

Secure Machine Learning Hardware: Challenges and Progress

Kyungmi Lee[§], *Member, IEEE*, Maitreyi Ashok[§], *Student Member, IEEE*, Saurav Maji, *Member, IEEE*, Rashmi Agrawal, *Member, IEEE*, Ajay Joshi, *Senior Member, IEEE*, Mengjia Yan, *Member, IEEE*, Joel S. Emer, *Fellow, IEEE*, and Anantha P. Chandrakasan, *Fellow, IEEE*

Abstract—With the rising adoption of deep neural networks (DNNs) for commercial and high-stakes applications that process sensitive user data and make critical decisions, security concerns are paramount. An adversary can undermine the confidentiality of user input or a DNN model, mislead a DNN to make wrong predictions, or even render a machine learning application unavailable to valid requests. While security vulnerabilities that enable such exploits can exist across multiple levels of the technology stack that supports machine learning applications, the hardware-level vulnerabilities can be particularly problematic.

In this article, we provide a comprehensive review of the hardware-level vulnerabilities affecting domain-specific DNN inference accelerators and recent progress in secure hardware design to address these. As domain-specific DNN accelerators have a number of differences compared to general-purpose processors and cryptographic accelerators where the hardware-level vulnerabilities have been thoroughly investigated, there are unique challenges and opportunities for secure machine learning hardware. We first categorize the hardware-level vulnerabilities into three scenarios based on an adversary's capability: 1) an adversary can only attack the off-chip components, such as the off-chip DRAM and the data bus, 2) an adversary can directly attack the on-chip structures in a DNN accelerator, and 3) an adversary can insert hardware trojans during the manufacturing and design process. For each category, we survey recent studies on attacks that pose practical security challenges to DNN accelerators. Then, we present recent advances in the defense solutions for DNN accelerators, addressing those security challenges with circuit-, architecture-, and algorithm-level techniques.

Index Terms—Hardware security, DNN accelerators, Side-channel attacks, Fault injection attacks, Memory security, Hardware trojan

I. INTRODUCTION

K. Lee, M. Ashok, M. Yan, and A. P. Chandrakasan are with the Department of Electrical Engineering & Computer Science, Massachusetts Institute of Technology, Cambridge, MA 02139 USA (e-mail: kyungmi@mit.edu, maitreyi@mit.edu, mengjiay@mit.edu, anantha@mit.edu).

S. Maji was with the Department of Electrical Engineering & Computer Science, Massachusetts Institute of Technology, Cambridge, MA 02139 USA.. He is currently with Intel Labs, Hillsboro, OR 97124 USA (majisaurav1@gmail.com).

R. Agrawal is with Boston University, Boston, MA 02215 USA, and Massachusetts Institute of Technology, Cambridge, MA 02139 USA (rashmi23@mit.edu).

A. Joshi is with the Electrical and Computer Engineering Department, Boston University, Boston, MA 02215 USA (joshi@bu.edu).

J. S. Emer is with the Department of Electrical Engineering and Computer Science, Massachusetts Institute of Technology, Cambridge, MA 02139 USA, and also with Nvidia Corporation, Westford, MA 01886 USA. (e-mail: jsemer@mit.edu)

[§]K. Lee and M. Ashok equally contributed to this work.

Manuscript received July 16, 2024; revised Oct 10, 2024.

THE past decade has witnessed the remarkable success of deep neural networks (DNNs) in a wide range of applications and benchmarks. Beyond the success in academia and research, DNNs have been adopted for numerous commercial and real-world applications, such as chatbots [1], medical imaging [2], drug discovery [3], autonomous driving [4], and circuit design [5]. However, with this wide usage, there is a growing concern over the *security* of DNNs, especially when they are deployed for high-stakes applications.

Security has three key aspects – confidentiality, integrity, and availability – each of which has direct connections to machine learning applications (Fig. 1). First, *confidentiality* of user-provided input data is crucial to ensure data privacy. For example, in biomedical applications [2], user input data can contain private health information that should not be accessed by unauthorized users and is protected under government regulations [6]. Also, confidentiality is important for the intellectual property of proprietary DNNs. Developing modern DNNs can require significant resources, such as a proprietary data set, high-performance computing systems, and custom training algorithms [7]. Thus, the model architectures and parameters of DNNs can be important assets for the developers, and they can be the target of an adversary who attempts to create a functional copy of the DNN.

Second, *integrity* of both user input data and DNN model parameters is essential for trustworthy AI applications. In order for a user to trust the prediction made by a DNN, there has to be a guarantee that both user-provided data and DNN model parameters are authentic. Recent work showed that the prediction of DNNs can be dramatically changed by the addition of a small amount of noise in either input data [8], [9] or model parameters [10]–[13], illustrating the importance of integrity for DNNs. Furthermore, errors that are introduced during the computation of a DNN, either random or adversarial, can also undermine the trustworthiness [14]–[17].

Finally, *availability* of machine learning applications can be compromised when an adversary wages denial-of-service attacks. By preventing the accelerator from performing necessary operations, the entire system can be slowed down or completely halted [18], [19]. For latency-critical applications, such as autonomous driving, this can result in real harm from not being able to make decisions on-demand.

These security concerns pose novel challenges for domain-specific accelerators targeting machine learning applications. Crucially, *hardware-level vulnerabilities* can be exploited to undermine the security required for machine learning appli-

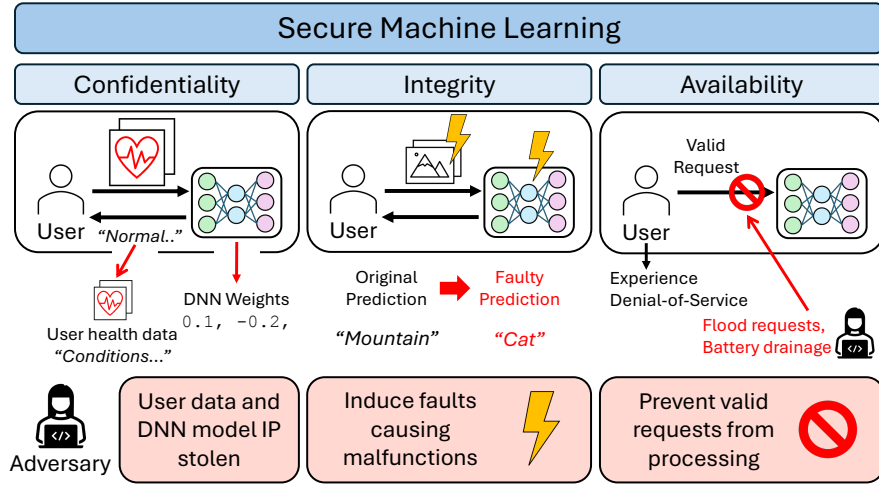


Fig. 1: Three aspects of security are confidentiality, integrity, and availability. Each aspect is crucial for secure machine learning and can be undermined by an adversary, as summarized in this figure.

cations. Hardware-level vulnerabilities often arise from inherent or intended characteristics of hardware designs, such as data-dependent power consumption [20] or electromagnetic radiation [21], interference [22], [23] and remanence [24], [25] properties of memory cells, and resource sharing across multiple tenants and processes [26]–[28]. As they are inherent to hardware designs, not architecture- or implementation-level flaws, these vulnerabilities cannot be simply removed. Furthermore, adversaries can embed hardware trojans in the manufacturing and design process of semiconductor chips [29]–[37], creating backdoor channels and causing malicious behavior without the knowledge of the original designer. These vulnerabilities motivate the need for *secure* DNN accelerators equipped with the appropriate defense solutions.

As such, there is a growing interest in securing DNN accelerators from hardware-level vulnerabilities. This article aims to provide a comprehensive overview of recent progress in secure machine learning hardware, specifically focusing on domain-specific accelerators targeting DNN inference. These accelerators service inference requests on user inputs with trained and fixed DNN model parameters, and all three aspects of security are pertinent to this scenario. In this article, we first categorize the source of hardware-level vulnerabilities in DNN accelerators and analyze the unique characteristics of DNN accelerators that affect hardware security (Sec. II). Then, we survey recent works exploiting hardware-level vulnerabilities to undermine the security of DNN accelerators (Sec. III). Next, we discuss recent progress in secure DNN accelerator designs, covering diverse circuit-, architecture-, and algorithm-level techniques proposed for the security (Sec. IV).

II. UNDERSTANDING THE SOURCES OF VULNERABILITIES

Hardware systems face various security vulnerabilities. The characteristics of underlying circuits, such as currents fluctuating depending on the data being processed [20], [21] and properties of memory bit cells [23], [38], can leak information or allow manipulation by an adversary. Often, an adversary can

insert hardware trojans. These hardware-level security vulnerabilities have been extensively studied for conventional general-purpose microprocessors or cryptographic accelerators.

Emerging domain-specific DNN accelerators [39]–[44] share many of the conventional hardware-level vulnerabilities. However, it is important to recognize the characteristics of DNN accelerators that distinguish the attacks targeting them from conventional hardware attacks. DNN accelerators employ high levels of parallelism for the arithmetic operations and typically have regular and predetermined data orchestration across their memory hierarchies [40], [45]–[47]. In this section, we first describe the sources of hardware-level vulnerabilities for DNN accelerators from the perspective of conventional hardware security (Sec. II-A). Next, we discuss the implications of these unique characteristics of DNN accelerators on the hardware-level vulnerabilities and attacks (Sec. II-B).

While this paper primarily focuses on hardware-level vulnerabilities affecting DNN accelerators, it is crucial to acknowledge vulnerabilities originating from the host processor when DNN accelerators are used as co-processors. A host processor with compromised system software, such as hypervisors or operating systems, can allow an adversary to gain root privileges [48]–[51]. Similarly, hardware attacks targeting the host processor can enable privilege escalation [23], [27], [28], [52]. Since the host processor off-loads DNN workloads to accelerators, an adversary with root privileges can potentially access data originating from the host, without further sophisticated attacks targeting the accelerators we describe in this section. Thus, the security of the host processor and the potential for an adversary to gain root access represent another layer of vulnerabilities affecting the overall system. Although our discussion focuses on hardware-level vulnerabilities specific to DNN accelerators, readers are directed to [23], [27], [28], [48], [51], [53]–[55] for a comprehensive overview of host processor security issues and software-level vulnerabilities.

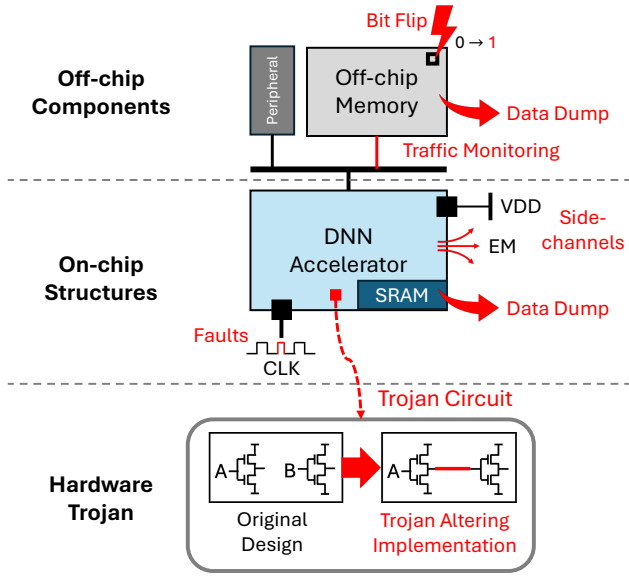


Fig. 2: The sources of hardware-level vulnerabilities affecting DNN accelerators.

A. Conventional Taxonomy for Hardware-level Vulnerabilities

We can categorize the hardware-level vulnerabilities according to the scope of what can be trusted in hardware systems (Fig. 2). In the first scenario, an adversary cannot directly observe or manipulate a victim DNN accelerator, thus the on-chip structures in a victim DNN accelerator can be trusted to be secure. However, the off-chip components, such as the main memory and peripherals that are connected to a victim DNN accelerator, can still be vulnerable to an adversary (Sec. II-A1). Alternatively, there can be a scenario where the on-chip structures are exposed to an adversary, and an adversary can obtain direct information from a victim DNN accelerator (Sec. II-A2).

The above two scenarios assumed that a DNN accelerator is manufactured faithfully according to the original design. This assumption fails when an adversary can insert hardware trojans to the design, such that the behavior of a victim DNN accelerator cannot be trusted (Sec. II-A3).

In the rest of this section, we provide an overview of vulnerabilities for each scenario. These vulnerabilities have been conventionally investigated for general-purpose CPUs/GPUs and cryptographic accelerators. As DNN accelerators share similar security challenges, where appropriate, we will also describe related work on CPUs, GPUs, and cryptographic accelerators that we believe will be applicable to DNN accelerators.

1) *Vulnerabilities of the Off-chip Components:* Since DNNs have a large memory footprint for their model parameters and intermediate tensors, an off-chip memory element such as DRAM acts as the main memory for a DNN accelerator by holding the data that cannot fit on the limited on-chip memory [40], [46], [47]. However, this off-chip memory element and a bus connecting it to the DNN accelerator can be targets of an adversary [24], [25], [38]. Several works [25], [38] showed that an adversary can read out data stored in a DRAM by executing an abnormal booting sequence or

by physically removing a DRAM and connecting it to an adversary-controlled machine. These attacks are often referred to as ‘cold-boot’ attacks, and they exploit the characteristic of a DRAM that data remains for a certain period of time even after the power is disconnected. These attacks undermine the confidentiality of data stored in off-chip memory.

Off-chip memory is also vulnerable to fault injection attacks that undermine integrity [11], [12], [22], [23]. An adversary with physical access to an off-chip memory can use electromagnetic pulses [56] and lasers [57] to induce bit flips in data stored in a DRAM. Furthermore, bit flips can be remotely induced by an adversary who exploits interference between bit cells in a DRAM, which causes one bit cell’s activity to affect its neighboring bit cells [22], [23]. For example, Rowhammer [22] is an attack that induces predictable bit flips in a DRAM by repeatedly accessing the neighboring rows of the victim bit cells, without physical access to a DRAM.

As an off-chip memory is on a separate die and sometimes a separate package from the DNN accelerator, a bus connecting those two can be susceptible to both confidentiality and integrity breaches. An adversary who can probe this bus can monitor and modify the data traffic and the metadata such as addresses and request types (e.g., read or write) for each transaction [58]–[60]. Therefore, even when an adversary lacks the capability to directly attack a victim DNN accelerator, the off-chip components can be vulnerable to attacks undermining confidentiality and integrity.

We briefly note that the emerging trend of 2.5D/3D integration for off-chip memory elements can reduce some of the vulnerabilities associated with traditional off-chip DRAMs [61]–[63]. 3D integration allows stacking memory elements and processors vertically within the same package. From the security perspective, memory elements residing in the same package reduce the attack surface as bus snooping and tampering attacks are not applicable. As such, some threat models for GPUs that use HBMs assume that the on-package HBMs can be considered trusted, similar to the on-chip structures [61], [64], unlike conventional assumptions on off-chip memory elements.

2) *Vulnerabilities of the On-chip Structures:* Data-dependent characteristics of a victim DNN accelerator, such as the power consumption [20], timing information [65], and electromagnetic radiation [21] of an accelerator can act as *side-channels* that leak information. An adversary exploits the fact that these characteristics often have a strong correlation with data and operations of the accelerator. For example, the power consumption of an electronic circuit depends on the activity factor and thus the data being processed, and an adversary who can observe the power consumption of an accelerator can reverse-engineer the data [20]. Similarly, timing information can be exploited if instructions have conditional branches that require different number of cycles to complete an operation (e.g., multiplications, nonlinear functions) depending on the data [66]–[68]. Spatially distributed currents can also induce electromagnetic radiation that can be measured non-invasively [69].

An adversary can learn about confidential data, such as user inputs [70], [71] and model parameters [72]–[75], using

these side-channel leakages. When fine-grained side-channel leakages, such as those at the single instruction level, are available, an adversary can obtain powerful insights into exact data values being processed. However, such fine-grained information can be challenging to be obtained in some scenarios. For example, many computations are performed in parallel across several arithmetic units in GPUs (also similarly in DNN accelerators), and often an adversary can only observe aggregate features, such as memory allocation sizes and hardware counter values [76], [77]. Still, these coarse-grained leakages can pose significant security risks as we elaborate in Sec. III-B1 of this article.

Moreover, an adversary can manipulate the power or clock supply to the device to induce faulty operations in a DNN accelerator [15]–[17]. Clock glitches can induce skipping of instructions in microprocessors [78] or timing violations in combinatorial logic [15]. Glitches in the power supply can also cause timing violations [16], [17] as the timing requirement of a circuit depends on the supply voltage [16]. High-precision electromagnetic pulses and lasers can cause bit flips in the internal state of an accelerator [56], [79]. As faults induced by these physical manipulation methods alter the result of operations, they can undermine the integrity of DNN accelerators.

Fully invasive attacks can also affect the privacy and integrity of these accelerators. Attacks such as voltage micro-probing [80] can allow attackers to find the values stored in specific memory locations or communicated in buses. Furthermore, IC delayering and imaging [81] can give full knowledge of the integrated circuit layout, which has successfully been used to recover the entire design.

Often, these attacks on a DNN accelerator require an adversary to have physical access to an accelerator to measure the side-channel leakages or interfere with the hardware. However, recent work demonstrated these attacks can be waged remotely by leveraging a software interface developed for power monitoring [82], micro-architectural side-channels available from resource sharing [83], or a monitoring circuit that is co-located with a victim DNN accelerator when multiple users share a remote FPGA [16], [17], [84]–[87]. Thus, the security of the on-chip structures is relevant for both edge and cloud environments.

Finally, on-chip memory elements, such as SRAMs that act as buffers and scratchpad memory in DNN accelerators, can be also subjected to data readout attacks similar to the off-chip memory. Recent work from [88] showed that an adversary can exploit the power domain separation in a processor (i.e., SRAMs have a different supply voltage from other on-chip components for energy-efficiency) to induce data retention of SRAMs similar to the cold-boot attacks against the off-chip DRAM. Also, another work [89] demonstrated that a secret stored in the on-chip SRAM can be imprinted to the bit cells, causing it to be leaked. While these exploits are only demonstrated for general-purpose microprocessors, the same methodology can be potentially applied to on-chip SRAMs of DNN accelerators.

3) *Hardware Trojans*: Finally, an adversary can insert hardware trojans, which are circuits that cause malicious effects

on the operation of a DNN accelerator without knowledge of the original designer [29]–[31]. These hardware trojans can operate as backdoor channels to leak information [90], [91], induce bit flips and rewire logic gates to generate faulty outputs [30], [37], and drain the battery of a victim accelerator to make it unavailable to users [92]. Thus, hardware trojans are powerful attacks that can broadly undermine the confidentiality, integrity, and availability of a DNN accelerator.

An adversary can insert hardware trojans at various stages throughout the design and fabrication process of a DNN accelerator. As the semiconductor supply chain gets more complex, with third-party hardware IP blocks and fabrication processes outsourced to foundries, the risk of hardware trojans is also increasing [93], [94]. While normal functional testing after the fabrication of a semiconductor chip exists, a carefully designed hardware trojan can evade such tests by only rarely triggering its malicious behavior [29], [32], [34].

B. Unique Characteristics of DNN Accelerators

While the conventional taxonomy for hardware security provides an excellent primer to understanding various hardware-level vulnerabilities affecting a DNN accelerator, it is important to recognize the unique characteristics of DNN accelerators and their ramifications on the vulnerabilities. These unique characteristics often make exploiting hardware-level vulnerabilities more challenging to an adversary targeting DNN accelerators.

First, DNN accelerators generally use predetermined and structured control across their memory hierarchy and processing elements, without using conditional branches in their instructions [39], [40], [47]. For example, the timing variation observed for floating-point arithmetic operations and ReLU activation functions (i.e., negative values are clipped to zero) in some micro-controllers [68] is not present in DNN accelerators as they can be implemented with a dedicated circuit that completes in a fixed number of cycles. As a result, timing side-channel leakages are limited to coarse-grained patterns, such as the total cycles required to process a single layer in a DNN workload [95], and obtaining fine-grained timing information for each arithmetic operation or data access is challenging. Thus, an adversary has a much more limited opportunity for exploiting timing side channels in DNN accelerators.

Second, DNN accelerators have many processing elements that perform arithmetic operations running in parallel. An adversary observes the mixture of power consumption or electromagnetic radiation patterns from multiple processing elements that compute with different input and weight pairs. This large amount of parallelism can be an obstacle for an adversary trying to exploit side-channel leakages, especially when an adversary does not know the exact architecture and scheduling of operations in the accelerator [73], [74].

Third, as DNNs have a large memory footprint, the amount of secrets that an adversary has to recover from DNN accelerators is much larger than the 128-bit or 256-bit secret keys in cryptographic accelerators [20], [96]. Combined with the two above-mentioned characteristics that limit exploitable side-channel leakages in DNN accelerators, it is practically

challenging to fully recover the fine-grained information of a victim DNN accelerator. Instead, many attacks on DNN accelerators target high-level features such as the model architecture of a victim DNN instead of every single weight parameter [95], [97]–[99], and the class information of user input instead of the exact input features [100]. These high-level features still present significant security challenges. DNN model architectures are often core intellectual property for developers, while input class information can be sensitive, particularly in biomedical applications where a DNN classification might involve detecting health conditions or diseases, making the class information itself personal health information that has to be protected.

Fourth, the robustness of DNNs to small random perturbations can increase the difficulty of fault injection attacks. In cryptographic accelerators, even a single bit flip in a secret key has an avalanche effect across all operations using the perturbed key [101], [102]. Often, a bit flip in the most significant bit (MSB) of the exponent field in a floating-point number can also have detrimental impacts on DNNs [12]. However, in general, the performance of a DNN is not significantly affected by small random noise added to its model parameters or input data [8], [10], [103], which is a characteristic often leveraged to design efficient hardware accelerators using approximate [104], [105] or analog computing [106]–[109]. Therefore, an adversary must algorithmically identify the worst-case faults for its victim DNN, as random faults can be less effective [10], [13].

Finally, an adversary's prior knowledge about a victim DNN has to be carefully considered. Cryptographic accelerators usually implement a standardized cipher (i.e., AES [110]) whose algorithm is publicly known. Thus, the adversary can compute the expected results given known inputs to perform statistical side-channel analysis [20]. However, DNNs have a wide variety of model architectures [111]–[114], numerical precision [115]–[119], and sparsity patterns [120]–[123], and there is no single 'standard' DNN algorithm. Thus, an adversary who has no knowledge about a victim DNN (i.e., 'black-box' attack) has to identify this high-level information before attempting to recover the model parameters [70], [71], [73]. Also, some fault-injection attacks require full knowledge of a victim DNN, including its model parameters (i.e., 'white-box' attack) [10], [13], [124]. Therefore, an adversary's prior knowledge of a victim DNN becomes a key component of a threat model.

III. RECENT ADVANCES IN ATTACKS TARGETING DNN ACCELERATORS

In this section, we provide a comprehensive survey of recent advances in attack methodologies exploiting hardware-level vulnerabilities of DNN accelerators. For each category of vulnerabilities introduced in Sec. II-A, we outline how attacks can undermine confidentiality, integrity, and availability aspects of DNN applications. These advances demonstrate that hardware-level vulnerabilities can be a significant threat to DNN accelerators.

A. Attacks on the Off-chip Memory and Bus

1) *Confidentiality Attacks*: Cold-boot attacks [125], [126] are a well-known methodology for extracting the data stored in an off-chip memory element. When a DNN accelerator uses an off-chip memory as its main memory, this attack methodology can be applied to steal a victim DNN's model architecture and parameters. A recent study [25] performed cold-boot attacks on a commercially available edge DNN accelerator. The attack targets the RAM on the host processor using low temperatures (utilizing the fact that data remains for a longer period of time when the temperature is low) and the abnormal booting sequence of the host processor. In this attack demonstration, key prior knowledge required for an adversary is the storage format of a victim DNN's model architecture and parameters, which can be publicly available if a DNN accelerator uses an open-source protocol to interface with the host processor [127]. Then, an adversary can leverage this knowledge to determine bits corresponding to the model architecture and the parameters, and correct bit errors to reconstruct a functional DNN.

Furthermore, the metadata of off-chip memory traffic, such as the requested addresses, types (i.e., read or write), and timestamps, can act as a side-channel that an adversary can use to reverse engineer a victim DNN's model architecture [58]–[60]. Recall that DNN accelerators use structured and pre-determined control for accessing their off-chip memories. While this characteristic can reduce fine-grained timing side-channel leakages (Sec. II-B), the off-chip memory access pattern still leaks relevant coarse-grained information about a victim DNN. For example, one layer's output tensor is used as the next layer's input tensor in many DNNs, and identifying the read-after-write dependencies across the off-chip memory access patterns can reveal the intermediate tensor size and the timing 'boundary' of when one layer starts and ends [58]. Combined with another property of a DNN that the model parameters are usually read-only during the inference, [58], [60] showed that reverse engineering of a DNN's model architecture is feasible with only the memory traffic patterns.

2) *Integrity Attacks: Fault Injections*: DNNs are generally robust to small amounts of random noise, with some exceptions on the exponent field of a floating-point number [12]. Thus, fault injection attacks have to be more precise in terms of which bits to target (Sec. II-B). Recent studies proposed approaches to identify the most harmful bit flips in the model parameters of a DNN [10], [13], [124] and showed that only a handful of bit flips is sufficient to completely degrade the performance of a victim DNN. [10] used the gradient information of each weight parameter to determine the most important bits in a linearly quantized DNN. [10] chose the bits with a large magnitude of gradients and showed that only 10-20 bit flips are required to degrade the performance of a victim DNN to the random guess level. [13] leveraged the characteristics of sparse matrices and the compressed storage formats to attack pruned DNNs with fine-grained sparsity in their weights. In particular, [13] exploited that the compressed storage formats for sparse matrices separately store the nonzero values and their positions in the original matrices, and that a bit flip in the positions leads

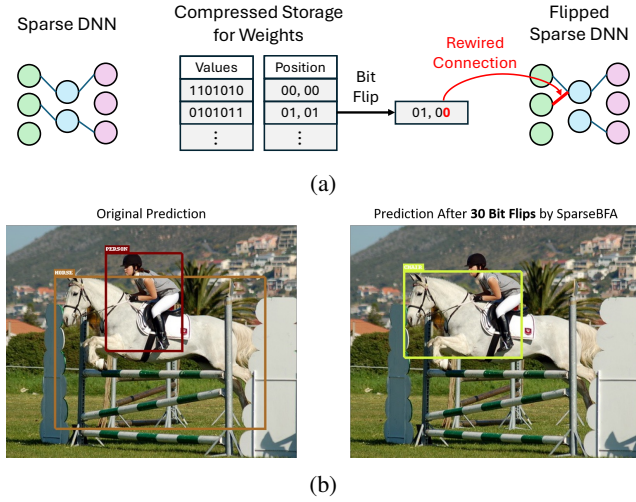


Fig. 3: (a) In [13], the position of nonzero values in a sparse matrix is subjected to bit flip attacks, resulting in rewiring of connections in a victim DNN. (b) The original object detection model correctly predicts that there are a human and a horse in an input sample. However, after flipping 30 bits using [13], the model wrongly predicts that there is a chair in this example. This image is reproduced from [13] (© IEEE 2022).

to rewiring of the connections while not changing the nonzero value directly (Fig. 3). [13] also showed that flipping only a small percentage of the total memory footprint, less than 0.00005%, of the compressed model parameters is sufficient to cause significant performance degradation as illustrated in Fig. 3b. Note that these algorithms do require an adversary to have full knowledge of a victim DNN, such as the model architecture, parameters, and storage format.

Using these algorithms to identify a few worst-case bit flip targets, an adversary can use fault injection methodologies against the off-chip main memory of a DNN accelerator. Recent works [11], [12] showed that Rowhammer can be practically utilized to induce bit flips in the target bit cells in a DRAM storing the model parameters of a victim DNN. For example, in [11], flipping ~ 20 bits took ≤ 100 seconds, demonstrating that Rowhammer can be a practical threat for the off-chip DRAM. Finally, the faults in an off-chip memory element can be used to induce targeted misclassification only for certain inputs, making the attack more stealthy [124].

These integrity attacks are particularly challenging for inference accelerators. Inference accelerators typically do not update the parameters of DNNs they are servicing, and corrupted bits in the parameters can impact all subsequent inference requests if there is no mechanism to detect these attacks. In contrast, DNN *training* accelerators that update the parameters can often recover from corrupted bits, and the impact of integrity attacks can be less damaging [128].

B. Attacks on the On-chip Structures of a DNN Accelerator

1) Confidentiality Attacks: Side-channel Attacks (SCAs):

As we discussed in Sec. II-B, timing side-channels are limited for DNN accelerators that do not have data-dependent

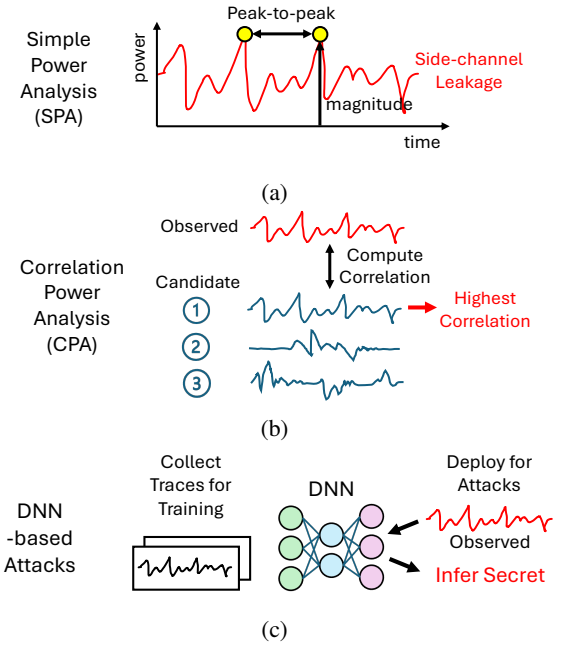


Fig. 4: Physical side-channel attacks analyze the physical-layer leakages of a victim DNN accelerator, such as its power traces or electromagnetic radiations. An adversary can use statistics and pattern recognition algorithms to reverse-engineer secrets from the leakages, such as (a) a Simple Power Analysis (SPA), (b) a Correlation Power Analysis (CPA), and (c) a DNN.

branch instructions. As such, many works on side-channel attacks (SCAs) on DNN accelerators focused on other types of physical side-channel leakages, such as power consumption and electromagnetic radiation from a victim DNN accelerator. After an adversary collects those traces, an adversary uses statistics or pattern recognition algorithms to reverse-engineer secrets from those traces (Fig. 4). For example, a Simple Power Analysis (SPA) passively monitors the magnitude and peaks in the traces to infer high-level secrets like the DNN model architecture [95], [97]–[99], [129], [130] (Fig. 4a). Another type of statistical analysis is a Correlation Power Analysis (CPA) [131], [132] (Fig. 4b). For a CPA attack, an adversary algorithmically computes the power model of a victim circuit for several candidate values for the secret. Then, an adversary correlates the actual side-channel traces with an algorithmic model using statistical analysis and chooses the candidate value that has the highest correlation. Compared to a SPA attack, a CPA attack generally requires more information for an adversary, such as the inputs given to a victim circuit to generate the traces.

Recently, some SCAs leveraged powerful pattern recognition capabilities of DNNs to strengthen their attacks [133]–[136] (Fig. 4c). For example, [133] used deep learning to predict the key of AES cryptographic implementations with a much smaller attack time compared to traditional attack methods like CPA attacks. Furthermore, this was applicable across multiple devices, where the power SCA variation between devices is much higher than that between different key values on the same device. Similarly, [134] used neural networks

to predict the data being converted in an analog to digital converter (ADC) based on power SCA data. Both multilayer perceptrons and convolutional neural networks (CNNs) were highly successful at learning the relationship between time-series spikes in power consumption from capacitors charging and discharging to the output digital bits, with over 99% accuracy for all 12 bits of an unprotected commercial ADC. In [134], CNNs could even reverse-engineer some information from an ADC equipped with protection mechanisms.

These DNN-methods use a profiled form of attack [137], where training data is required with the SCA traces corresponding to a particular known computation rather than just a simplified power model as in conventional SPA and CPA attacks. On the other hand, much less data is required during the actual attack [138], which can even be done in real-time without having active control over an input to repeat the operation enough times for successful statistical analysis as in CPA attacks [131], [132]. This tradeoff is valuable for an attacker, who has the ability to buy a device that is nominally identical to the target through regular commercial channels, and train a model that can apply with high accuracy even with process variations.

While physical SCAs have been well studied in cryptographic circuits [139]–[142], their presence in DNN accelerators [143]–[145] has only been considered in the last decade. In this subsection, we introduce recent advances in physical SCAs against DNN accelerators, which aim to reverse-engineer user input data, DNN model parameters, and architectures.

Recovering User Input Data Recent studies [70], [71] demonstrated the use of physical SCAs to recover user input data for simple image processing workloads such as MNIST [146]. Such attacks can be particularly relevant in machine learning systems where an on-chip sensor and ADC generates the inputs for the accelerator, rather than being transmitted from off-chip. They leverage the first layer of a DNN that directly operates on user input data and note any significant changes of the power consumption during the processing of this first layer. These changes correlate to large differences between nearby pixels in a user input, which reveals the silhouette of the image. Note that an adversary requires the knowledge of a victim DNN's model architecture to determine when an accelerator processes the first layer, but not the model parameters. This attack is currently limited to relatively simple user input data, such as MNIST or medical images that have a simple image on top of a plain background [70]. However, the success rate of this attack has been shown to be higher if an adversary can actively collect power traces from a DNN accelerator with known input data to build a template, which is then used to reconstruct secret user input data [70]. Furthermore, similar attacks can be applied to other layers of a DNN. While the intermediate activations might not directly leak the user input, they can be leveraged to infer the DNN model parameters.

Recovering Model Parameters The next type of attack attempts to recover the DNN model parameters [72]–[75]. The threat model for these attacks assume that a victim DNN's model architecture is known to an adversary, and an adversary

can wage CPA attacks with the known input data to a victim DNN accelerator.

In [73], a proof-of-concept CPA attack on a DNN accelerator using a spatial architecture using an array of processing elements [44] was demonstrated. This work computed the correlation after observing a chain of sequential computations in one processing element within the array, which improves the attack accuracy. However, this attack was demonstrated on a relatively simple accelerator architecture compared to commercial accelerators. Perhaps most importantly, [74] is the first to consider realistic hardware architectures, by attempting to perform SCA on a commercially available edge DNN accelerator, whose accelerator architecture is unknown. While this attack is not entirely successful even on very basic neural network structures, it reveals a real threat that could be further exploited with more complex analysis methods. In summary, many of these demonstrated attacks have difficulties in recovering inputs and model parameters with high accuracy without very limiting assumptions.

Recovering Model Architecture Instead, many recent works focused solely on recovering the DNN model architecture [95], [97]–[99]. Unlike CPA attacks described above, these attacks utilize SPA attacks that only require passive SCA measurements without any control over inputs and weights, and can even be performed remotely [98], [99] or on larger scale commercial accelerators [95], [99].

For example, [97] used electromagnetic radiation measurements to recover the number of layers and the number of parameters in each layer in a victim DNN. More recent works [95], [98], [99] trained machine learning models that infer a victim DNN's model architecture, including the number of layers and tensor shapes in each layer, from the side-channel measurements. While these attacks do not have 100% accuracy, they can be combined with prior knowledge of typical architecture hyperparameters (e.g., for image processing applications, convolutional neural networks with 2-dimensional kernels are typically used) and some fine tuning [97], [98] to achieve similar accuracy to the original network. Finally, we note that these attacks on recovering the model architecture can be combined with the above-mentioned attacks on recovering the model parameters to reconstruct the entire DNN model without any prior knowledge [147].

In-memory computing (IMC) So far, we discussed physical SCAs affecting DNN accelerators with conventional architectures, which have computing processing elements that interact with each other and a separate central bank of memory. However, DNN accelerators using the emerging in-memory computing (IMC) technique [148], [149], also referred to as compute in memory (CIM) have been proposed. These have unique sources of SCA vulnerabilities [150]–[152], as discussed in this section.

First, IMC accelerators depend on intrinsic analog properties of SRAM or non-volatile memory (NVM) bit cells to perform multiply-and-accumulate operations. This is in contrast to near memory compute accelerators, where the memory is placed close to the compute through methods such as 3D stacking [153]. While the computation itself is in the analog domain, IMC accelerators can also have side-channel leakages from

the additional required peripheral analog-to-digital converter (ADC) components [134].

Second, unlike logic gates, the memory bit cells in IMC accelerators have a regular structure, which makes the electro-magnetic side-channel leakages easier to obtain [154]. Also, this regular structure can easily reveal which regions of the memory are operating for a specific DNN layer, and prior work showed that the model architecture can be reverse-engineered using laser imaging of an IMC accelerator [155], [156].

Finally, the physical and circuit properties specific to the memory cells can be exploited to recover the data values. For example, the difference in currents when reading and writing 0s and 1s to the memory bit cells can be significantly large for some NVM types [157]–[159], allowing an adversary to easily distinguish the data values. Another example is memristor devices, where leakage currents of cells depend on the nearby stored data values [160].

2) *Confidentiality Attacks: Invasive Readout:* Aside from observing side-channel leakages, invasive attacks that directly tamper with the packaged devices can cause harm to data confidentiality. Voltage micro-probing can be used to directly tap memory cells or buses on the top metal layers of the IC after decapsulating the circuit [80], [161]. This is particularly harmful for non-volatile memory based accelerators which will hold the model parameters even after powering off for the packaging removal process [162]. Furthermore, the entire design of an accelerator can be reconstructed by an attacker to either gain knowledge of the architecture or to produce counterfeit products. This is done through successive delay-er and imaging of the IC, after which image processing tools reconstruct the netlist from GDS information [81], [163]–[165].

3) *Integrity Attacks: Fault Injections:* Fault injection attacks on DNN accelerators can cause malfunctioning of a victim DNN [15]–[17]. For example, [15] used clock glitches to induce timing violation in the logic of a DNN accelerator implemented with an FPGA. Even without prior knowledge of a victim DNN model, they showed that a small amount of glitches (e.g., affecting 0.1% of clock cycles) can significantly degrade the accuracy. Other works [16], [17] further demonstrated the possibility of remote power glitch attacks for a cloud FPGA that hosts multiple users. When an adversary can co-locate the malicious circuit with a victim DNN accelerator, remote power glitches can induce either random faults or duplication faults, both resulting in erroneous computation results [16], [17]. Other modalities such as laser fault injection can also be used for targeted attacks on portions of the circuitry [166].

Another potential attack model leverages physical fault injection attacks to reverse-engineer DNN model parameters [167]. When an adversary has prior knowledge of a victim DNN's model architecture and can observe the computation results of an accelerator, [167] showed that an adversary can compare the computation result using the original input and the fault injected data to obtain model parameters. Therefore, while integrity is the primary concern for fault injection attacks, confidentiality can also be undermined.

C. Inserting Hardware Trojans

Hardware trojans are circuits added to cause malicious effects on the operation without the knowledge of the original designer and can be inserted in DNN accelerators to target various operations (Fig. 5). They generally contain a trigger, which waits for some rare event to start the harmful operation while escaping detection during normal functional testing. In addition, the payload can have a variety of harmful effects like leaking information, draining the battery, or rewiring logic gates. As opposed to neural trojans that are implemented at a software level on the trained model, here we survey trojans specific to the hardware implementation.

1) *Harmful Payloads Affecting Confidentiality, Integrity, and Availability :* Many attacks [29]–[31] implement targeted changes to the weight parameter value at the circuit level similar to the bit flip attacks from Rowhammer [168]. Others have similar effects to changing weights but instead target the computation logic circuitry or the input/output of the block [32]–[35]. All of these can severely affect the integrity of the accelerator, especially if the model is known and can be used to create strong adversarial attacks. Other targets of the trojan that can cause integrity or availability issues include intentional glitches in the clock [36] to cause timing violations or reconfiguring the inter-processing element connections of the network-on-chip [37] to change the network structure.

Furthermore, trojan payloads can be targeted to alternate architectures such as IMC, where the analog properties can be utilized to cause failures. For example, internal voltages can be modified to break the sense amplifier margin, or transient bounces on the ground or power rail can cause read and write failures [169], [170]. In other cases, the analog IMC structure can be exploited to affect confidentiality, by directly relating the power consumption of large blocks such as the ADC to secret weight parameter information [171]. In addition, these payloads can target peripheral portions of the accelerator such as the memory controller, where the layer output activations being written back to off-chip DRAM can be modified [172].

2) *Triggers :* For all these payloads, there needs to be a method to determine when the trojan should start its operation. While some of the trojans are always on [31], [33], this is rare since the RTL IP and post-fabrication functionality tests are likely to see the effect of any such attack. These specific works get past this issue by having a limited payload that only targets certain input images. Otherwise, many works focus on a trigger from specific neural network inputs, since we assume the attacker has some control over this. Thus, they utilize a specific rare data value or input image that is unlikely to occur during standard testing and can be accurately detected but still does not seem obviously malicious [29], [30], [36], [37], [169], [172]. This undetectability during testing can be strengthened by relying on a sequence of such rare inputs [32]. At the same time, other works generate a trojan trigger based on accelerator control signals. For example, while typical accelerators will cycle between different addresses and have a standardized memory access pattern based on the chosen dataflow, some triggers can depend on the same address being accessed many times [34], [170]. While digital accelerators would need to detect this using dedicated circuitry in the timing control logic,

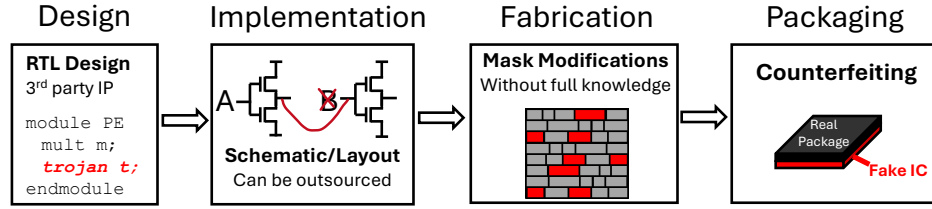


Fig. 5: Steps of DNN accelerator production where hardware trojan insertion can occur.

analog IMC adds further motivation for such attacks. Since RRAM resistances drift with many reads, the effect of repeated accesses can be directly read out as a change in delay or voltage of the readout [170].

IV. RECENT ADVANCES IN DESIGNING SECURE DNN ACCELERATORS

In this section, we present recent progress in defending DNN accelerators from various hardware-level exploits discussed in Sec. III. We present protections that provide security for the off-chip memory (Sec. IV-A), counter physical side-channel attacks (Sec. IV-B) and detect fault injections (Sec. IV-C) on the on-chip structures, and detect hardware trojans (Sec. IV-E). In addition, we present defense solutions for emerging in-memory computing accelerators (Sec. IV-D), which have their unique challenges compared to conventional accelerator architectures. Finally, we present a brief description of Fully Homomorphic Encryption (FHE) technology (Sec. IV-F), which ensures the privacy of both model parameters and input data by performing computations on encrypted data, without having to trust any hardware components.

A. Providing Security for the Off-chip Memory

1) *Authenticated Encryption for the Off-chip Memory*: One approach to defending an off-chip memory element from data readout attacks (e.g., cold-boot attack [125], [126] discussed in Sec. III-A1) and fault injection attacks (Sec. III-A2) is to *encrypt and authenticate* all data stored in an off-chip memory element using cryptographic primitives (Fig. 6). This approach is adopted for memory security in a trusted execution environment (TEE) for general-purpose CPUs [55], [173]–[179] and GPUs [61], which provides a secure ‘enclave’ in the memory. Encrypting the off-chip data using strong cryptographic primitives, such as AES [110], prevents an adversary from recovering the actual plaintext data even if data readout attacks successfully extract the data. Furthermore, authenticating all data transfers with cryptographic hashes enables integrity verification to detect any bit flips in the off-chip data, since the computed hashes for perturbed data will not match the original values [173]. Adding the ‘timestamp’ information (also known as ‘counters’), such as the number of updates performed on a certain address, when generating cryptographic hashes also prevents an adversary from ‘replaying’ the old data and hashes [173]. Note that the security of this authenticated encryption approach depends on the confidentiality of a secret key used for cryptographic operations, which can be stored on a special register on-chip.

Recent studies investigated adapting TEEs for DNN accelerators [180]–[184]. The structured and predetermined data access patterns of DNN accelerators are leveraged to reduce the overhead of supporting a TEE in these works [180], [182]. A key observation was that managing the timestamps, which used to be a major source of performance overhead in CPUs, can be simplified with this structured data access pattern. As a result, the major overhead of this approach reduces to cryptographic encryption and authentication operations required for all off-chip data. In addition, recent work from [184] proposed a software-defined approach to tailor a TEE for diverse DNN accelerator deployment conditions.

For authentication, a cryptographic hash is computed for a block of data and all data in this authentication block is needed for integrity verification. This can lead to additional off-chip data traffic when we only need a portion of data in one authentication block. In DNN accelerators, this property introduces a challenge when authentication blocks misalign with ‘tiles’ of data tensors, where a tile is a basic granularity of data movement in DNN accelerators chosen to optimize the data reuse and performance. All data in authentication blocks need to be fetched although they do not belong to a tile a DNN accelerator requested, causing performance degradation due to the additional off-chip data traffic and cryptographic operations.

In particular, [183] pointed out that cross-layer dependency in a DNN (i.e., one layer’s output is the next layer’s input) complicates the optimal authentication block assignment that minimizes the additional off-chip traffic. Conventionally, tile assignment is typically done independently for each layer. Consequently, the same tensor data can have different tile assignments when it is used as an output tensor of one layer (e.g., 1×3 tiles in Fig. 7a) and an input tensor of the next layer (e.g., 2×2 tiles in Fig. 7a). Suppose we use each tile in an output tensor as an authentication block when we compute the first layer. Then, when the accelerator requests one tile in an input tensor to be fetched for processing the next layer, two output tensor tiles need to be fetched due to authentication, incurring a large amount of redundant reads (Fig. 7b).

To overcome this challenge, [183] proposed an algorithm that determines the optimal authentication block and tile assignment for DNN accelerators. In order to identify the optimal authentication block assignment, [183] applied an exhaustive search and analytically computed the amount of the additional off-chip traffic for each assignment. [183] showed that even for the same DNN accelerator architecture, the performance can be improved by up to 33% with the optimal

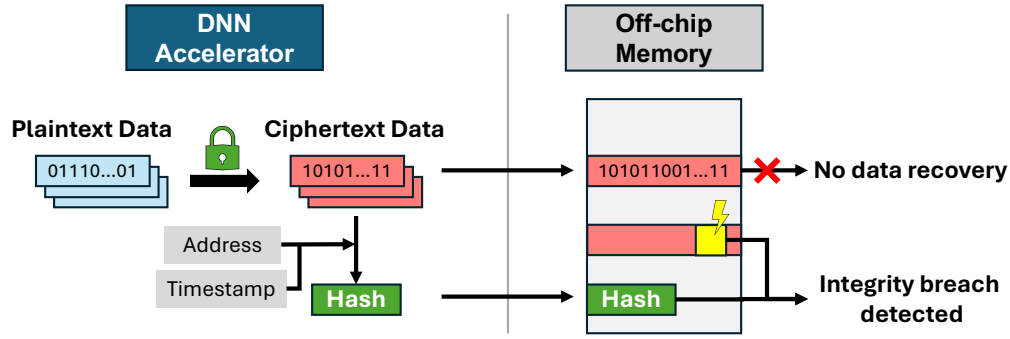


Fig. 6: Authenticated encryption can provide confidentiality and integrity of the off-chip data. The data is always encrypted before being written to the off-chip memory, and a cryptographic hash is assigned for a block of data for integrity verification.

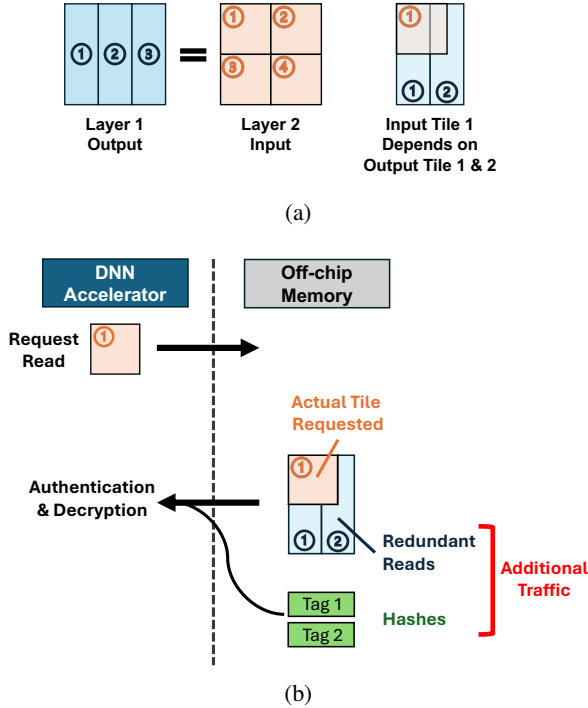


Fig. 7: Cross-layer dependency in a DNN complicates the authentication block assignment.

assignment. In summary, several works [180]–[184] proposed techniques to reduce the overhead of encrypting and authenticating all off-chip data, which provides a strong hardware-level security guarantee.

Comparison with GPU TEEs TEEs for GPUs and DNN accelerators share some key similarities in their security requirements [61], [64]. Both must establish secure communication channels with their host processors, even when the host processor’s system software can be compromised. Also, they need to protect against bus snooping and tampering when they are connected to the host processor using PCIe.

There can be differences regarding the assumptions of off-chip memory security. GPUs often integrate off-chip memory elements within the same package using vertical integration, and some work on TEEs for GPUs consider this on-package memory to be trusted [61]. However, this assumption cannot

be shared among all GPUs, when many consumer-grade GPUs opt for conventional DRAMs on separate packages. As such, other work on GPU TEEs [185] considered hardware-level off-chip memory protection, similar to CPUs and DNN accelerators. Although DNN accelerators can benefit from the vertical integration of off-chip memory in the future, especially for high-performance datacenters [44], DNN accelerators also choose to use separate off-package DRAMs and low-power DRAMs, which require TEEs to provide off-chip memory security.

Finally, DNN accelerators are still evolving, with recent active research into multi-tenancy and virtualization [184], [186], and TEEs for DNN accelerators may need to address new challenges for memory management in the future. For example, TEEs for GPUs require complex virtual-physical address translation to ensure memory allocated for a secure process is isolated [61]. While current DNN accelerators use simple memory control based on regular patterns, future accelerator designs might have to adopt more complex memory management under TEEs.

2) *Alternative Approaches to Memory Encryption:* Instead of encrypting and authenticating all off-chip data, alternative approaches have been proposed to selectively encrypt a few important weights [187]–[189]. While the rest of the model parameters and intermediate tensors that are not encrypted are vulnerable to an adversary, a few encrypted values prevent an adversary from reconstructing a functional DNN with comparable performance to the original victim DNN.

Another approach is to shuffle elements in weight tensors of a DNN and use shuffling information as a secret to provide the confidentiality of DNN model parameters [190]–[192]. To reduce the overhead of decrypting the shuffled data, shuffling can be performed coarsely by swapping rows and columns in a matrix [191], [192] or selectively applied to important layers [190]. These approaches benefit applications requiring the confidentiality of DNN model parameters with low-cost protection, although it does not provide the confidentiality of user input data or protection against fault injection attacks.

Finally, [193] proposed a low overhead authentication technique using algorithmic/parity checksums [194]–[196]. Compared to cryptographic hashes, algorithmic/parity checksums can have low computational and storage overhead. For ex-

ample, [194] set the sum of weight parameters in a row as a checksum, without using more sophisticated cryptographic operations. [195] incorporated parity-based checksums where the least significant bit of each model parameter acts as a parity bit. This approach eliminates the storage overhead for hashes, while its impact on the performance of a DNN is negligible. Similarly, [196] used a 2-bit algorithmic checksum for each group of model parameters. Conventionally, algorithmic and parity checksums [194]–[196] can be susceptible to an adversary who can flip multiple bits or flip the checksum bits. [193] addressed this susceptibility by encrypting the model parameters along with a parity-based checksum [195], reducing the probability of an adversary successfully counterfeiting checksums. Note that the overhead of these checksums also depends on the precision used for DNN inference. For example, the overhead of a single- or 2-bit checksum can potentially become more expensive for lower precision weight parameters and inputs.

3) *Hiding Memory Traffic Patterns*: Recall that an adversary can reverse-engineer a victim DNN’s model architecture, such as the number of layers and tensor shapes in each layer, by observing the metadata like addresses and request types of the off-chip traffic [58]–[60] (Sec. III-A1). Since the metadata is still visible to an adversary even when all off-chip data is encrypted, different protection mechanisms are needed for this attack. In general-purpose processors, Oblivious RAM [197], [198] was proposed to hide the data access pattern of a program. While this technique provides a theoretical guarantee that the resulting data access pattern is independent of the true pattern, it incurs a large overhead in the off-chip access latency and bandwidth requirement [197], [198].

In DNN accelerators, recent works [191], [199], [200] instead proposed more lightweight defenses that aim to obfuscate the read-after-write dependency of the intermediate tensors and the total number of data accesses that reveal the tensor sizes, instead of completely anonymizing the data access pattern as in Oblivious RAM. For example, [199] shuffles all data accesses around layer boundaries (i.e., when a DNN accelerator completes processing one layer and moves on to the next layer) so that an adversary cannot observe the actual start and end of each layer. [191] breaks a layer in a DNN into multiple smaller sub-layers such that an adversary is misled to observe a wrong number of layers in a DNN. [191] also proposes a mechanism to reduce or increase the total number of accesses, by either partially caching the intermediate tensors (i.e., the cached data does not have to be written and read back from the off-chip memory) or issuing dummy requests.

Another approach is entirely algorithm dependent, where an obfuscated DNN that preserves the functionality of the original DNN but with a different model architecture is designed [201]. [201] used a careful search algorithm to ensure that the obfuscated DNN architecture has a small hardware performance overhead (e.g., the total number of computations). Overall, these works [191], [199]–[201] showed that an adversary is misled to the wrong DNN model architecture and ends up with a reconstructed DNN with worse performance than the original victim DNN.

4) *Limitations*: There are many potential research topics related to the memory security of DNN accelerators. First, the performance of a DNN accelerator using memory encryption and authentication (Sec. IV-A1) can be throttled by cryptographic operations required for all data traffic. While [180]–[184] proposed several techniques to minimize the overhead and optimize the trade-off considering cryptographic operations, scaling this approach to future memory technologies offering a substantially higher bandwidth can be challenging. For example, having multiple cryptographic accelerators running in parallel to match a high off-chip memory bandwidth can prevent the performance slowdown at a significant cost for area and energy. As DNNs can be memory-intensive and memory technologies are rapidly improving, scaling the defense solution is an important challenge.

Second, many approaches with non-cryptographic defenses often lack the theoretical guarantee of security. Although they have a low hardware overhead for defenses, such a benefit should be understood in the context of a trade-off with the security level. While the empirical results can provide a practical demonstration of security [194]–[196], [202]–[204], a well-principled analysis on those approaches will be useful for future research.

B. Physical-layer Side-channel Attack (SCA) Protections

As physical SCAs broadly affect the confidentiality of DNN accelerators (Sec. III-B1), there is a growing interest in defending DNN accelerators from physical SCAs. The core idea for the defense is to decorrelate the circuit currents from the data so that side-channel leakages (e.g., power consumption, electromagnetic radiation) do not contain any information about the data being processed. There are two proposed approaches to achieve this decorrelation, masking [75], [205]–[209] and shuffling [209], [210].

1) *Masking*: Masking splits up data into multiple shares, each of which is independently unrelated to the original value [211]–[213]. Here we denote the original value as $\bar{x}$, and n independent shares as $x^1, x^2, \dots, x^n$. In order to recover the original value, all n shares are needed, and an adversary who only knows $\leq n - 1$ shares cannot recover the true value. Computing on these shares is performed in such a way that every output value generated during the computation is independent of at least one input share, so the power consumption for each function is unrelated to the original data $\bar{x}$.

There are two sharing techniques investigated in recent works. First, arithmetic masking [214] uses modular summation to split up each data value, where K is a modulus:

$$x^1 + x^2 + \dots + x^n \bmod K \equiv \bar{x} \quad (1)$$

Note that simple linear arithmetic operations like additions can be performed over each share separately without having to combine them in intermediate steps.

Alternatively, Boolean masking [211], [212], [215] uses bitwise exclusive OR (XOR) operations:

$$x^1 \oplus x^2 \dots \oplus x^n = \bar{x} \quad (2)$$

Unlike arithmetic sharing, Boolean masking is more geared towards operations requiring bit-wise manipulations, such as comparing sign bits. However, Boolean masking can result in more complicated implementations for arithmetic operations [208], [209], since each bit-level logic gate has to be modified.

Recent work adopted Boolean sharing or chose to use a combination of techniques for different functions (e.g., arithmetic masking for multiply-and-accumulate operations but Boolean masking for sign bit comparison) [205], [207]. While the latter can be more optimized for each function, it requires secure conversion between the two types of sharing and incurs additional cost [216], [217].

Recently, the first SCA defense for a binary neural network accelerator was proposed, which does not require multiplications and only needs to support additions [75]. They used arithmetic masking to split input tensors into two independent shares, which are passed to two adder trees separately to compute the results for each share. However, recall that arithmetic sharing involves modular operations (Eq. (1)), while additions required for DNN computations are standard arithmetic. This difference results in the leakage of the sign bit in [75]’s approach, and [75] augmented arithmetic sharing with a circuit-level technique that equalizes the power consumption regardless of the data for sign bit computations. Furthermore, the ReLU activation function that clips negative values to zero requires comparing the sign bits of data, which can be better supported using Boolean masking. Thus, [75] alternates between arithmetic masking for additions and Boolean masking for the ReLU function, with an additional overhead for secure conversion.

More recent work [205]–[207] tried to push the limitations of [75]. [205], [206] modified a DNN algorithm to use modular arithmetic for computations such that arithmetic sharing can be directly applied without the leakage issue, but they require careful selection of the modulus K and significant algorithm modifications. [207] further extended the arithmetic masking solutions to multi-bit fixed point precision, thus improving the practicality of SCA defenses.

Other recent studies instead adopted Boolean masking for all computations, including multiply-and-accumulate operations, in DNN accelerators [208], [209], [218]. However, as we discussed earlier, Boolean masking requires all arithmetic units to be modified due to bit-level sharing and incurs a large performance, energy, and area overhead. For example, Boolean masking increases the energy consumption of a DNN accelerator by 37.9x when Threshold Implementation (TI), a type of Boolean masking that split data into three independent shares to guarantee a high-level of security [212], is used for the multiply-and-accumulate circuitry [209] (Fig. 8). So far, these works target a binarized neural network that does not require multipliers [208] or limit the model parameters of a DNN to be powers-of-2 (e.g., $2^0, 2^1, \dots$) [209]. In particular, the power-of-2 approach of [209] reduces a multiply-and-accumulate operation to a shift and XOR operation, which is linear at the bit-level and easy to implement with Boolean masking. Furthermore, limiting the model parameters to powers-of-2 has negligible impact on the classification accuracy of DNN

models for applications such as ECG arrhythmia detection or epileptic seizure recognition.

Nevertheless, implementing a multi-bit accumulator with TI introduces several challenges. The TI requirements stipulate that the Boolean shares of data should be uniformly distributed and mutually independent of each other [219]–[221]. However, the intermediate computation results of the shift and accumulate operations can violate this. [209] carefully designed its accumulator circuit to not combine shares that are not jointly uniform and to add fresh randomness when combining dependent shares is unavoidable. Random bits are supplied using a lightweight cipher that is also implemented with the TI methodology. Overall, these techniques reduced the overhead of Boolean masking to only 64% more area and a 5.5x increase in the energy consumption for a DNN accelerator supporting multi-bit precision in [209].

2) *Shuffling*: Another approach to decorrelate side-channel measurements from the actual data is to *shuffle* the computations temporally and spatially [210]. A random sequence generator determines the temporal order of the multiplications in a DNN accelerator, and an adversary who does not know this random sequence cannot recover the true order of the data. Also, the spatial allocation of multiplications to each processing element in a DNN accelerator is randomly determined, such that an adversary utilizing high-resolution electromagnetic side-channel observations can be deterred as well. In addition to this randomization, [210] introduces a compensation mechanism to achieve overall constant power consumption, using a regression model that estimates the dynamic and static power of the processing elements. Similarly, [209] also takes advantage of temporal shuffling for input security when activations are processed by a DSP unit before MAC computations.

3) *Limitations*: While this set of works shows a promising start to securing DNN accelerators from physical SCA, there remains significant necessary progress. Many of the accelerator defenses incur significant overheads, at around 2x or higher area and power for iso-latency [205], [208], [210]. The masking and shuffling schemes also depend on a constant supply of random bits, necessitating an on-chip pseudo-random number generator. While works like [207], [209] start to address the reduction of this requirement through careful design and reuse of randomness, further exploration is necessary into the security trade-offs.

Furthermore, as with the side-channel attack research, the prior work mostly focuses on smaller architectures and models, particularly those implemented on an FPGA. We can not expect the overheads and optimizations to directly scale to more practical and larger designs with floating point precision. Particularly, some requirements such as constant conversion between the arithmetic and Boolean masking domains may become prohibitively expensive with many DNN layers of larger sizes.

C. On-chip Fault Injection Detection

When an adversary attacks the power and clock supply to an accelerator chip, all modules in this chip sharing the same

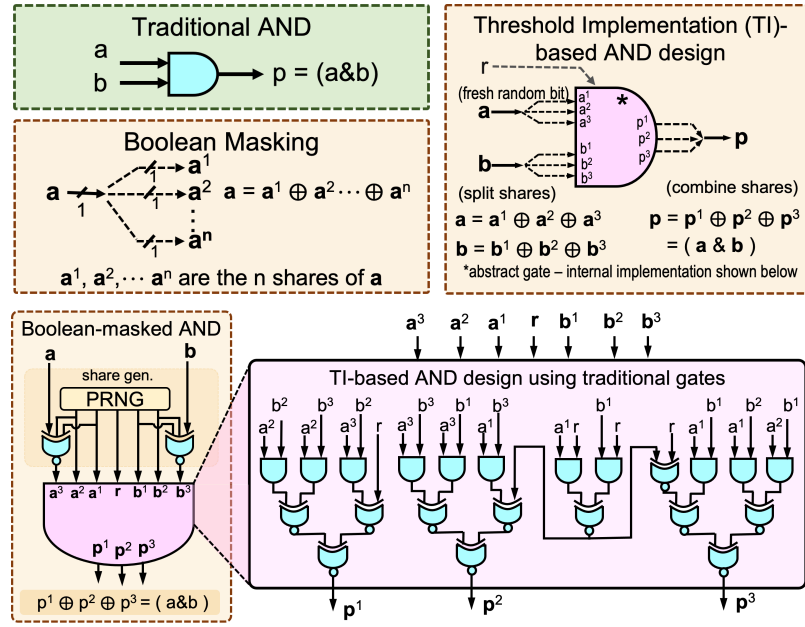


Fig. 8: An AND gate using Threshold Implementation (TI) that splits the data into three Boolean shares. TI is adopted for a side-channel secure DNN accelerator implementation in [209]. This figure is reproduced from [209] (© IEEE 2023).

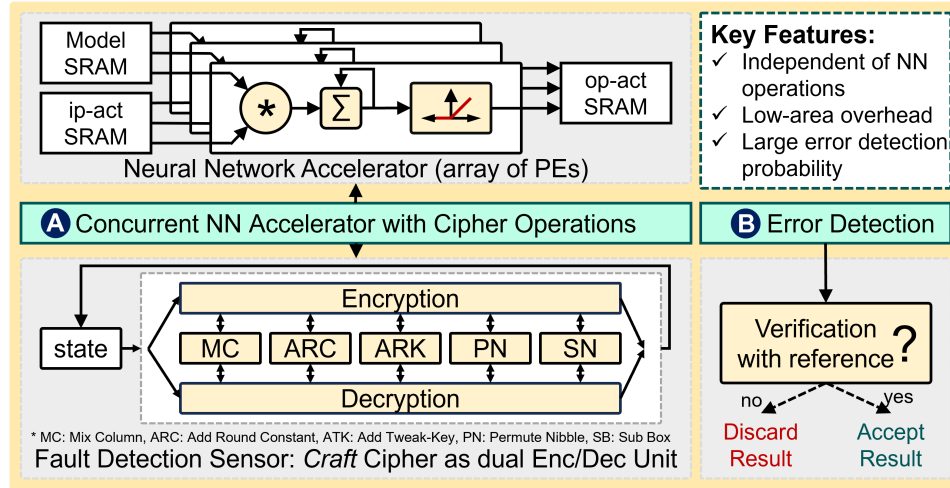


Fig. 9: An on-chip Craft cryptographic engine as a fault detection sensor for EM fault injections and clock glitches [193]. The figure is reproduced from [193] (© IEEE 2023).

power and clock supply will be affected. Similarly, electro-magnetic pulses with low spatial resolution will affect many modules in a chip. Using this property, [193] used an on-chip cryptographic engine performing encryptions and decryptions as a sensor to detect clock glitches and electromagnetic faults. They observed that cryptographic ciphers are highly sensitive to faults [222], and any glitch that affects a DNN accelerator will most likely also affect an on-chip cryptographic engine. A lightweight Craft cipher [223] was used to sequentially perform encryption and decryption over an internal state, such that the internal state looped back to the original value after a certain number of encryption-decryption cycles (Fig. 9). If the internal state does not match the original value, it is likely to indicate an anomaly due to faults and processing can be halted. [193] demonstrated that their detection method only incurs an

area overhead of 5.9% and successfully identifies all 3×10^6 clock glitches and 0.1×10^6 electromagnetic glitches during testing.

1) *Limitations:* While [193] proposed an effective detection method for physical fault injection attacks that affect multiple modules in an accelerator, it cannot detect faults with higher spatial resolution that precisely target a victim component without affecting a detection circuit. Different detection methods have been studied in the context of cryptographic accelerators [224]–[227], which can be further explored for DNN accelerators. In addition to detecting faults, correcting errors due to fault injection attacks can be necessary for applications requiring real-time processing without interruptions.

D. Defense Solutions for In-Memory Compute (IMC)

1) *Side-channel Defenses for IMC*: IMC accelerators often have thousands or millions of parallel multiply-and-accumulate operations, and the overhead of securing these systems can be significant. Unfortunately, there has not been as much research in securing IMC accelerators from physical SCAs, and these are mostly limited to qualitative descriptions of potential defenses when novel attacks were proposed [151], [157]. For example, [151] suggests using dummy memristors and data to hide the DNN model architecture during operation at the expense of significant area, power, and latency.

[228] demonstrated the first thorough defense for digital IMC accelerators using Boolean sharing. Recall that Boolean sharing can be efficiently implemented for operations that are bit-wise linear, such as XOR and XNOR. [228] leveraged linear XNOR gates for multiplication instead of the conventional AND gates to reduce the overhead of Boolean sharing. Additionally, for the adder logic, [228] observed that full adders have better uniformity of outputs compared to half adders. Since uniformity is a key requirement of Boolean sharing and fresh randomness is required if it is not satisfied, utilizing full adders instead of half adders can improve efficiency and security. To this end, [228] used carry-save addition in the adder trees and bit-serial accumulators to eliminate half adder logic.

2) *Data Remanence in IMC*: IMC accelerators have a unique problem with data readout attacks when they use NVM cells that retain the data even after the system has been powered off. An adversary can simply read all the data stored in NVM cells with cold-boot attacks [125], [126]. Furthermore, even for volatile memories, IMC accelerators store the model parameters for long periods of time in regular structures on-chip, and they are susceptible to voltage microprobing attacks [81], [229], [230].

What further complicates the defense solutions for IMC accelerators is that compute and memory are no longer separable. Thus, while the data has to be stored in encrypted ciphertexts for memory security, we cannot simply decrypt the data before the computation (e.g., when the data is fetched from the off-chip memory) as in conventional DNN accelerators. Recent works proposed several approaches to address challenges with memory encryption in IMC accelerators.

The first approach is to integrate decryption with the computation with a simple encryption scheme [202]–[204], [231], [232]. For example, [231], [232] XOR-ed the model parameters with secrets before storing them in the bit cells of an IMC accelerator. Then, the inverse of those secrets is applied to the input tensors before the computation, such that the secrets will be canceled out in the computation outputs. Thus, these works can perform decryption on the fly without having to explicitly do this operation on the model parameters. However, the major source of overhead in this scheme is generating those secrets that have the same size as the model parameters of the entire DNN, which was achieved by using a stream cipher in [232]. Instead of generating highly-secure secrets that are only used once as in [232], other works often used less secure but more efficient schemes to generate

secrets, such as using the intrinsic startup value of bit cells as secrets [202] and flipping entire rows and columns of the memory instead of XOR-ing the data with secrets [203], [204].

The second approach is to leverage arithmetic sharing such that IMC accelerators store and compute only one share of the model parameters [233]. The host secure processor operates on the other share of the model parameters and combines the computation results of the IMC accelerator to obtain the final outputs [233]. However, this approach incurs the overhead of twice the compute, as the computations have to be performed on two arithmetic shares.

Third, similarly to its use in physical SCA defenses (Sec. IV-B2), shuffling can also be applied to IMC accelerators for security [162], [234], [235]. The model parameters can be mapped to IMC bit cells in a permuted manner, and the computation can correct the permutation at the final step.

Lastly, some works [234], [236] leveraged the intrinsic physical characteristics of the analog components in IMC accelerators, such as the circuit offsets [234] and retention limitations [236], to provide the security. For example, [234] proposed fine-tuning a per-chip model that reflects the unavoidable analog-to-digital converter offset, such that the model parameters are specific to each chip and an adversary cannot achieve the same performance in another chip using the extracted model parameters.

3) *Limitations*: The defense solutions for IMC accelerators are limited except for very recent work on a digital IMC accelerator [228], and analysis of how such protections translate to more traditional analog IMC accelerators without affecting compute SNR and performance remains an open question.

E. Hardware Trojan Detection

There have been many methods proposed for hardware trojan detection. At a high level, these utilize methods such as functional or side-channel testing to detect the existence of the trojan, or modify the design to make trojan insertion difficult through methods such as redundancy and operation obfuscation.

Functional testing is likely to find trojans that have a high probability of being triggered with a payload that has an obvious impact on the chip's operation [237]–[239]. However, more complex trigger designs that aim for a low chance of accidental triggering will also make it unlikely for it to be detected with this form of testing.

On the other hand, side-channel testing uses unintentional effects of the addition of trigger and payload circuitry on the power consumption and currents to detect the presence of this additional circuitry [240]–[244]. However, since it is quite difficult to have an accurate power model for complex CMOS circuits, many such methods depend on a golden circuit that is guaranteed to be trojan-free to compare against. Prior work has shown the possibility of combining high spatial resolution and wide field of view novel sensing methods with unsupervised deep learning for unbiased and golden circuit-free trojan detection at high sensitivity [245].

Many of the accelerator trojan triggers and payloads have quite small footprints while also being rare enough to escape

detection via traditional functional verification. However, some work has considered how to find such trojans in the context of machine learning accelerators [157], [246]. For example, some trojan triggers specifically wait for data around the three-sigma limit of the input distribution, which is learned based on the validation dataset provided to the hardware user for functional verification [35], [247]. One low overhead but not fool-proof method to avoid this changes the distribution of the inputs in the validation set from that of the expected data [247].

In addition, redundancy can be used for trojan detection by having an additional processing element that is nominally identical to each of the regular elements at different times [248]. Thus, any deviation between two outputs indicates some malicious change. However, similarly to ECC [249], [250], the amount of redundancy creates an obvious tradeoff between the significant overheads and the number of errors it can detect or correct.

Similarly, obfuscation can be used to shuffle images at the pixel and block levels and to apply random mapping between pixel values to hide the input image data [251]. Thus, triggers that depend on specific patterns of data in a certain order or location are difficult to target by an attacker, but this does not apply to triggers that just detect a certain value.

1) *Limitations:* Most importantly, many of these security or detection mechanisms are incumbent on a specific type of trigger or payload being present, which cannot be known a priori. A general technique that can apply to many varieties of trojans is necessary. One solution that has been explored is a verifiable ASIC, which provides a computationally secure way to detect any trojan that causes a change in the outputs [252]. By providing an interactive proof specialized for matrix multiplication [253] and non-linear activations, there is a negligible chance that a trojan-inserted chip can give the correct response to all the verifier's queries [252]. Some optimizations can be added to exploit neural network-specific parameters such as high sparsity and reuse of existing convolutional operations for the proof [252], but this does still add overheads to the original chip design. Furthermore, we need a detection scheme that is also verifiable against attacks with effects other than modifying the final outputs. To this end, a hybrid solution of functional and side-channel testing, which has minimal dependence on on-chip detection circuitry seems most appropriate.

F. Fully Homomorphic Encryption for Privacy-preserving Computing

Fully Homomorphic Encryption (FHE) has emerged as a mainstream technology in privacy-preserving computing over the past decade. FHE enables cloud operators to perform computations on *encrypted* data without ever requiring to decrypt it. The result of such computation is also in an encrypted form. This encrypted result can only be decrypted by the data owner who holds the private/secret key. An illustrative use case of how a data owner can outsource computation on private data to an untrusted third-party cloud platform is shown in Figure 10.

Since the data remains encrypted throughout the computation process, physical SCAs (Sec. III-B1) on a processor or data readout attacks on an accelerator memory (Sec. III-A1)

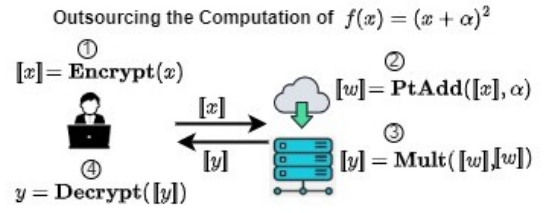


Fig. 10: Third-party cloud platform with outsourced FHE-based computing. The figure is reproduced from [254].

fail to disclose meaningful information about the plaintext data or intermediate computations. To guarantee the confidentiality of data, FHE does not require any trusted hardware, such as a trusted execution environment (Sec. IV-A1) either.

However, despite its robust data privacy guarantees, FHE is not widely adopted. This is primarily due to its inherent slowness when performing computations on encrypted data, which stems from its compute and memory-intensive nature [254]–[256]. This performance issue persists across different FHE schemes, including BGV [257], BFV [258], CKKS [259], and TFHE [260]. All of these schemes, by mathematical construction, support operations on either integers, real numbers, or bits, which enables a wide variety of real-world applications.

The CKKS scheme is currently the most popular FHE scheme for DNNs. While the CKKS scheme can support DNNs, it does require an expensive operation known as bootstrapping to be performed frequently. The noise added from the FHE scheme grows with each homomorphic encryption being performed. Once it grows beyond a certain level, it is impossible to continue computing any further with correctness. This bootstrapping operation de-noises the ciphertext and depending upon the scheme parameters, can take several seconds [261] to several minutes [256].

There have been many recent algorithmic [262]–[265], software [261], [266], [267], and hardware [268]–[273] optimizations to the CKKS primitive and bootstrapping operations to help reduce the bootstrapping execution time. However, bootstrapping remains a performance bottleneck as it consumes up to 95% of the total application execution time [261], [273] when performing deep neural network training and inference. This underscores the critical need for accelerating bootstrapping while performing training and inference using encrypted data.

There have been initiatives to accelerate CKKS on both FPGA and ASIC platforms. HEAX [274] emerged as an early FPGA-based accelerator for FHE which accelerated only CKKS encrypted multiplication, while all other operations were delegated to a host processor. One of the more recent FPGA implementations, FAB [268], implemented packed CKKS bootstrapping for the first time on an FPGA with support for practical parameter sets. The FAB design balances the compute and memory requirements of the encrypted deep neural networks, while being cognizant of the compute and on-chip memory constraints of the underlying FPGA. Even though FAB outperforms all prior CPU/GPU implementations by $9.5\times$ to $456\times$, the performance was still limited by the

bootstrapping implementation, which could not be parallelized on multiple FPGAs.

The first ASIC design was proposed by Samardzic et al. [269], and further optimized in their subsequent work, CraterLake [271]. In parallel, Kim et al. proposed three other ASIC designs namely BTS [270], ARK [272], and SHARP [273]. These designs incorporate various optimizations and have demonstrated remarkable performance in bootstrapping. These ASIC solutions are promising as they outperform CPU/GPU/FPGA implementations and narrow the gap between the performance of computing on plaintext vs. ciphertext when evaluating encrypted deep neural networks. However, to achieve this performance improvement, they make use of a large number of custom modular multipliers, large register files, and large on-chip memory. To enable such architectures, these ASIC implementations must use expensive advanced technology nodes like 7nm or 12nm, making these solutions extremely expensive.

1) *Limitations:* All the aforementioned research works contribute to ongoing efforts aimed at improving the practicality and scalability of FHE schemes, particularly in the context of deep neural networks. Further FHE algorithmic improvements will play a significant role in this direction, demonstrating promising progress towards making secure computations on encrypted data more efficient and viable for real-world applications.

V. FUTURE WORK AND CONCLUSION

Looking forward, the security of DNN accelerators will continue to be an important challenge. Also, as DNNs are rapidly evolving, new challenges and vulnerabilities can arise.

One particular topic for future research can be on the impact of large language models (LLMs) on hardware security. On one hand, LLMs can be exploited by adversaries to automate attacks, such as generating hardware trojans when provided with the source RTL code [275]. However, LLMs can also offer potential benefits to hardware designers. For example, they can be used to automatically identify vulnerabilities in hardware designs [276], [277].

Rapidly evolving DNN workloads that have increasingly large model sizes, such as LLMs, also present a dual challenge for DNN accelerators. DNN accelerators should not only be secured from hardware-level vulnerabilities, but should also provide efficient support for increasingly compute- and memory-intensive workloads. As such, the scalability of current defense solutions can become a critical challenge, as they have to protect more complex models without significant sacrifice in performance.

In summary, this article provide a comprehensive survey into recent advances in hardware security of DNN accelerators. Compared to general-purpose microprocessors or cryptographic accelerators for which hardware security has been thoroughly investigated, DNN accelerators have similarities (Sec. II-A) and unique differences (Sec. II-B) in their hardware-level vulnerabilities. We provide a comprehensive review of recent advances in attacking DNN accelerators exploiting those vulnerabilities (Sec. III) and defense solutions for DNN accelerators (Sec. IV). In addition to recent

progress, there are many potential future directions spanning across circuits, architecture, and algorithms for secure machine learning hardware that satisfies confidentiality, integrity, and availability.

ACKNOWLEDGMENT

This work was funded in part by MIT-IBM Watson AI Lab, Samsung Electronics, Red Hat Collaboratory, Korea Foundation for Advanced Studies, NSF GRFP (Grant No. 1745302), and MathWorks Engineering Fellowship.

REFERENCES

- [1] OpenAI, "Gpt-4 technical report," 2024.
- [2] F. Isensee, P. F. Jaeger, S. A. Kohl, J. Petersen, and K. H. Maier-Hein, "nnu-net: a self-configuring method for deep learning-based biomedical image segmentation," *Nature methods*, vol. 18, no. 2, pp. 203–211, 2021.
- [3] J. M. Jumper, R. Evans, A. Pritzel, T. Green, M. Figurnov, O. Ronneberger, K. Tunyasuvunakool, R. Bates, A. Zidek, A. Potapenko, A. Bridgland, C. Meyer, S. A. A. Kohl, A. Ballard, A. Cowie, B. Romera-Paredes, S. Nikolov, R. Jain, J. Adler, T. Back, S. Petersen, D. A. Reiman, E. Clancy, M. Zielinski, M. Steinegger, M. Pacholska, T. Berghammer, S. Bodenstein, D. Silver, O. Vinyals, A. W. Senior, K. Kavukcuoglu, P. Kohli, and D. Hassabis, "Highly accurate protein structure prediction with alphafold," *Nature*, vol. 596, pp. 583 – 589, 2021. [Online]. Available: <https://api.semanticscholar.org/CorpusID:235959867>
- [4] S. Grigorescu, B. Trasnea, T. Cocias, and G. Macesanu, "A survey of deep learning techniques for autonomous driving," *Journal of field robotics*, vol. 37, no. 3, pp. 362–386, 2020.
- [5] A. Mirhoseini, A. Goldie, M. Yazgan, J. W. Jiang, E. Songhori, S. Wang, Y.-J. Lee, E. Johnson, O. Pathak, A. Nazi et al., "A graph placement methodology for fast chip design," *Nature*, vol. 594, no. 7862, pp. 207–212, 2021.
- [6] *The Health Insurance Portability and Accountability Act (HIPAA)*. U.S. Dept. of Labor, Employee Benefits Security Administration, 2004.
- [7] E. Strickland, "15 graphs that explain the state of ai in 2024: The ai index tracks the generative ai boom, model costs, and responsible ai use," *IEEE Spectrum*, Apr 2024. [Online]. Available: <https://spectrum.ieee.org/ai-index-2024>
- [8] C. Szegedy, W. Zaremba, I. Sutskever, J. Bruna, D. Erhan, I. Goodfellow, and R. Fergus, "Intriguing properties of neural networks," 2013.
- [9] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu, "Towards deep learning models resistant to adversarial attacks," *arXiv preprint arXiv:1706.06083*, 2017.
- [10] A. S. Rakin, Z. He, and D. Fan, "Bit-flip attack: Crushing neural network with progressive bit search," in *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, 2019, pp. 1211–1220.
- [11] F. Yao, A. S. Rakin, and D. Fan, "Deephammer: Depleting the intelligence of deep neural networks through targeted chain of bit flips," in *29th USENIX Security Symposium (USENIX Security 20)*. USENIX Association, Aug. 2020, pp. 1463–1480. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity20/presentation/yao>
- [12] S. Hong, P. Frigo, Y. Kaya, C. Giuffrida, and T. Dumitras, "Terminal brain damage: Exposing the graceless degradation in deep neural networks under hardware fault attacks," in *Proceedings of the 28th USENIX Conference on Security Symposium*, ser. SEC'19. USA: USENIX Association, 2019, p. 497–514.
- [13] K. Lee and A. P. Chandrakasan, "Sparsebfa: Attacking sparse deep neural networks with the worst-case bit flips on coordinates," in *ICASSP 2022 - 2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2022, pp. 4208–4212.
- [14] G. Li, S. K. S. Hari, M. Sullivan, T. Tsai, K. Pattabiraman, J. Emer, and S. W. Keckler, "Understanding error propagation in deep learning neural network (dnn) accelerators and applications," in *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, ser. SC '17. New York, NY, USA: Association for Computing Machinery, 2017. [Online]. Available: <https://doi.org/10.1145/3126908.3126964>
- [15] W. Liu, C.-H. Chang, F. Zhang, and X. Lou, "Imperceptible Misclassification Attack on Deep Learning Accelerator by Glitch Injection," in *ACM/IEEE Design Automation Conference (DAC)*, 2020.

- [16] Y. Luo, C. Gongye, Y. Fei, and X. Xu, "DeepStrike: Remotely-Guided Fault Injection Attacks on DNN Accelerator in Cloud-FPGA," in *ACM/IEEE Design Automation Conference*, 2021, pp. 295–300.
- [17] A. S. Rakin, Y. Luo, X. Xu, and D. Fan, "Deep-Dup: An adversarial weight duplication attack framework to crush deep neural network in Multi-Tenant FPGA," in *USENIX Security Symposium*, 2021.
- [18] M. Sinha, P. Bhattacharyya, S. S. Rout, N. B. Prakriya, and S. Deb, "Securing an accelerator-rich system from flooding-based denial-of-service attacks," *IEEE Transactions on Emerging Topics in Computing*, vol. 10, no. 2, pp. 855–869, 2022.
- [19] S. Charles, Y. Lyu, and P. Mishra, "Real-time detection and localization of dos attacks in noc based socs," in *2019 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2019, pp. 1160–1165.
- [20] P. Kocher, J. Jaffe, and B. Jun, "Differential power analysis," in *Advances in Cryptology—CRYPTO'99: 19th Annual International Cryptology Conference Santa Barbara, California, USA, August 15–19, 1999 Proceedings 19*. Springer, 1999, pp. 388–397.
- [21] D. Agrawal, B. Archambeault, J. R. Rao, and P. Rohatgi, "The em side—channel(s)," in *Cryptographic Hardware and Embedded Systems - CHES 2002*, B. S. Kaliski, ç. K. Koç, and C. Paar, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2003, pp. 29–45.
- [22] Y. Kim, R. Daly, J. Kim, C. Fallin, J. H. Lee, D. Lee, C. Wilkerson, K. Lai, and O. Mutlu, "Flipping bits in memory without accessing them: An experimental study of dram disturbance errors," in *2014 ACM/IEEE 41st International Symposium on Computer Architecture (ISCA)*, 2014, pp. 361–372.
- [23] O. Mutlu and J. S. Kim, "Rowhammer: A retrospective," *Trans. Comp.-Aided Des. Integ. Cir. Sys.*, vol. 39, no. 8, p. 1555–1571, Aug. 2020. [Online]. Available: <https://doi.org/10.1109/TCAD.2019.2915318>
- [24] L. Szekeres, M. Payer, T. Wei, and D. Song, "Sok: Eternal war in memory," in *2013 IEEE Symposium on Security and Privacy*, 2013, pp. 48–62.
- [25] Y.-S. Won, S. Chatterjee, D. Jap, A. Basu, and S. Bhasin, "Deepfreeze: Cold boot attacks and high fidelity model recovery on commercial edgml device," in *2021 IEEE/ACM International Conference On Computer Aided Design (ICCAD)*, 2021, pp. 1–9.
- [26] F. Liu, Y. Yarom, Q. Ge, G. Heiser, and R. B. Lee, "Last-level cache side-channel attacks are practical," in *2015 IEEE Symposium on Security and Privacy*, 2015, pp. 605–622.
- [27] M. Lipp, M. Schwarz, D. Gruss, T. Prescher, W. Haas, A. Fogh, J. Horn, S. Mangard, P. Kocher, D. Genkin, Y. Yarom, and M. Hamburg, "Meltdown: Reading kernel memory from user space," in *27th USENIX Security Symposium (USENIX Security 18)*. Baltimore, MD: USENIX Association, Aug. 2018, pp. 973–990. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity18/presentation/lipp>
- [28] P. Kocher, J. Horn, A. Fogh, D. Genkin, D. Gruss, W. Haas, M. Hamburg, M. Lipp, S. Mangard, T. Prescher, M. Schwarz, and Y. Yarom, "Spectre attacks: Exploiting speculative execution," in *2019 IEEE Symposium on Security and Privacy (SP)*, 2019, pp. 1–19.
- [29] J. Clements and Y. Lao, "Hardware trojan design on neural networks," in *2019 IEEE International Symposium on Circuits and Systems (ISCAS)*, 2019, pp. 1–5.
- [30] V. Venceslai, A. Marchisio, I. Alouani, M. Martina, and M. Shafique, "Neuroattack: Undermining spiking neural networks security through externally triggered bit-flips," in *2020 International Joint Conference on Neural Networks (IJCNN)*, 2020, pp. 1–8.
- [31] A. Warnecke, J. Speith, J.-N. Möller, K. Rieck, and C. Paar, "Evil from within: Machine learning backdoors through hardware trojans," 2023.
- [32] Z. Liu, J. Ye, X. Hu, H. Li, X. Li, and Y. Hu, "Sequence triggered hardware trojan in neural network accelerator," in *2020 IEEE 38th VLSI Test Symposium (VTS)*, 2020, pp. 1–6.
- [33] H. Xu, D. Liu, C. Merkel, and M. Zuzak, "Exploiting logic locking for a neural trojan attack on machine learning accelerators," in *Proceedings of the Great Lakes Symposium on VLSI 2023*, ser. GLSVLSI '23. New York, NY, USA: Association for Computing Machinery, 2023, p. 351–356. [Online]. Available: <https://doi.org/10.1145/3583781.3590242>
- [34] S.-H. Huang, W.-C. Cheng, and J.-F. Li, "Hardware trojans of computing-in-memories: Issues and methods," in *2023 IEEE International Symposium on Defect and Fault Tolerance in VLSI and Nanotechnology Systems (DFT)*, 2023, pp. 1–6.
- [35] T. A. Odetola, F. Khalid, H. Mohammed, T. C. Sandefur, and S. R. Hasan, "Feshi: Feature map-based stealthy hardware intrinsic attack," *IEEE Access*, vol. 9, pp. 115 370–115 387, 2021.
- [36] R. Mukherjee and R. S. Chakraborty, "Novel hardware trojan attack on activation parameters of fpga-based dnn accelerators," *IEEE Embedded Systems Letters*, vol. 14, no. 3, pp. 131–134, 2022.
- [37] C. Yang, J. Hou, M. Wu, K. Mei, and L. Geng, "Hardware trojan attacks on the reconfigurable interconnections of convolutional neural networks accelerators," in *2020 IEEE 15th International Conference on Solid-State & Integrated Circuit Technology (ICSICT)*, 2020, pp. 1–3.
- [38] J. A. Halderman, S. D. Schoen, N. Heninger, W. Clarkson, W. Paul, J. A. Calandrino, A. J. Feldman, J. Appelbaum, and E. W. Felten, "Lest we remember: Cold boot attacks on encryption keys," in *17th USENIX Security Symposium (USENIX Security 08)*. San Jose, CA: USENIX Association, Jul. 2008. [Online]. Available: <https://www.usenix.org/conference/17th-usenix-security-symposium/lest-we-remember-cold-boot-attacks-encryption-keys>
- [39] Y.-H. Chen, T. Krishna, J. S. Emer, and V. Sze, "Eyeriss: An energy-efficient reconfigurable accelerator for deep convolutional neural networks," *IEEE Journal of Solid-State Circuits*, vol. 52, no. 1, pp. 127–138, 2017.
- [40] V. Sze, Y.-H. Chen, T.-J. Yang, and J. S. Emer, "Efficient Processing of Deep Neural Networks: A Tutorial and Survey," *Proceedings of the IEEE*, vol. 105, no. 12, pp. 2295–2329, 2017.
- [41] T. Chen, Z. Du, N. Sun, J. Wang, C. Wu, Y. Chen, and O. Temam, "Diannao: A small-footprint high-throughput accelerator for ubiquitous machine-learning," in *Proceedings of the 19th International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS '14. New York, NY, USA: Association for Computing Machinery, 2014, p. 269–284. [Online]. Available: <https://doi.org/10.1145/2541940.2541967>
- [42] S. Han, X. Liu, H. Mao, J. Pu, A. Pedram, M. A. Horowitz, and W. J. Dally, "Eie: Efficient inference engine on compressed deep neural network," in *Proceedings of the 43rd International Symposium on Computer Architecture*, ser. ISCA '16. IEEE Press, 2016, p. 243–254. [Online]. Available: <https://doi.org/10.1109/ISCA.2016.30>
- [43] B. Moons and M. Verhelst, "An energy-efficient precision-scalable convnet processor in 40-nm cmos," *IEEE Journal of Solid-State Circuits*, vol. 52, no. 4, pp. 903–914, 2017.
- [44] N. P. Jouppi, C. Young, N. Patil, D. Patterson, G. Agrawal, R. Bajwa, S. Bates, S. Bhatia, N. Boden, A. Borchers, R. Boyle, P.-I. Cantin, C. Chao, C. Clark, J. Coriell, M. Daley, M. Dau, J. Dean, B. Gelb, T. V. Ghaemmhami, R. Gottipati, W. Gulland, R. Hagmann, C. R. Ho, D. Hogberg, J. Hu, R. Hundt, D. Hurt, J. Ibarz, A. Jaffey, A. Jaworski, A. Kaplan, H. Khaitan, D. Killebrew, A. Koch, N. Kumar, S. Lacy, J. Laudon, J. Law, D. Le, C. Leary, Z. Liu, K. Lucke, A. Lundin, G. MacKean, A. Maggiore, M. Mahony, K. Miller, R. Nagarajan, R. Narayanaswami, R. Ni, K. Nix, T. Norrie, M. Omernick, N. Penukonda, A. Phelps, J. Ross, M. Ross, A. Salek, E. Samadiani, C. Severn, G. Sizikov, M. Snellman, J. Souter, D. Steinberg, A. Swing, M. Tan, G. Thorson, B. Tian, H. Toma, E. Tuttle, V. Vasudevan, R. Walter, W. Wang, E. Wilcox, and D. H. Yoon, "In-datacenter performance analysis of a tensor processing unit," in *Proceedings of the 44th Annual International Symposium on Computer Architecture*, ser. ISCA '17. New York, NY, USA: Association for Computing Machinery, 2017, p. 1–12. [Online]. Available: <https://doi.org/10.1145/3079856.3080246>
- [45] Y.-H. Chen, J. Emer, and V. Sze, "Eyeriss: A spatial architecture for energy-efficient dataflow for convolutional neural networks," in *2016 ACM/IEEE 43rd Annual International Symposium on Computer Architecture (ISCA)*, 2016, pp. 367–379.
- [46] M. Pellauer, Y. S. Shao, J. Clemons, N. Crago, K. Hegde, R. Venkatesan, S. W. Keckler, C. W. Fletcher, and J. Emer, "Buffets: An efficient and composable storage idiom for explicit decoupled data orchestration," in *Proceedings of the Twenty-Fourth International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS '19. New York, NY, USA: Association for Computing Machinery, 2019, p. 137–151. [Online]. Available: <https://doi.org/10.1145/3297858.3304025>
- [47] A. Parashar, P. Raina, Y. S. Shao, Y.-H. Chen, V. A. Ying, A. Mukkara, R. Venkatesan, B. Khailany, S. W. Keckler, and J. Emer, "Timeloop: A systematic approach to dnn accelerator evaluation," in *2019 IEEE International Symposium on Performance Analysis of Systems and Software (ISPASS)*, 2019, pp. 304–315.
- [48] D. Perez-Botero, J. Szefer, and R. B. Lee, "Characterizing hypervisor vulnerabilities in cloud computing servers," in *Proceedings of the 2013 International Workshop on Security in Cloud Computing*, ser. Cloud Computing '13. New York, NY, USA: Association for Computing Machinery, 2013, p. 3–10. [Online]. Available: <https://doi.org/10.1145/2484402.2484406>
- [49] P. Colp, M. Nanavati, J. Zhu, W. Aiello, G. Coker, T. Deegan, P. Loscocco, and A. Warfield, "Breaking up is hard to do: security and functionality in a commodity hypervisor," in *Proceedings*

- of the *Twenty-Third ACM Symposium on Operating Systems Principles*, ser. SOSP '11. New York, NY, USA: Association for Computing Machinery, 2011, p. 189–202. [Online]. Available: <https://doi.org/10.1145/2043556.2043575>
- [50] J. Szefer, E. Keller, R. B. Lee, and J. Rexford, “Eliminating the hypervisor attack surface for a more secure cloud,” in *Proceedings of the 18th ACM Conference on Computer and Communications Security*, ser. CCS '11. New York, NY, USA: Association for Computing Machinery, 2011, p. 401–412. [Online]. Available: <https://doi.org/10.1145/2046707.2046754>
- [51] H. Chen, Y. Mao, X. Wang, D. Zhou, N. Zeldovich, and M. F. Kaashoek, “Linux kernel vulnerabilities: state-of-the-art defenses and open problems,” in *Proceedings of the Second Asia-Pacific Workshop on Systems*, ser. APSys '11. New York, NY, USA: Association for Computing Machinery, 2011. [Online]. Available: <https://doi.org/10.1145/2103799.2103805>
- [52] Y. Xiao, X. Zhang, Y. Zhang, and R. Teodorescu, “One bit flips, one cloud flows: Cross-VM row hammer attacks and privilege escalation,” in *25th USENIX Security Symposium (USENIX Security 16)*. Austin, TX: USENIX Association, Aug. 2016, pp. 19–35. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity16/technical-sessions/presentation/xiao>
- [53] M. D. Hill, J. Masters, P. Ranganathan, P. Turner, and J. L. Hennessy, “On the spectre and meltdown processor security vulnerabilities,” *IEEE Micro*, vol. 39, no. 2, pp. 9–19, 2019.
- [54] Y. Yarom and K. Falkner, “FLUSH+RELOAD: A high resolution, low noise, l3 cache Side-Channel attack,” in *23rd USENIX Security Symposium (USENIX Security 14)*. San Diego, CA: USENIX Association, Aug. 2014, pp. 719–732. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity14/technical-sessions/presentation/yarom>
- [55] V. Costan and S. Devadas, “Intel sgx explained,” *Cryptology ePrint Archive*, 2016.
- [56] Q. Liu, L. Guo, and H. Tang, “Fault model analysis of dram under electromagnetic fault injection attack,” in *2023 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2023, pp. 1–6.
- [57] B. Colombier, A. Menu, J.-M. Dutertre, P.-A. Moëllic, J.-B. Rigaud, and J.-L. Danger, “Laser-induced single-bit faults in flash memory: Instructions corruption on a 32-bit microcontroller,” in *2019 IEEE International Symposium on Hardware Oriented Security and Trust (HOST)*, 2019, pp. 1–10.
- [58] W. Hua, Z. Zhang, and G. E. Suh, “Reverse engineering convolutional neural networks through side-channel information leaks,” in *Proceedings of the 55th Annual Design Automation Conference*, ser. DAC '18. New York, NY, USA: Association for Computing Machinery, 2018. [Online]. Available: <https://doi.org/10.1145/3195970.3196105>
- [59] X. Hu, L. Liang, S. Li, L. Deng, P. Zuo, Y. Ji, X. Xie, Y. Ding, C. Liu, T. Sherwood, and Y. Xie, “DeepSniffer: A dnn model extraction framework based on learning architectural hints,” in *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS '20. New York, NY, USA: Association for Computing Machinery, 2020, p. 385–399. [Online]. Available: <https://doi.org/10.1145/3373376.3378460>
- [60] D. Yang, P. J. Nair, and M. Lis, “Huffduff: Stealing pruned dnns from sparse accelerators,” in *Proceedings of the 28th ACM International Conference on Architectural Support for Programming Languages and Operating Systems, Volume 2*, ser. ASPLOS 2023. New York, NY, USA: Association for Computing Machinery, 2023, p. 385–399. [Online]. Available: <https://doi.org/10.1145/3575693.3575738>
- [61] S. Volos, K. Vaswani, and R. Bruno, “Graviton: Trusted execution environments on GPUs,” in *13th USENIX Symposium on Operating Systems Design and Implementation (OSDI 18)*. Carlsbad, CA: USENIX Association, Oct. 2018, pp. 681–696. [Online]. Available: <https://www.usenix.org/conference/osdi18/presentation/volos>
- [62] H. Jun, J. Cho, K. Lee, H.-Y. Son, K. Kim, H. Jin, and K. Kim, “Hbm (high bandwidth memory) dram technology and architecture,” in *2017 IEEE International Memory Workshop (IMW)*, 2017, pp. 1–4.
- [63] JEDEC, “High bandwidth memory (hbm) dram,” JESD235D, Technical Report, 2021.
- [64] I. Jang, A. Tang, T. Kim, S. Sethumadhavan, and J. Huh, “Heterogeneous isolated execution for commodity gpus,” in *Proceedings of the Twenty-Fourth International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS '19. New York, NY, USA: Association for Computing Machinery, 2019, p. 455–468. [Online]. Available: <https://doi.org/10.1145/3297858.3304021>
- [65] D. Brumley and D. Boneh, “Remote timing attacks are practical,” *Computer Networks*, vol. 48, no. 5, pp. 701–716, 2005, web Security. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1389128605000125>
- [66] C. Gongye, Y. Fei, and T. Wahl, “Reverse-engineering deep neural networks using floating-point timing side-channels,” in *2020 57th ACM/IEEE Design Automation Conference (DAC)*, 2020, pp. 1–6.
- [67] V. Duddu, D. Samanta, D. V. Rao, and V. E. Balas, “Stealing neural networks via timing side channels,” 2019.
- [68] S. Maji, U. Banerjee, and A. P. Chandrakasan, “Leaky nets: Recovering embedded neural network models and inputs through simple power and timing side-channels—attacks and defenses,” *IEEE Internet of Things Journal*, vol. 8, no. 15, pp. 12079–12092, 2021.
- [69] L. Batina, S. Bhasin, D. Jap, and S. Pickel, “CSI NN: Reverse engineering of neural network architectures through electromagnetic side channel,” in *28th USENIX Security Symposium (USENIX Security 19)*. Santa Clara, CA: USENIX Association, Aug. 2019, pp. 515–532. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity19/presentation/batina>
- [70] L. Wei, B. Luo, Y. Li, Y. Liu, and Q. Xu, “I know what you see: Power side-channel attack on convolutional neural network accelerators,” in *Proceedings of the 34th Annual Computer Security Applications Conference*, ser. ACSAC '18. New York, NY, USA: Association for Computing Machinery, 2018, p. 393–406. [Online]. Available: <https://doi.org/10.1145/3274694.3274696>
- [71] S. Moini, S. Tian, D. Holcomb, J. Szefer, and R. Tessier, “Remote power side-channel attacks on bnn accelerators in fpgas,” in *2021 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2021, pp. 1639–1644.
- [72] K. Yoshida, T. Kubota, M. Shiozaki, and T. Fujino, “Model-extraction attack against fpga-dnn accelerator utilizing correlation electromagnetic analysis,” in *2019 IEEE 27th Annual International Symposium on Field-Programmable Custom Computing Machines (FCCM)*, 2019, pp. 318–318.
- [73] K. YOSHIDA, M. SHIOZAKI, S. OKURA, T. KUBOTA, and T. FUJINO, “Model reverse-engineering attack against systolic-array-based dnn accelerator using correlation power analysis,” *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences*, vol. E104.A, no. 1, pp. 152–161, 2021.
- [74] Y.-S. Won, S. Chatterjee, D. Jap, A. Basu, and S. Bhasin, “Wac: First results on practical side-channel attacks on commercial machine learning accelerator,” in *Proceedings of the 5th Workshop on Attacks and Solutions in Hardware Security*, ser. ASHES '21. New York, NY, USA: Association for Computing Machinery, 2021, p. 111–114. [Online]. Available: <https://doi.org/10.1145/3474376.3487284>
- [75] A. Dubey, R. Cammarota, and A. Aysu, “Maskednet: The first hardware inference engine aiming power side-channel protection,” in *2020 IEEE International Symposium on Hardware Oriented Security and Trust (HOST)*. Los Alamitos, CA, USA: IEEE Computer Society, dec 2020, pp. 197–208. [Online]. Available: <https://doi.ieeecomputersociety.org/10.1109/HOST45689.2020.9300276>
- [76] H. Naghibijouybari, A. Neupane, Z. Qian, and N. Abu-Ghazaleh, “Rendered insecure: Gpu side channel attacks are practical,” in *Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security*, ser. CCS '18. New York, NY, USA: Association for Computing Machinery, 2018, p. 2139–2153. [Online]. Available: <https://doi.org/10.1145/3243734.3243831>
- [77] E. Karimi, Z. H. Jiang, Y. Fei, and D. Kaeli, “A timing side-channel attack on a mobile gpu,” in *2018 IEEE 36th International Conference on Computer Design (ICCD)*, 2018, pp. 67–74.
- [78] W. Liu, C.-H. Chang, F. Zhang, and X. Lou, “Imperceptible misclassification attack on deep learning accelerator by glitch injection,” in *2020 57th ACM/IEEE Design Automation Conference (DAC)*, 2020, pp. 1–6.
- [79] X. Hou, J. Breier, D. Jap, L. Ma, S. Bhasin, and Y. Liu, “Physical security of deep learning on edge devices: Comprehensive evaluation of fault injection attack vectors,” *Microelectronics Reliability*, vol. 120, p. 114116, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0026271421000822>
- [80] H. Wang, D. Forte, M. M. Tehranipoor, and Q. Shi, “Probing attacks on integrated circuits: Challenges and research opportunities,” *IEEE Design & Test*, vol. 34, no. 5, pp. 63–71, 2017.
- [81] M. T. Rahman, Q. Shi, S. Tajik, H. Shen, D. L. Woodard, M. Tehranipoor, and N. Asadizanjani, “Physical inspection & attacks: New frontier in hardware security,” in *2018 IEEE 3rd International Verification and Security Workshop (IVSW)*, 2018, pp. 93–102.
- [82] X. Zhang, A. A. Ding, and Y. Fei, “Deep-learning model extraction through software-based power side-channel,” in *2023 IEEE/ACM In-*

- ternational Conference on Computer Aided Design (ICCAD), 2023, pp. 1–9.
- [83] M. Yan, C. W. Fletcher, and J. Torrellas, “Cache telepathy: Leveraging shared resource attacks to learn DNN architectures,” in *29th USENIX Security Symposium (USENIX Security 20)*. USENIX Association, Aug. 2020, pp. 2003–2020. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity20/presentation/yan>
- [84] F. Schellenberg, D. R. Gnad, A. Moradi, and M. B. Tahoori, “An inside job: Remote power analysis attacks on fpgas,” in *2018 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2018, pp. 1111–1116.
- [85] S. Tian, S. Moini, D. Holcomb, R. Tessier, and J. Szefer, “A practical remote power attack on machine learning accelerators in cloud fpgas,” in *2023 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2023, pp. 1–6.
- [86] S. Moini, S. Tian, D. Holcomb, J. Szefer, and R. Tessier, “Power side-channel attacks on bnn accelerators in remote fpgas,” *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, vol. 11, no. 2, pp. 357–370, 2021.
- [87] Y. Zhang, R. Yasaei, H. Chen, Z. Li, and M. A. A. Faruque, “Stealing neural network structure through remote fpga side-channel analysis,” *IEEE Transactions on Information Forensics and Security*, vol. 16, pp. 4377–4388, 2021.
- [88] J. Mahmod and M. Hicks, “Sram has no chill: exploiting power domain separation to steal on-chip secrets,” in *Proceedings of the 27th ACM International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS ’22. New York, NY, USA: Association for Computing Machinery, 2022, p. 1043–1055. [Online]. Available: <https://doi.org/10.1145/3503222.3507710>
- [89] —, “Untrustzone: Systematic accelerated aging to expose on-chip secrets,” in *2024 IEEE Symposium on Security and Privacy (SP)*. Los Alamitos, CA, USA: IEEE Computer Society, may 2024, pp. 72–72. [Online]. Available: <https://doi.ieeecomputersociety.org/10.1109/SP54263.2024.00069>
- [90] L. Lin, M. Kasper, T. Güneysu, C. Paar, and W. Burleson, “Trojan side-channels: Lightweight hardware trojans through side-channel engineering,” in *Cryptographic Hardware and Embedded Systems - CHES 2009, 11th International Workshop, Lausanne, Switzerland, September 6-9, 2009, Proceedings*, ser. Lecture Notes in Computer Science, vol. 5747. Springer, 2009, pp. 382–395. [Online]. Available: <https://www.iacr.org/archive/ches2009/57470383/57470383.pdf>
- [91] A. De, M. Nasim Imtiaz Khan, K. Nagarajan, and S. Ghosh, “Hart-bleed: Using hardware trojans for data leakage exploits,” *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, vol. 28, no. 4, pp. 968–979, 2020.
- [92] G. T. Becker, A. Lakshminarasimhan, L. Lin, S. Srivathsa, V. B. Suresh, and W. Bureson, “Implementing hardware trojans: Experiences from a hardware trojan challenge,” in *2011 IEEE 29th International Conference on Computer Design (ICCD)*, 2011, pp. 301–304.
- [93] D. Knichel, T. Moos, and A. Moradi, “The risk of outsourcing: Hidden sea trojans in third-party ip-cores threaten cryptographic ics,” in *2020 IEEE European Test Symposium (ETS)*, 2020, pp. 1–6.
- [94] P. Gaikwad, J. Cruz, P. Chakraborty, S. Bhunia, and T. Hoque, “Hardware ip assurance against trojan attacks with machine learning and post-processing,” *J. Emerg. Technol. Comput. Syst.*, vol. 19, no. 3, jun 2023. [Online]. Available: <https://doi.org/10.1145/3592795>
- [95] N. Gupta, A. Jati, and A. Chattopadhyay, “Ai attacks ai: Recovering neural network architecture from nvla using ai-assisted side channel attack,” *Cryptology ePrint Archive*, Paper 2023/368, 2023, <https://eprint.iacr.org/2023/368>. [Online]. Available: <https://eprint.iacr.org/2023/368>
- [96] J. A. Halderman, S. D. Schoen, N. Heninger, W. Clarkson, W. Paul, J. A. Calandrino, A. J. Feldman, J. Appelbaum, and E. W. Felten, “Lest we remember: Cold-boot attacks on encryption keys,” *Commun. ACM*, vol. 52, no. 5, p. 91–98, may 2009. [Online]. Available: <https://doi.org/10.1145/1506409.1506429>
- [97] H. Yu, H. Ma, K. Yang, Y. Zhao, and Y. Jin, “Deepem: Deep neural networks model recovery through em side-channel information leakage,” in *2020 IEEE International Symposium on Hardware Oriented Security and Trust (HOST)*, 2020, pp. 209–218.
- [98] Y. Zhang, R. Yasaei, H. Chen, Z. Li, and M. A. A. Faruque, “Stealing neural network structure through remote fpga side-channel analysis,” *IEEE Transactions on Information Forensics and Security*, vol. 16, pp. 4377–4388, 2021.
- [99] X. Yan, X. Lou, G. Xu, H. Qiu, S. Guo, C. H. Chang, and T. Zhang, “Mercury: An automated remote side-channel attack to nvidia deep learning accelerator,” in *2023 International Conference on Field Programmable Technology (ICFPT)*, 2023, pp. 188–197.
- [100] S. Maji, K. Lee, and A. P. Chandrakasan, “Sparseleakynets: Classification prediction attack over sparsity-aware embedded neural networks using timing side-channel information,” *IEEE Computer Architecture Letters*, pp. 1–4, 2024.
- [101] H. Bar-El, H. Choukri, D. Naccache, M. Tunstall, and C. Whelan, “The sorcerer’s apprentice guide to fault attacks,” *Proceedings of the IEEE*, vol. 94, no. 2, pp. 370–382, 2006.
- [102] D. Boneh, R. A. DeMillo, and R. J. Lipton, “On the importance of checking cryptographic protocols for faults,” in *Advances in Cryptology — EUROCRYPT ’97*, W. Fumy, Ed. Berlin, Heidelberg: Springer Berlin Heidelberg, 1997, pp. 37–51.
- [103] D. Hendrycks and T. Dietterich, “Benchmarking neural network robustness to common corruptions and perturbations,” *Proceedings of the International Conference on Learning Representations*, 2019.
- [104] O. Temam, “A defect-tolerant accelerator for emerging high-performance applications,” in *2012 39th Annual International Symposium on Computer Architecture (ISCA)*, 2012, pp. 356–367.
- [105] Q. Zhang, T. Wang, Y. Tian, F. Yuan, and Q. Xu, “Approxann: An approximate computing framework for artificial neural network,” in *2015 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2015, pp. 701–706.
- [106] W. Haensch, T. Gokmen, and R. Puri, “The next generation of deep learning hardware: Analog computing,” *Proceedings of the IEEE*, vol. 107, no. 1, pp. 108–122, 2019.
- [107] A. Shafiee, A. Nag, N. Muralimanohar, R. Balasubramanian, J. P. Strachan, M. Hu, R. S. Williams, and V. Srikumar, “Isaac: a convolutional neural network accelerator with in-situ analog arithmetic in crossbars,” in *Proceedings of the 43rd International Symposium on Computer Architecture*, ser. ISCA ’16. IEEE Press, 2016, p. 14–26. [Online]. Available: <https://doi.org/10.1109/ISCA.2016.12>
- [108] M. Onen, N. Emond, B. Wang, D. Zhang, F. M. Ross, J. Li, B. Yildiz, and J. A. Del Alamo, “Nanosecond protonic programmable resistors for analog deep learning,” *Science*, vol. 377, no. 6605, pp. 539–543, 2022.
- [109] S. Ambrogio, P. Narayanan, H. Tsai, R. M. Shelby, I. Boybat, C. Di Nolfo, S. Sidler, M. Giordano, M. Bodini, N. C. Farinha et al., “Equivalent-accuracy accelerated neural-network training using analogue memory,” *Nature*, vol. 558, no. 7708, pp. 60–67, 2018.
- [110] M. Dworkin, E. Barker, J. Nechvatal, J. Foti, L. Bassham, E. Roback, and J. A. Dray, “Advanced encryption standard (aes),” 2001–11-26 2001.
- [111] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “Imagenet classification with deep convolutional neural networks,” in *Advances in Neural Information Processing Systems*, F. Pereira, C. Burges, L. Bottou, and K. Weinberger, Eds., vol. 25. Curran Associates, Inc., 2012. [Online]. Available: <https://proceedings.neurips.cc/paper/2012/file/c399862d3b9d6b76c8436e924a68c45b-Paper.pdf>
- [112] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, Jun 2016. [Online]. Available: <http://dx.doi.org/10.1109/CVPR.2016.90>
- [113] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, and L.-C. Chen, “Mobilenetv2: Inverted residuals and linear bottlenecks,” 2019.
- [114] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. u. Kaiser, and I. Polosukhin, “Attention is all you need,” in *Advances in Neural Information Processing Systems*, I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, Eds., vol. 30. Curran Associates, Inc., 2017. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf
- [115] S. Han, H. Mao, and W. J. Dally, “Deep compression: Compressing deep neural networks with pruning, trained quantization and Huffman coding,” 2016.
- [116] M. Rastegari, V. Ordonez, J. Redmon, and A. Farhadi, “Xnor-net: Imagenet classification using binary convolutional neural networks,” in *European conference on computer vision*. Springer, 2016, pp. 525–542.
- [117] M. Courbariaux, I. Hubara, D. Soudry, R. El-Yaniv, and Y. Bengio, “Binarized neural networks,” in *Neural Information Processing Systems*, 2016. [Online]. Available: <https://api.semanticscholar.org/CorpusID:14796162>
- [118] B. Jacob, S. Kligys, B. Chen, M. Zhu, M. Tang, A. Howard, H. Adam, and D. Kalenichenko, “Quantization and training of neural networks for efficient integer-arithmetic-only inference,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2018, pp. 2704–2713.

- [119] C. Zhu, S. Han, H. Mao, and W. J. Dally, "Trained ternary quantization," in *International Conference on Learning Representations*, 2017. [Online]. Available: https://openreview.net/forum?id=S1_pAu9xl
- [120] S. Han, J. Pool, J. Tran, and W. J. Dally, "Learning both weights and connections for efficient neural networks," 2015.
- [121] S. Anwar, K. Hwang, and W. Sung, "Structured pruning of deep convolutional neural networks," *J. Emerg. Technol. Comput. Syst.*, vol. 13, no. 3, feb 2017. [Online]. Available: <https://doi.org/10.1145/3005348>
- [122] J.-H. Luo, J. Wu, and W. Lin, "Thinet: A filter level pruning method for deep neural network compression," in *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, Oct 2017.
- [123] P. Molchanov, A. Mallya, S. Tyree, I. Frosio, and J. Kautz, "Importance estimation for neural network pruning," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- [124] A. S. Rakin, Z. He, J. Li, F. Yao, C. Chakrabarti, and D. Fan, "T-bfa: Targeted bit-flip adversarial weight attack," 2021.
- [125] M. Gruhn and T. Müller, "On the practicability of cold boot attacks," in *2013 International Conference on Availability, Reliability and Security*, 2013, pp. 390–397.
- [126] J. A. Halderman, S. D. Schoen, N. Heninger, W. Clarkson, W. Paul, J. A. Calandrino, A. J. Feldman, J. Appelbaum, and E. W. Felten, "Lest we remember: cold-boot attacks on encryption keys," *Commun. ACM*, vol. 52, no. 5, p. 91–98, may 2009. [Online]. Available: <https://doi.org/10.1145/1506409.1506429>
- [127] "Intel openvino toolkit." [Online]. Available: <https://www.intel.com/content/www/us/en/developer/tools/opencv-toolkit/overview.html>
- [128] Y. He, M. Hutton, S. Chan, R. De Gruijl, R. Govindaraju, N. Patil, and Y. Li, "Understanding and mitigating hardware failures in deep learning training systems," in *Proceedings of the 50th Annual International Symposium on Computer Architecture*, ser. ISCA '23. New York, NY, USA: Association for Computing Machinery, 2023. [Online]. Available: <https://doi.org/10.1145/3579371.3589105>
- [129] P. Kocher, J. Jaffe, and B. Jun, "Differential power analysis," in *Advances in Cryptology — CRYPTO' 99*, M. Wiener, Ed. Berlin, Heidelberg: Springer Berlin Heidelberg, 1999, pp. 388–397.
- [130] T. Messerges, E. Dabbish, and R. Sloan, "Examining Smart-Card Security under the Threat of Power Analysis Attacks," *IEEE Transactions on Computers*, vol. 51, no. 5, pp. 541–552, 2002.
- [131] E. Brier, C. Clavier, and F. Olivier, "Correlation power analysis with a leakage model," in *Cryptographic Hardware and Embedded Systems - CHES 2004*, M. Joye and J.-J. Quisquater, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2004, pp. 16–29.
- [132] D. Hwang, K. Tiri, A. Hodjat, B.-C. Lai, S. Yang, P. Schaumont, and I. Verbauwhede, "AES-Based Security Coprocessor IC in 0.18- μ m CMOS With Resistance to Differential Power Analysis Side-Channel Attacks," *IEEE Journal of Solid-State Circuits*, vol. 41, no. 4, pp. 781–792, 2006.
- [133] D. Das, A. Golder, J. Danial, S. Ghosh, A. Raychowdhury, and S. Sen, "X-deepsca: Cross-device deep learning side channel attack," in *2019 56th ACM/IEEE Design Automation Conference (DAC)*, 2019, pp. 1–6.
- [134] T. Jeong, A. P. Chandrakasan, and H.-S. Lee, "S2adc: A 12-bit, 1.25-ms/s secure sar adc with power side-channel attack resistance," *IEEE Journal of Solid-State Circuits*, vol. 56, no. 3, pp. 844–854, 2021.
- [135] B. Hettwer, S. Gehrler, and T. Guneyssu, "Applications of machine learning techniques in side-channel attacks: a survey," *Journal of Cryptographic Engineering*, vol. 10, no. 2, pp. 135–162, 2020.
- [136] Z. Martinasek, J. Hajny, and L. Malina, "Optimization of power analysis using neural network," in *Smart Card Research and Advanced Applications*, A. Francillon and P. Rohatgi, Eds. Cham: Springer International Publishing, 2014, pp. 94–107.
- [137] S. Chari, J. R. Rao, and P. Rohatgi, "Template attacks," in *Cryptographic Hardware and Embedded Systems - CHES 2002*, B. S. Kaliski, ç. K. Koç, and C. Paar, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2003, pp. 13–28.
- [138] L. Weissbart, S. Picek, and L. Batina, "One trace is all it takes: Machine learning-based side-channel attack on eddsa," in *Security, Privacy, and Applied Cryptography Engineering*, S. Bhasin, A. Mendelson, and M. Nandi, Eds. Cham: Springer International Publishing, 2019, pp. 86–105.
- [139] S. Ors, F. Gurkaynak, E. Oswald, and B. Preneel, "Power-analysis attack on an ASIC AES implementation," in *International Conference on Information Technology: Coding and Computing, 2004. Proceedings. ITCC 2004.*, vol. 2, 2004, pp. 546–552 Vol.2.
- [140] M. Alioto, L. Giancane, G. Scotti, and A. Trifiletti, "Leakage Power Analysis Attacks: A Novel Class of Attacks to Nanometer Cryptographic Circuits," *IEEE Transactions on Circuits and Systems I: Regular Papers*, vol. 57, no. 2, pp. 355–367, 2010.
- [141] Y.-I. Hayashi, N. Homma, T. Mizuki, T. Aoki, H. Sone, L. Sauvage, and J.-L. Danger, "Analysis of Electromagnetic Information Leakage From Cryptographic Devices With Different Physical Structures," *IEEE Transactions on Electromagnetic Compatibility*, vol. 55, no. 3, pp. 571–580, 2013.
- [142] A. Ghosh, D. Das, J. Danial, V. De, S. Ghosh, and S. Sen, "An EM/Power SCA-Resilient AES-256 with Synthesizable Signature Attenuation Using Digital-Friendly Current Source and RO-Bleed-Based Integrated Local Feedback and Global Switched-Mode Control," in *2021 IEEE International Solid-State Circuits Conference (ISSCC)*, vol. 64, 2021, pp. 499–501.
- [143] M. Méndez Real and R. Salvador, "Physical side-channel attacks on embedded neural networks: A survey," *Applied Sciences*, vol. 11, no. 15, 2021. [Online]. Available: <https://www.mdpi.com/2076-3417/11/15/6790>
- [144] H. Chabanne, J.-L. Danger, L. Guiga, and U. Kühne, "Side channel attacks for architecture extraction of neural networks," *CAAI Transactions on Intelligence Technology*, vol. 6, no. 1, pp. 3–16, 2021. [Online]. Available: <https://ietresearch.onlinelibrary.wiley.com/doi/abs/10.1049/cit.2.12026>
- [145] V. Meyers, D. Gnad, and M. Tahoori, "Active and passive physical attacks on neural network accelerators," *IEEE Design & Test*, vol. 40, no. 5, pp. 70–85, 2023.
- [146] L. Deng, "The mnist database of handwritten digit images for machine learning research [best of the web]," *IEEE Signal Processing Magazine*, vol. 29, no. 6, pp. 141–142, 2012.
- [147] Y. Gao, H. Ma, M. Yan, J. He, Y. Zhao, and Y. Jin, "Nnleak: An ai-oriented dnn model extraction attack through multi-stage side channel analysis," in *2023 Asian Hardware Oriented Security and Trust Symposium (AsianHOST)*, 2023, pp. 1–6.
- [148] M. Kang, M.-S. Keel, N. R. Shanbhag, S. Eilert, and K. Curewitz, "An energy-efficient vlsi architecture for pattern recognition via deep embedding of computation in sram," in *2014 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2014, pp. 8326–8330.
- [149] N. Verma, H. Jia, H. Valavi, Y. Tang, M. Ozatay, L.-Y. Chen, B. Zhang, and P. Deaville, "In-memory computing: Advances and prospects," *IEEE Solid-State Circuits Magazine*, vol. 11, no. 3, pp. 43–55, 2019.
- [150] Z. Wang, Y. Wu, Y. Park, S. Yoo, X. Wang, J. K. Eshraghian, and W. D. Lu, "Powergan: A machine learning approach for power side-channel attack on compute-in-memory accelerators," *Advanced Intelligent Systems*, vol. 5, no. 12, p. 2300313, 2023. [Online]. Available: <https://onlinelibrary.wiley.com/doi/abs/10.1002/aisy.202300313>
- [151] Z. Wang, F.-H. Meng, Y. Park, J. K. Eshraghian, and W. D. Lu, "Side-channel attack analysis on in-memory computing architectures," *IEEE Transactions on Emerging Topics in Computing*, vol. 12, no. 1, pp. 109–121, 2024.
- [152] F. Staudigl, F. Merchant, and R. Leupers, "A survey of neuromorphic computing-in-memory: Architectures, simulators, and security," *IEEE Design & Test*, vol. 39, no. 2, pp. 90–99, 2022.
- [153] G. Singh, L. Chelini, S. Corda, A. Javed Awan, S. Stuijk, R. Jordans, H. Corporaal, and A.-J. Boonstra, "A review of near-memory computing architectures: Opportunities and challenges," in *2018 21st Euromicro Conference on Digital System Design (DSD)*, 2018, pp. 608–617.
- [154] M. T. Arafin and Z. Lu, "Security challenges of processing-in-memory systems," in *Proceedings of the 2020 on Great Lakes Symposium on VLSI*, ser. GLSVLSI '20. New York, NY, USA: Association for Computing Machinery, 2020, p. 229–234. [Online]. Available: <https://doi.org/10.1145/3386263.3411365>
- [155] J. Sepulveda, C. Reinbrecht, and J.-P. Diguët, "Security aspects of neuromorphic mpocs," in *2018 IEEE/ACM International Conference on Computer-Aided Design (ICCAD)*, 2018, pp. 1–6.
- [156] D. Samyde, S. Skorobogatov, R. Anderson, and J.-J. Quisquater, "On a new way to read data from memory," in *First International IEEE Security in Storage Workshop, 2002. Proceedings.*, 2002, pp. 65–69.
- [157] M. N. I. Khan and S. Ghosh, "Comprehensive study of security and privacy of emerging non-volatile memories," *Journal of Low Power Electronics and Applications*, vol. 11, no. 4, 2021. [Online]. Available: <https://www.mdpi.com/2079-9268/11/4/36>
- [158] A. Chakraborty, A. Mondal, and A. Srivastava, "Correlation power analysis attack against stt-mram based cyptosystems," *Cryptology*

- ePrint Archive, Paper 2017/413, 2017, <https://eprint.iacr.org/2017/413>. [Online]. Available: <https://eprint.iacr.org/2017/413>
- [159] M. N. I. Khan, S. Bhasin, A. Yuan, A. Chattopadhyay, and S. Ghosh, "Side-channel attack on stram based cache for cryptographic application," in *2017 IEEE International Conference on Computer Design (ICCD)*, 2017, pp. 33–40.
 - [160] V. Kommareddy, B. Zhang, F. Yao, R. Ewetz, and A. Awad, "Are crossbar memories secure? new security vulnerabilities in crossbar memories," *IEEE Computer Architecture Letters*, vol. 18, no. 02, pp. 174–177, jul 2019.
 - [161] S. Skorobogatov, *Physical Attacks and Tamper Resistance*. New York, NY: Springer New York, 2012, pp. 143–173. [Online]. Available: https://doi.org/10.1007/978-1-4419-8080-9_7
 - [162] Y. Wang, S. Jin, and T. Li, "A low cost weight obfuscation scheme for security enhancement of reram based neural network accelerators," in *2021 26th Asia and South Pacific Design Automation Conference (ASP-DAC)*, 2021, pp. 499–504.
 - [163] M. Holler, M. Guizar-Sicairos, E. H. R. Tsai, R. Dinapoli, E. Müller, O. Bunk, J. Raabe, and G. Aeppli, "High-resolution non-destructive three-dimensional imaging of integrated circuits," *Nature*, vol. 543, no. 7645, pp. 402–406, Mar 2017. [Online]. Available: <https://doi.org/10.1038/nature21698>
 - [164] Y.-Y. Tee, X. Hong, D. Cheng, C.-S. Chee, Y. Shi, T. Lin, and B.-H. Gwee, "Patch-based adversarial training for error-aware circuit annotation of delayed ic images," *IEEE Transactions on Circuits and Systems II: Express Briefs*, vol. 70, no. 9, pp. 3694–3698, 2023.
 - [165] D. Cheng, Y. Shi, T. Lin, B.-H. Gwee, and K.-A. Toh, "Hybrid K - means clustering and support vector machine method for via and metal line detections in delayed ic images," *IEEE Transactions on Circuits and Systems II: Express Briefs*, vol. 65, no. 12, pp. 1849–1853, 2018.
 - [166] J. Breier, X. Hou, D. Jap, L. Ma, S. Bhasin, and Y. Liu, "Practical fault attack on deep neural networks," in *Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security*, ser. CCS '18. New York, NY, USA: Association for Computing Machinery, 2018, p. 2204–2206. [Online]. Available: <https://doi.org/10.1145/3243734.3278519>
 - [167] J. Breier, D. Jap, X. Hou, S. Bhasin, and Y. Liu, "SNIFF: Reverse Engineering of NBeural Networks With Fault Attacks," *IEEE Transactions on Reliability*, vol. 71, no. 4, pp. 1527–1539, 2021.
 - [168] O. Mutlu and J. S. Kim, "Rowhammer: A retrospective," *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, vol. 39, no. 8, pp. 1555–1571, 2020.
 - [169] M. N. Imtiaz Khan, K. Nagarajan, and S. Ghosh, "Hardware trojans in emerging non-volatile memories," in *2019 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2019, pp. 396–401.
 - [170] K. Nagarajan, M. N. I. Khan, and S. Ghosh, "Entt: A family of emerging nvm-based trojan triggers," in *2019 IEEE International Symposium on Hardware Oriented Security and Trust (HOST)*, 2019, pp. 51–60.
 - [171] L. Wu, R. Sreekumar, R. Sharifi, K. Skadron, M. R. Stant, and A. Venkat, "Hardware trojans in envm neuromorphic devices," in *2023 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2023, pp. 1–6.
 - [172] Y. Zhao, X. Hu, S. Li, J. Ye, L. Deng, Y. Ji, J. Xu, D. Wu, and Y. Xie, "Memory trojan attack on neural network accelerators," in *2019 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2019, pp. 1415–1420.
 - [173] B. Gassend, G. Suh, D. Clarke, M. van Dijk, and S. Devadas, "Caches and hash trees for efficient memory integrity verification," in *The Ninth International Symposium on High-Performance Computer Architecture, 2003. HPCA-9 2003. Proceedings.*, 2003, pp. 295–306.
 - [174] C. Yan, D. Engländer, M. Prvulovic, B. Rogers, and Y. Solihin, "Improving cost, performance, and security of memory encryption and authentication," in *33rd International Symposium on Computer Architecture (ISCA'06)*, 2006, pp. 179–190.
 - [175] B. Rogers, S. Chhabra, M. Prvulovic, and Y. Solihin, "Using address independent seed encryption and bonsai merkle trees to make secure processors os- and performance-friendly," in *40th Annual IEEE/ACM International Symposium on Microarchitecture (MICRO 2007)*, 2007, pp. 183–196.
 - [176] M. Taassori, A. Shafiee, and R. Balasubramonian, "Vault: Reducing paging overheads in sgx with efficient integrity verification structures," in *Proceedings of the Twenty-Third International Conference on Architectural Support for Programming Languages and Operating Systems*, ser. ASPLOS '18. New York, NY, USA: Association for Computing Machinery, 2018, p. 665–678. [Online]. Available: <https://doi.org/10.1145/3173162.3177155>
 - [177] G. Saileshwar, P. J. Nair, P. Ramrakhiani, W. Elsassser, J. A. Joao, and M. K. Qureshi, "Morphable counters: Enabling compact integrity trees for low-overhead secure memories," in *2018 51st Annual IEEE/ACM International Symposium on Microarchitecture (MICRO)*, 2018, pp. 416–427.
 - [178] D. Lee, D. Kohlbrenner, S. Shinde, K. Asanović, and D. Song, "Keystone: An open framework for architecting trusted execution environments," in *Proceedings of the Fifteenth European Conference on Computer Systems*, ser. EuroSys '20. New York, NY, USA: Association for Computing Machinery, 2020. [Online]. Available: <https://doi.org/10.1145/3342195.3387532>
 - [179] T. Alves and D. Felton, "Trustzone: Integrated hardware and software security," 01 2004.
 - [180] W. Hua, M. Umar, Z. Zhang, and G. E. Suh, "Mgx: Near-zero overhead memory protection for data-intensive accelerators," in *Proceedings of the 49th Annual International Symposium on Computer Architecture*, ser. ISCA '22. New York, NY, USA: Association for Computing Machinery, 2022, p. 726–741. [Online]. Available: <https://doi.org/10.1145/3470496.3527418>
 - [181] —, "Guardnn: Secure accelerator architecture for privacy-preserving deep learning," in *Proceedings of the 59th ACM/IEEE Design Automation Conference*, ser. DAC '22. New York, NY, USA: Association for Computing Machinery, 2022, p. 349–354. [Online]. Available: <https://doi.org/10.1145/3489517.3530439>
 - [182] S. Lee, J. Kim, S. Na, J. Park, and J. Huh, "Tnpu: Supporting trusted execution with tree-less integrity protection for neural processing unit," in *2022 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2022, pp. 229–243.
 - [183] K. Lee, M. Yan, J. Emer, and A. Chandrakasan, "Secureloop: Design space exploration of secure dnn accelerators," in *Proceedings of the 56th Annual IEEE/ACM International Symposium on Microarchitecture*, ser. MICRO '23. New York, NY, USA: Association for Computing Machinery, 2023, p. 194–208. [Online]. Available: <https://doi.org/10.1145/3613424.3614273>
 - [184] S. Banerjee, S. Wei, P. Ramrakhiani, and M. Tiwari, "Triton: Software-defined threat model for secure multi-tenant ml inference accelerators," in *Proceedings of the 12th International Workshop on Hardware and Architectural Support for Security and Privacy*, ser. HASP '23. New York, NY, USA: Association for Computing Machinery, 2023, p. 19–28. [Online]. Available: <https://doi.org/10.1145/3623652.3623672>
 - [185] S. Na, S. Lee, Y. Kim, J. Park, and J. Huh, "Common counters: Compressed encryption counters for secure gpu memory," in *2021 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2021, pp. 1–13.
 - [186] S. Kim, H. Genc, V. V. Nikiforov, K. Asanović, B. Nikolić, and Y. S. Shao, "Moca: Memory-centric, adaptive execution for multi-tenant deep neural networks," in *2023 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2023, pp. 828–841.
 - [187] Q. Liu, W. Wen, and Y. Wang, "Concurrent weight encoding-based detection for bit-flip attack on neural network accelerators," in *Proceedings of the 39th International Conference on Computer-Aided Design*, ser. ICCAD '20. New York, NY, USA: Association for Computing Machinery, 2020. [Online]. Available: <https://doi.org/10.1145/3400302.3415726>
 - [188] Y. Cai, X. Chen, L. Tian, Y. Wang, and H. Yang, "Enabling secure in-memory neural network computing by sparse fast gradient encryption," in *2019 IEEE/ACM International Conference on Computer-Aided Design (ICCAD)*, 2019, pp. 1–8.
 - [189] P. Zuo, Y. Hua, L. Liang, X. Xie, X. Hu, and Y. Xie, "Sealing neural network models in encrypted deep learning accelerators," in *2021 58th ACM/IEEE Design Automation Conference (DAC)*, 2021, pp. 1255–1260.
 - [190] N. Lin, X. Chen, H. Lu, and X. Li, "Chaotic weights: A novel approach to protect intellectual property of deep neural networks," *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, vol. 40, no. 7, pp. 1327–1339, 2021.
 - [191] Y. Che and R. Wang, "Dnncloak: Secure dnn models against memory side-channel based reverse engineering attacks," in *2022 IEEE 40th International Conference on Computer Design (ICCD)*, 2022, pp. 89–96.
 - [192] B. F. Goldstein, V. C. Patil, V. C. Ferreira, A. S. Nery, F. M. G. França, and S. Kundu, "Preventing dnn model ip theft via hardware obfuscation," *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, vol. 11, no. 2, pp. 267–277, 2021.
 - [193] S. Maji, K. Lee, C. Gongye, Y. Fei, and A. P. Chandrakasan, "An Energy-Efficient Neural Network Accelerator with Improved Protec-

- tions Against Fault-Attacks,” in *European Solid State Circuits Conference*, 2023, pp. 233–236.
- [194] E. Ozen and A. Orailoglu, “Low-Cost Error Detection in Deep Neural Network Accelerators with Linear Algorithmic Checksums,” *J. Electron. Test.*, vol. 36, no. 6, pp. 703–718, 2020.
- [195] S. Burel, A. Evans, and L. Anghel, “Zero-Overhead Protection for CNN Weights,” in *IEEE International Symposium on Defect and Fault Tolerance in VLSI and Nanotechnology Systems*, 2021.
- [196] J. Li, A. S. Rakin, Z. He, D. Fan, and C. Chakrabarti, “RADAR: Runtime Adversarial Weight Attack Detection and Accuracy Recovery,” in *Design, Automation & Test in Europe Conference & Exhibition*, 2021, pp. 790–795.
- [197] E. Stefanov, M. van Dijk, E. Shi, C. Fletcher, L. Ren, X. Yu, and S. Devadas, “Path oram: an extremely simple oblivious ram protocol,” in *Proceedings of the 2013 ACM SIGSAC Conference on Computer & Communications Security*, ser. CCS ’13. New York, NY, USA: Association for Computing Machinery, 2013, p. 299–310. [Online]. Available: <https://doi.org/10.1145/2508859.2516660>
- [198] M. Maas, E. Love, E. Stefanov, M. Tiwari, E. Shi, K. Asanovic, J. Kubiatowicz, and D. Song, “Phantom: practical oblivious computation in a secure processor,” in *Proceedings of the 2013 ACM SIGSAC Conference on Computer & Communications Security*, ser. CCS ’13. New York, NY, USA: Association for Computing Machinery, 2013, p. 311–324. [Online]. Available: <https://doi.org/10.1145/2508859.2516692>
- [199] Y. Liu, D. Dachman-Soled, and A. Srivastava, “Mitigating reverse engineering attacks on deep neural networks,” in *2019 IEEE Computer Society Annual Symposium on VLSI (ISVLSI)*, 2019, pp. 657–662.
- [200] X. Wang, R. Hou, Y. Zhu, J. Zhang, and D. Meng, “Npufort: a secure architecture of dnn accelerator against model inversion attack,” in *Proceedings of the 16th ACM International Conference on Computing Frontiers*, ser. CF ’19. New York, NY, USA: Association for Computing Machinery, 2019, p. 190–196. [Online]. Available: <https://doi.org/10.1145/3310273.3323070>
- [201] T. Zhou, S. Ren, and X. Xu, “Obfunas: A neural architecture search-based dnn obfuscation approach,” in *Proceedings of the 41st IEEE/ACM International Conference on Computer-Aided Design*, ser. ICCAD ’22. New York, NY, USA: Association for Computing Machinery, 2022. [Online]. Available: <https://doi.org/10.1145/3508352.3549429>
- [202] W. Li, Y. Wang, H. Li, and X. Li, “P3m: a pim-based neural network model protection scheme for deep learning accelerator,” in *Proceedings of the 24th Asia and South Pacific Design Automation Conference*, ser. ASPDAC ’19. New York, NY, USA: Association for Computing Machinery, 2019, p. 633–638. [Online]. Available: <https://doi.org/10.1145/3287624.3287695>
- [203] A. Malhotra, C. Wang, and S. K. Gupta, “Bnn-flip: Enhancing the fault tolerance and security of compute-in-memory enabled binary neural network accelerators,” in *Proceedings of the 29th Asia and South Pacific Design Automation Conference*, ser. ASPDAC ’24. IEEE Press, 2024, p. 146–152. [Online]. Available: <https://doi.org/10.1109/ASP-DAC58780.2024.10473947>
- [204] M. Zou, J. Zhou, X. Cui, W. Wang, and S. Kvatinisky, “Enhancing security of memristor computing system through secure weight mapping,” in *2022 IEEE Computer Society Annual Symposium on VLSI (ISVLSI)*, 2022, pp. 182–187.
- [205] A. Dubey, A. Ahmad, M. A. Pasha, R. Cammarota, and A. Aysu, “Modulonet: Neural networks meet modular arithmetic for efficient hardware masking,” *IACR Transactions on Cryptographic Hardware and Embedded Systems*, vol. 2022, no. 1, p. 506–556, Nov. 2021. [Online]. Available: <https://tches.iacr.org/index.php/TCHES/article/view/9306>
- [206] A. Dubey and A. Aysu, “A full-stack approach for side-channel secure ml hardware,” in *2023 IEEE International Test Conference (ITC)*, 2023, pp. 186–195.
- [207] M. Brosch, M. Probst, M. Glaser, and G. Sigl, “A masked hardware accelerator for feed-forward neural networks with fixed-point arithmetic,” *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, vol. 32, no. 02, pp. 231–244, feb 2024.
- [208] A. Dubey, R. Cammarota, and A. Aysu, “Bomanet: boolean masking of an entire neural network,” in *Proceedings of the 39th International Conference on Computer-Aided Design*, ser. ICCAD ’20. New York, NY, USA: Association for Computing Machinery, 2020. [Online]. Available: <https://doi.org/10.1145/3400302.3415649>
- [209] S. Maji, U. Banerjee, S. H. Fuller, and A. P. Chandrakasan, “A threshold implementation-based neural network accelerator with power and electromagnetic side-channel countermeasures,” *IEEE Journal of Solid-State Circuits*, vol. 58, no. 1, pp. 141–154, 2023.
- [210] Q. Fang, L. Lin, H. Zhang, T. Wang, and M. Alioto, “Voltage scaling-agnostic counteraction of side-channel neural net reverse engineering via machine learning compensation and multi-level shuffling,” in *2023 IEEE Symposium on VLSI Technology and Circuits (VLSI Technology and Circuits)*, 2023, pp. 1–2.
- [211] O. Reparaz, B. Bilgin, S. Nikova, B. Gierlichs, and I. Verbauwhede, “Consolidating masking schemes,” *Cryptology ePrint Archive*, Paper 2015/719, 2015, <https://eprint.iacr.org/2015/719>. [Online]. Available: <https://eprint.iacr.org/2015/719>
- [212] S. Nikova, C. Rechberger, and V. Rijmen, “Threshold Implementations Against Side-Channel Attacks and Glitches,” in *ICICS*, vol. 4307. Springer, 2006, pp. 529–545.
- [213] S. Chari, C. S. Jutla, J. R. Rao, and P. Rohatgi, “Towards sound approaches to counteract power-analysis attacks,” in *Advances in Cryptology - CRYPTO ’99, 19th Annual International Cryptology Conference, Santa Barbara, California, USA, August 15-19, 1999, Proceedings*, ser. Lecture Notes in Computer Science, vol. 1666. Springer, 1999, pp. 398–412.
- [214] T. S. Messerges, “Securing the aes finalists against power analysis attacks,” in *Fast Software Encryption*, G. Goos, J. Hartmanis, J. van Leeuwen, and B. Schneier, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2001, pp. 150–164.
- [215] H. Gross, S. Mangard, and T. Korak, “Domain-oriented masking: Compact masked hardware implementations with arbitrary protection order,” in *Proceedings of the 2016 ACM Workshop on Theory of Implementation Security*, ser. TIS ’16. New York, NY, USA: Association for Computing Machinery, 2016, p. 3. [Online]. Available: <https://doi.org/10.1145/2996366.2996426>
- [216] J.-S. Coron, J. Großschädl, and P. K. Vadnala, “Secure conversion between boolean and arithmetic masking of any order,” in *CHES*. Springer, 2014, pp. 188–205. [Online]. Available: <https://www.iacr.org/archive/ches2014/87310157/87310157.pdf>
- [217] J.-S. Coron and L. Goubin, “On boolean and arithmetic masking against differential power analysis,” in *Cryptographic Hardware and Embedded Systems — CHES 2000*, Ç. K. Koç and C. Paar, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2000, pp. 231–237.
- [218] A. Dubey, R. Cammarota, V. Suresh, and A. Aysu, “Guarding machine learning hardware against physical side-channel attacks,” *J. Emerg. Technol. Comput. Syst.*, vol. 18, no. 3, apr 2022. [Online]. Available: <https://doi.org/10.1145/3465377>
- [219] T. Beyne and B. Bilgin, “Uniform first-order threshold implementations,” *Cryptology ePrint Archive*, Paper 2016/715, 2016, <https://eprint.iacr.org/2016/715>. [Online]. Available: <https://eprint.iacr.org/2016/715>
- [220] B. Bilgin, B. Gierlichs, S. Nikova, V. Nikov, and V. Rijmen, “Higher-order threshold implementations,” *Cryptology ePrint Archive*, Paper 2014/751, 2014, <https://eprint.iacr.org/2014/751>. [Online]. Available: <https://eprint.iacr.org/2014/751>
- [221] B. Bilgin, J. Daemen, V. Nikov, S. Nikova, V. Rijmen, and G. Van Assche, “Efficient and first-order dpa resistant implementations of keccak,” in *Smart Card Research and Advanced Applications*, A. Francillon and P. Rohatgi, Eds. Cham: Springer International Publishing, 2014, pp. 187–199.
- [222] A. Baksi, S. Bhasin, J. Breier, D. Jap, and D. Saha, “A Survey on Fault Attacks on Symmetric Key Cryptosystems,” *ACM Comput. Surv.*, vol. 55, no. 4, 2022.
- [223] C. Beierle, G. Leander, A. Moradi, and S. Rasoolzadeh, “Craft: Lightweight tweakable block cipher with efficient protection against dfa attacks,” *IACR Trans. Symmetric Cryptol.*, vol. 2019, pp. 5–45, 2019. [Online]. Available: <https://api.semanticscholar.org/CorpusID:75135098>
- [224] K. Matsuda, T. Fujii, N. Shoji, T. Sugawara, K. Sakiyama, Y.-I. Hayashi, M. Nagata, and N. Miura, “A 286F²/Cell Distributed Bulk-Current Sensor and Secure Flush Code Eraser Against Laser Fault Injection Attack,” in *IEEE International Solid-State Circuits Conference*, 2018, pp. 352–354.
- [225] R. Kumar, A. Varna, C. Tokunaga, S. Taneja, V. De, and S. Mathew, “A 100Gbps Fault-Injection Attack Resistant AES-256 Engine with 99.1-to-99.99% Error Coverage in Intel 4 CMOS,” in *IEEE International Solid-State Circuits Conference*, 2023.
- [226] S. Song, S. G. Tell, B. Zimmer, S. S. Kudva, N. Nedovic, and C. T. Gray, “An FLL-Based Clock Glitch Detector for Security Circuits in a 5nm FINFET Process,” in *IEEE Symposium on VLSI Technology and Circuits*, 2022, pp. 146–147.
- [227] S. Mathew, M. Anders, R. Krishnamurthy, and S. Borkar, “A 6.5GHz 54mW 64-bit Parity-Checking Adder for 65nm Fault-Tolerant Micro-

- processor Execution Cores,” in *IEEE Symposium on VLSI Circuits*, 2007, pp. 46–47.
- [228] M. Ashok, S. Maji, X. Zhang, J. Cohn, and A. P. Chandrakasan, “A secure digital in-memory compute (imc) macro with protections for side-channel and bus probing attacks,” in *2024 IEEE Custom Integrated Circuits Conference (CICC)*, 2024, pp. 1–2.
- [229] S. Skorobogatov, “How microprobing can attack encrypted memory,” in *2017 Euromicro Conference on Digital System Design (DSD)*, 2017, pp. 244–251.
- [230] C. Boit, C. Helfmeier, and U. Kerst, “Security risks posed by modern ic debug and diagnosis tools,” in *2013 Workshop on Fault Diagnosis and Tolerance in Cryptography*, 2013, pp. 3–11.
- [231] W. Li, S. Huang, X. Sun, H. Jiang, and S. Yu, “Secure-rram: A 40nm 16kb compute-in-memory macro with reconfigurability, sparsity control, and embedded security,” in *2021 IEEE Custom Integrated Circuits Conference (CICC)*, 2021, pp. 1–2.
- [232] S. Huang, H. Jiang, X. Peng, W. Li, and S. Yu, “Secure xor-cim engine: Compute-in-memory sram architecture with embedded xor encryption,” *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, vol. 29, no. 12, pp. 2027–2039, 2021.
- [233] W. Xiong, L. Ke, D. Jankov, M. Kounavis, X. Wang, E. Northup, J. A. Yang, B. Acun, C.-J. Wu, P. T. Peter Tang, G. Edward Suh, X. Zhang, and H.-H. S. Lee, “Secndp: Secure near-data processing with untrusted memory,” in *2022 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*, 2022, pp. 244–258.
- [234] S. Huang, X. Peng, H. Jiang, Y. Luo, and S. Yu, “New security challenges on machine learning inference engine: Chip cloning and model reverse engineering,” 2020.
- [235] M. Zou, Z. Zhu, Y. Cai, J. Zhou, C. Wang, and Y. Wang, “Security enhancement for rram computing system through obfuscating crossbar row connections,” in *2020 Design, Automation & Test in Europe Conference & Exhibition (DATE)*, 2020, pp. 466–471.
- [236] C. Yang, B. Liu, H. Li, Y. Chen, M. Barnell, Q. Wu, W. Wen, and J. Rajendran, “Thwarting replication attack against memristor-based neuromorphic computing system,” *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, vol. 39, no. 10, pp. 2192–2205, 2020.
- [237] R. S. Chakraborty, S. Narasimhan, and S. Bhunia, “Hardware trojan: Threats and emerging solutions,” in *2009 IEEE International High Level Design Validation and Test Workshop*. San Francisco, CA USA: IEEE, 2009, pp. 166–171.
- [238] R. S. Chakraborty, F. Wolff, S. Paul, C. Papachristou, and S. Bhunia, “Mero: A statistical approach for hardware trojan detection,” in *Proceedings of the 11th International Workshop on Cryptographic Hardware and Embedded Systems*. Lausanne, Switzerland: Springer, 2009, p. 396–410.
- [239] K. Xiao, D. Forte, Y. Jin, R. Karri, S. Bhunia, and M. M. Tehranipoor, “Hardware Trojans: Lessons Learned after One Decade of Research,” *ACM Transactions on Design Automation of Electronic Systems*, vol. 22, no. 1, pp. 1–23, May 2016.
- [240] J. Balasch, B. Gierlichs, and I. Verbauwhede, “Electromagnetic circuit fingerprints for Hardware Trojan detection,” in *2015 IEEE International Symposium on Electromagnetic Compatibility (EMC)*. Dresden, Germany: IEEE, Aug. 2015, pp. 246–251.
- [241] S. Narasimhan, D. Du, R. S. Chakraborty, S. Paul, F. G. Wolff, C. A. Papachristou, K. Roy, and S. Bhunia, “Hardware Trojan Detection by Multiple-Parameter Side-Channel Analysis,” *IEEE Transactions on Computers*, vol. 62, no. 11, pp. 2183–2195, Nov. 2013.
- [242] S. Bhunia, M. S. Hsiao, M. Banga, and S. Narasimhan, “Hardware trojan attacks: Threat analysis and countermeasures,” *Proceedings of the IEEE*, vol. 102, no. 8, pp. 1229–1247, Aug. 2014.
- [243] Y. Qin and T. Xia, “Sensitivity analysis of ring oscillator based hardware Trojan detection,” in *2017 IEEE 17th International Conference on Communication Technology (ICCT)*. Chengdu, China: IEEE, Oct. 2017, pp. 1979–1983.
- [244] N. B. Gunti and K. Lingasubramanian, “Efficient static power based side channel analysis for Hardware Trojan detection using controllable sleep transistors,” in *SoutheastCon 2015*. Fort Lauderdale, FL, USA: IEEE, Apr. 2015, pp. 1–6.
- [245] M. Ashok, M. J. Turner, R. L. Walsworth, E. V. Levine, and A. P. Chandrakasan, “Hardware trojan detection using unsupervised deep learning on quantum diamond microscope magnetic field images,” *J. Emerg. Technol. Comput. Syst.*, vol. 18, no. 4, oct 2022. [Online]. Available: <https://doi.org/10.1145/3531010>
- [246] K. I. Gubbi, I. Kaur, A. Hashem, S. M. P. D. H. Homayoun, A. Sasan, and S. Salehi, “Securing ai hardware: Challenges in detecting and mitigating hardware trojans in ml accelerators,” in *2023 IEEE 66th International Midwest Symposium on Circuits and Systems (MWSCAS)*, 2023, pp. 821–825.
- [247] T. A. Odetola, H. R. Mohammed, and S. R. Hasan, “A stealthy hardware trojan exploiting the architectural vulnerability of deep learning architectures: Input interception attack (iia),” 2021.
- [248] P. Sun, B. Halak, and T. Kazmierski, “Towards hardware trojan resilient design of convolutional neural networks,” in *2022 IEEE 35th International System-on-Chip Conference (SOCC)*, 2022, pp. 1–6.
- [249] R. W. Hamming, “Error detecting and error correcting codes,” *The Bell System Technical Journal*, vol. 29, no. 2, pp. 147–160, 1950.
- [250] C. L. Chen and M. Y. Hsiao, “Error-correcting codes for semiconductor memory applications: A state-of-the-art review,” *IBM Journal of Research and Development*, vol. 28, no. 2, pp. 124–134, 1984.
- [251] H. Yu, P. Sun, B. Halak, K. Shanthakumar, and T. Kazmierski, “Tamper resistant design of convolutional neural network hardware accelerator,” in *2023 Asian Hardware Oriented Security and Trust Symposium (AsianHOST)*, 2023, pp. 1–5.
- [252] M. I. Mera Collantes, Z. Ghodsi, and S. Garg, “Safetpu: A verifiably secure hardware accelerator for deep neural networks,” in *2020 IEEE 38th VLSI Test Symposium (VTS)*, 2020, pp. 1–6.
- [253] J. Thaler, “Time-optimal interactive proofs for circuit evaluation,” in *Advances in Cryptology – CRYPTO 2013*, R. Canetti and J. A. Garay, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2013, pp. 71–89.
- [254] L. de Castro, R. Agrawal, R. Yazicigil, A. Chandrakasan, V. Vaikuntanathan, C. Juvekar, and A. Joshi, “Does fully homomorphic encryption need compute acceleration?” 2021. [Online]. Available: <https://arxiv.org/abs/2112.06396>
- [255] F. Turan, S. S. Roy, and I. Verbauwhede, “Heaws: An accelerator for homomorphic encryption on the amazon aws fpga,” *IEEE Transactions on Computers*, vol. 69, no. 8, pp. 1185–1196, 2020.
- [256] S. Halevi and V. Shoup, “Algorithms in helib,” in *Advances in Cryptology – CRYPTO 2014*, J. A. Garay and R. Gennaro, Eds. Berlin, Heidelberg: Springer Berlin Heidelberg, 2014, pp. 554–571.
- [257] Z. Brakerski, C. Gentry, and V. Vaikuntanathan, “(Leveled) fully homomorphic encryption without bootstrapping,” in *ITCS ’12*, 2012.
- [258] S. Halevi, Y. Polyakov, and V. Shoup, “An improved rms variant of the bfv homomorphic encryption scheme,” in *Topics in Cryptology – CT-RSA 2019 - The Cryptographers’ Track at the RSA Conference 2019, Proceedings*, M. Matsui, Ed. Germany: Springer Verlag, 2019, pp. 83–105.
- [259] J. H. Cheon, A. Kim, M. Kim, and Y. Song, “Homomorphic encryption for arithmetic of approximate numbers,” in *Advances in Cryptology – ASIACRYPT 2017*, T. Takagi and T. Peyrin, Eds. Cham: Springer International Publishing, 2017, pp. 409–437.
- [260] I. Chillotti, N. Gama, M. Georgieva, and M. Izabachène, “Tfhe: fast fully homomorphic encryption over the torus,” *Journal of Cryptology*, vol. 33, no. 1, pp. 34–91, 2020.
- [261] W. Jung, S. Kim, J. H. Ahn, J. H. Cheon, and Y. Lee, “Over 100x faster bootstrapping in fully homomorphic encryption through memory-centric optimization with gpus,” *IACR Transactions on Cryptographic Hardware and Embedded Systems*, vol. 2021, no. 4, p. 114–148, Aug. 2021. [Online]. Available: <https://tches.iacr.org/index.php/TCHES/article/view/9062>
- [262] H. Chen, I. Chillotti, and Y. Song, “Improved bootstrapping for approximate homomorphic encryption,” in *Advances in Cryptology – EUROCRYPT 2019*, Y. Ishai and V. Rijmen, Eds. Cham: Springer International Publishing, 2019, pp. 34–54.
- [263] K. Han and D. Ki, “Better bootstrapping for approximate homomorphic encryption,” in *Topics in Cryptology – CT-RSA 2020*, S. Jarecki, Ed. Cham: Springer International Publishing, 2020, pp. 364–390.
- [264] R. Agrawal, L. de Castro, C. Juvekar, A. Chandrakasan, V. Vaikuntanathan, and A. Joshi, “Mad: Memory-aware design techniques for accelerating fully homomorphic encryption,” *MICRO ’23, October 28–November 1, 2023, Toronto, ON, Canada*, 2023.
- [265] R. Agrawal and A. Joshi, *On architecting fully homomorphic encryption-based computing systems*. Springer, 2023.
- [266] J.-P. Bossuat, C. Mouchet, J. Troncoso-Pastoriza, and J.-P. Hubaux, “Efficient bootstrapping for approximate homomorphic encryption with non-sparse keys,” in *Advances in Cryptology – EUROCRYPT 2021*, A. Canteaut and F.-X. Standaert, Eds. Cham: Springer International Publishing, 2021, pp. 587–617.
- [267] K. Shrivdhar, Y. Bao, R. Agrawal, M. Shen, G. Jonatan, E. Mora, A. Ingare, N. Livesay, J. L. Abellán, J. Kim, A. Joshi, and D. Kaeli, “Gme: Gpu-based microarchitectural extensions to accelerate homomorphic encryption,” in *Proceedings of the 56th Annual IEEE/ACM International Symposium on Microarchitecture*, 2023, p. 670–684.

- [268] R. Agrawal, L. de Castro, G. Yang, C. Juvekar, R. Yazicigil, A. Chandrakasan, V. Vaikuntanathan, and A. Joshi, "Fab: An fpga-based accelerator for bootstrappable fully homomorphic encryption," in *2023 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*. IEEE, 2023, pp. 882–895.
- [269] N. Samardzic, A. Feldmann, A. Krastev, S. Devadas, R. Dreslinski, C. Peikert, and D. Sanchez, "F1: A fast and programmable accelerator for fully homomorphic encryption," in *MICRO-54: 54th Annual IEEE/ACM International Symposium on Microarchitecture*, ser. MICRO '21. New York, NY, USA: Association for Computing Machinery, 2021, p. 238–252. [Online]. Available: <https://doi.org/10.1145/3466752.3480070>
- [270] S. Kim, J. Kim, M. J. Kim, W. Jung, M. Rhu, J. Kim, and J. H. Ahn, "Bts: An accelerator for bootstrappable fully homomorphic encryption," *arXiv preprint arXiv:2112.15479*, 2021.
- [271] N. Samardzic, A. Feldmann, A. Krastev, N. Manohar, N. Genise, S. Devadas, K. Eldefrawy, C. Peikert, and D. Sanchez, "Craterlake: a hardware accelerator for efficient unbounded computation on encrypted data," in *Proceedings of the 49th Annual International Symposium on Computer Architecture*, 2022, pp. 173–187.
- [272] J. Kim, G. Lee, S. Kim, G. Sohn, M. Rhu, J. Kim, and J. H. Ahn, "Ark: Fully homomorphic encryption accelerator with runtime data generation and inter-operation key reuse," in *2022 55th IEEE/ACM International Symposium on Microarchitecture (MICRO)*, 2022, pp. 1237–1254.
- [273] J. Kim, S. Kim, J. Choi, J. Park, D. Kim, and J. H. Ahn, "Sharp: A short-word hierarchical accelerator for robust and practical fully homomorphic encryption," in *Proceedings of the 50th Annual International Symposium on Computer Architecture*, 2023, pp. 1–15.
- [274] M. S. Riaz, K. Laine, B. Pelton, and W. Dai, "Heax: An architecture for computing on encrypted data," in *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems*, 2020, pp. 1295–1309.
- [275] G. Kokolakis, A. Moschos, and A. D. Keromytis, "Harnessing the power of general-purpose llms in hardware trojan design," in *Applied Cryptography and Network Security Workshops: ACNS 2024 Satellite Workshops, AIBlock, AIHWS, AIoTS, SCI, AAC, SiMLA, LLE, and CIMSS, Abu Dhabi, United Arab Emirates, March 5–8, 2024, Proceedings, Part I*. Berlin, Heidelberg: Springer-Verlag, 2024, p. 176–194. [Online]. Available: https://doi.org/10.1007/978-3-031-61486-6_11
- [276] B. Ahmad, S. Thakur, B. Tan, R. Karri, and H. Pearce, "On hardware security bug code fixes by prompting large language models," *IEEE Transactions on Information Forensics and Security*, vol. 19, pp. 4043–4057, 2024.
- [277] D. Saha, K. Yahyaei, S. K. Saha, M. Tehranipoor, and F. Farahmandi, "Empowering hardware security with llm: The development of a vulnerable hardware database," in *2024 IEEE International Symposium on Hardware Oriented Security and Trust (HOST)*. Los Alamitos, CA, USA: IEEE Computer Society, may 2024, pp. 233–243. [Online]. Available: <https://doi.ieeecomputersociety.org/10.1109/HOST55342.2024.10545393>



Kyungmi Lee (Member, IEEE) received the B.S. degree in Electrical and Computer Engineering from Seoul National University, Seoul, South Korea in 2018, and the S.M. and the Ph.D. degree in Electrical Engineering and Computer Science from the Massachusetts Institute of Technology (MIT), Cambridge, MA, USA in 2020 and 2024, respectively.

She is currently a Postdoctoral Associate at MIT, supervised by Professor Anantha P. Chandrakasan.

Her research focuses on hardware security of deep neural network accelerators and energy-efficient hardware and software design for machine learning applications.

She received the MIT Jacobs Presidential Fellowship in 2018, the Siebel Scholars in 2020, and the Doctoral Fellowship from Korea Foundation for Advanced Studies from 2018 to 2023. She received the Bob Owens Best Student Paper Award at the IEEE International Workshop on Signal Processing Systems in 2021.



Maitreyi Ashok (Student Member, IEEE) received the B.S. degree in electrical engineering from the California Institute of Technology (Caltech), Pasadena, CA, USA in 2019, and the S.M. degree in electrical engineering and computer science from the Massachusetts Institute of Technology (MIT), Cambridge, MA, USA in 2021.

She is currently pursuing her PhD at MIT, supervised by Professor Anantha Chandrakasan. Her research focuses on hardware security in digital and mixed-signal integrated circuits against physical attacks including passive side-channels and hardware trojan insertion. She has previously interned at Intel Labs, Hillsboro, OR, USA in 2021 and Microsoft, Redmond, WA, USA in 2018 and 2019.

Maitreyi received the National Science Foundation Graduate Research Fellowship in 2019, the Analog Devices Fellowship in 2019, and the MathWorks Engineering Fellowship in 2021.



Saurav Maji (Member, IEEE) received the B.Tech. degree in Electronics and Electrical Communication Engineering from the Indian Institute of Technology (IIT) Kharagpur, India, in 2017, and the S.M. and Ph.D. degrees in Electrical Engineering and Computer Science from the Massachusetts Institute of Technology (MIT) in 2019 and 2023, respectively. Since July 2023, he has been affiliated as a Research Scientist with Circuits Research Lab, Intel Labs, Intel, Hillsboro, OR 97124, USA. His research interests include energy-efficient security solutions for IoT applications, side-channel security, and fault analysis for Machine Learning and Artificial Intelligence applications, and circuits for biomedical applications.

Dr. Maji was recipient of the President of India Gold Medal from IIT Kharagpur in 2017, the Analog Devices Fellowship from MIT in 2017, and the MIT EECS MathWorks Fellowship from MIT in 2022. He received the 2nd Place Winner Award in the Design Contest at the 2020 ACM/IEEE International Symposium on Low Power Electronics and Design (ISLPED). Additionally, he is the recipient of the Student Research Preview (SRP) Poster Award at the IEEE International Solid-State Circuits Conference (ISSCC) 2023.



Rashmi Agrawal (Member, IEEE) currently works as a Research Scientist at MIT and Boston University. She is also a co-founder of CipherSonic Labs, a venture dedicated to advancing data security and privacy solutions. Previously, Rashmi worked as a Research Scientist in Intel Labs' Security and Privacy Research group. She holds a PhD in Computer Engineering from Boston University, where her research focused on designing hardware accelerators for post-quantum cryptography and fully homomorphic encryption. This work aimed to enable post-

quantum secure, practical, and cost-effective private data analytics in cloud environments. Rashmi was honored with the EECS Rising Star Award in 2023 for her PhD research. In addition, she has received the best paper awards at ICCD 2020 and GLSVLSI 2020. During her PhD, Rashmi gained industry experience at Xilinx (AMD) and Analog Devices.



Ajay Joshi (Senior Member, IEEE) is a Professor in the ECE department at Boston University. Before joining BU, he got his PhD at Georgia Tech and then worked as a postdoc at MIT. He has worked as a Visiting Researcher at Google and as an Architect at Lightmatter Inc. He recently co-founded CipherSonic Labs Inc., which is focused on commercializing FHE-based computing. His research is in the areas of computer architecture and digital VLSI with a focus on security and privacy, machine learning, and photonic computing. He received the

NSF CAREER Award in 2012, the Boston University ECE Department's Award for Excellence in Teaching in 2014, Best Paper Awards at ASIACCS 2018 and HOST 2023, and Google Faculty Research Awards in 2018 and 2019. He currently serves as the Associate Editor for IEEE Transactions on VLSI Systems.



Mengjia Yan (Member, IEEE) is an Associate Professor in the Electrical Engineering and Computer Science Department, Massachusetts Institute of Technology. She received her Ph.D. degree from the University of Illinois at Urbana-Champaign (UIUC). Her research interests lie in computer architecture and hardware security, with a focus on side channel attacks and defenses. Mengjia received the IEEE/TCCA Young Computer Architect award in 2024, the NSF CAREER Award in 2021, the Intel Rising Star Faculty Award in 2022. Her work has

been recognized as an ACM Research Highlight, multiple MICRO TopPicks in Computer Architecture, and a MICRO best paper award.



Joel S. Emer (Fellow, IEEE) received the B.S. (honors) and M.S. degrees in electrical engineering from Purdue University, West Lafayette, IN, USA, in 1974 and 1975, respectively, and the Ph.D. degree in electrical engineering from the University of Illinois at Urbana-Champaign, Champaign, IL, USA, in 1979.

He was with Intel, where he was an Intel Fellow and the Director of Microarchitecture Research. At Intel, he led the VSSAD Group, which he had previously been a member of at Compaq and Digital

Equipment Corporation. He is currently a Senior Distinguished Research Scientist with Nvidia Architecture Research Group, Westford, MA, USA, where he is responsible for exploration of future architectures and modeling and analysis methodologies. He is also a Professor of the Practice at the Massachusetts Institute of Technology, Cambridge, MA, USA, where he teaches computer architecture and supervises graduate students. He has held various research and advanced development positions investigating processor microarchitecture and developing performance modeling and evaluation techniques. He has made architectural contributions to a number of VAX, Alpha, and X86 processors and is recognized as one of the developers of the widely employed quantitative approach to processor performance evaluation. He has been recognized for his contributions in the advancement of simultaneous multithreading technology, processor reliability analysis, cache organization, and spatial architectures for deep learning.

Dr. Emer is a Fellow of the Association for Computing Machinery (ACM). He has been a recipient of numerous public recognitions. In 2009, he received the Eckert-Mauchly Award and in 2023 the B Ramakrishna Rau Award for lifetime contributions in computer architecture, the Purdue University Outstanding Electrical and Computer Engineer Alumni Award, and the University of Illinois Electrical and Computer Engineering Distinguished Alumni Award in 2010 and 2011, respectively. His 1996 paper on simultaneous multithreading received the ACM/SIGARCH-IEEE-CS/TCCA: Most Influential Paper Award in 2011. He was named to the ISCA and Micro Halls of Fame in 2005 and 2015, respectively. He has had six papers selected for the IEEE Micro Top Picks in Computer Architecture, in 2003, 2004, 2007, 2013, 2015 and 2016. He was the Program Chair of ISCA in 2000 and of Micro in 2017.



Anantha P. Chandrakasan (Fellow, IEEE) received the B.S., M.S. and Ph.D. degrees in Electrical Engineering and Computer Sciences from the University of California, Berkeley, in 1989, 1990, and 1994 respectively. Since September 1994, he has been with the Massachusetts Institute of Technology, Cambridge, where he is currently the Vannevar Bush Professor of Electrical Engineering and Computer Science.

He was a co-recipient of several awards including the 2007 ISSCC Beatrice Winner Award for Editorial Excellence and the ISSCC Jack Kilby Award for Outstanding Student Paper (2007, 2008, 2009). He received the 2009 Semiconductor Industry Association (SIA) University Researcher Award and the 2013 IEEE Donald O. Pederson Award in Solid-State Circuits, an honorary doctorate from KU Leuven in 2016, the UC Berkeley EE Distinguished Alumni Award in 2017, and the 2019 IEEE Solid-State Circuits Society Distinguished Service Award. In 2015 he was elected to the National Academy of Engineering and in 2019 he was elected to the American Academy of Arts & Sciences. He was elected as fellow of the Association for Computing Machinery (ACM) in 2020. He is the recipient of the 2022 IEEE Mildred Dresselhaus Medal.

His research interests include ultra-low-power circuit and system design, energy harvesting, energy efficient RF circuits, and hardware security. He is a co-author of Low Power Digital CMOS Design (Kluwer Academic Publishers, 1995), Digital Integrated Circuits (Pearson Prentice-Hall, 2003, 2nd edition), and Sub-threshold Design for Ultra-Low Power Systems (Springer 2006).

He is an IEEE Fellow. He has served in various roles for the IEEE ISSCC including Program Chair, Signal Processing Sub-committee Chair, and Technology Directions Sub-committee Chair. He served as the Conference Chair of ISSCC from 2010 till 2018. He served as the Senior Technical Advisor to the Conference from ISSCC 2019 till ISSCC 2023. He was the Director of the MIT Microsystems Technology Laboratories from 2006 to 2011. From July 2011 - June 2017, he was head of the MIT Department of Electrical Engineering and Computer Science. Since July 2017, he has been Dean of the MIT School of Engineering and since January 2024, he has also held the position of MIT Chief Innovation and Strategy Officer.