Sontag. "On the computational power of neural nets." Proceedings of the fifth annual workshop on Computational learning theory. 1992.
- Tezcan, Cihangir. "Optimization of advanced encryption standard on graphics processing units." IEEE Access 9 (2021): 67315-67326.



SecLab

COMPUTER SECURITY LABORATORY



THE OHIO STATE
UNIVERSITY

TensorCrypt: Repurposing Neural Networks for Efficient Cryptographic Computation



University of
Massachusetts
Amherst