



Accelerating Sparse MTTKRP for Tensor Decomposition on FPGA

Sasindu Wijeratne

University of Southern California, USA
kangaram@usc.edu

Rajgopal Kannan

DEVCOM US Army Research Lab, USA
rajgopal.kannan.civ@army.mil

Ta-Yang Wang

University of Southern California, USA
tayangwa@usc.edu

Viktor Prasanna

University of Southern California, USA
prasanna@usc.edu

ABSTRACT

Sparse Matricized Tensor Times Khatri-Rao Product (spMTTKRP) is the most computationally intensive kernel in sparse tensor decomposition. In this paper, we propose a hardware-algorithm co-design on FPGA to minimize the execution time of spMTTKRP along all modes of an input tensor. We introduce FLYCOO, a novel tensor format that eliminates the communication of intermediate values to the FPGA external memory during the computation of spMTTKRP along all the modes. Our remapping of the tensor using FLYCOO also balances the workload among multiple Processing Engines (PEs). We propose a parallel algorithm that can concurrently process multiple partitions of the input tensor independent of each other. The proposed algorithm also orders the tensor dynamically during runtime to increase the data locality of the external memory accesses. We develop a custom FPGA accelerator design with (1) PEs consisting of a collection of pipelines that can concurrently process multiple elements of the input tensor and (2) memory controllers to exploit the spatial and temporal locality of the external memory accesses of the computation. Our work achieves a geometric mean of 8.8× and 3.8× speedup in execution time compared with the state-of-the-art CPU and GPU implementations on widely-used real-world sparse tensor datasets.

CCS CONCEPTS

• Computer systems organization → Reconfigurable computing; Special purpose systems.

KEYWORDS

Tensor Decomposition, Sparse MTTKRP, FPGA, Hardware Accelerators

ACM Reference Format:

Sasindu Wijeratne, Ta-Yang Wang, Rajgopal Kannan, and Viktor Prasanna. 2023. Accelerating Sparse MTTKRP for Tensor Decomposition on FPGA. In Proceedings of the 2023 ACM/SIGDA International Symposium on Field Programmable Gate Arrays (FPGA '23), February 12–14, 2023, Monterey, CA, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3543622.3573179>



This work is licensed under a Creative Commons Attribution International 4.0 License.

FPGA '23, February 12–14, 2023, Monterey, CA, USA
© 2023 Copyright held by the owner/author(s).
ACM ISBN 978-1-4503-9417-8/23/02.
<https://doi.org/10.1145/3543622.3573179>

1 INTRODUCTION

Tensor Decomposition (TD) transforms input tensors into a reduced latent space which can then be leveraged to learn salient features of the underlying data distribution. TD has been successfully employed in many fields including machine learning [2, 15, 18], signal processing [22], and network analysis [3]. One popular TD algorithm is Canonical Polyadic Decomposition via alternating least squares (CP-ALS) [6], which can be accelerated by optimizing spMTTKRP, its most expensive computation task. Since real-world tensors are generally sparse, keeping only nonzero values in memory is a natural way to reduce the memory footprint. However, there is a concomitant need to develop optimized sparse tensor formats that support highly irregular data access patterns of spMTTKRP [23].

Several tensor formats have been proposed for real-world sparse tensors [5, 12, 14, 16, 20]. Many of the proposed formats seek to alleviate the problem of irregular data access by either using the number of tensor copies proportional to the number of modes or through additional memory to save intermediate by-products of the computation. However, as the number of modes and the sparsity of the tensor increases, these approaches might introduce significant amount of memory overhead. One desirable solution to reduce the memory traffic to the external memory is to reduce the number of accesses to the external memory by increasing the data reuse. A cache can be used as an intermediate memory to achieve this goal. In addition, reordering the tensor based on space-filling curves has shown promising results. Tensor formats such as HiCOO [12] and ALTO [5] use variations of Z-Morton data ordering [12]. It brings the tensor elements with neighboring coordinates closer and increases the data reuse. However, these tensor formats still communicate a significant amount of intermediate values to the external memory. In this paper, we propose FLYCOO, a tensor format that eliminates the communication of intermediate values to the FPGA external memory during the execution time of spMTTKRP along all the modes of the input tensor. It also increases the locality of data used in spMTTKRP. Our work hides the cost of on-the-fly tensor remapping time by overlapping the spMTTKRP computations with the tensor remapping.

The key contributions of this paper are:

- We introduce FLYCOO, a novel tensor format that eliminates the communication of intermediate values generated during spMTTKRP computation to the FPGA external memory. FLYCOO balances the workload among multiple Processing Engines along all the output modes of the input tensor.
- We develop a novel FPGA accelerator for spMTTKRP, which consists of (1) PEs with multiple pipelines to concurrently

process input tensors in a streaming fashion, (2) cache sub-systems to exploit the locality of input factor matrices, (3) Direct Memory Access (DMA) buffers to load the input tensor partitions and store the output factor matrices in a streaming fashion, and (4) custom hardware to support on-the-fly tensor remapping.

- The memory controller design introduced in this paper reduces the total execution time by 4.5× compared with the traditional Direct Memory Access (DMA) buffer-based memory controllers. The memory controller further reduces 67% of the pipeline stalls caused by FPGA external memory accesses.
- On widely used real-world large tensor datasets, our FPGA accelerator achieves a geometric mean of 8.8× and 3.8× improvement in total execution time compared with the state-of-the-art CPU and GPU implementations, respectively.

2 BACKGROUND

2.1 Notations

An $\square$ -dimensional, real-valued sparse tensor is denoted by $X \in \mathbb{R}^{\square_0 \times \dots \times \square_{\square-1}}$. For a thorough review of tensors and tensor algebra, please refer to [10]. Table 1 summarizes the list of symbols used in this paper.

Symbol	Description
α	scalar (lowercase letter)
$\mathbf{v}$	vector (bold lowercase letter)
$\mathbf{M}$	matrix (bold capital letter)
$\mathbf{x}$	sparse tensor (Euler script letter)
$\square$	number of modes
$X_{(\square)}$	mode- $\square$ matricization of the tensor X
$\circ$	vector outer product
$\otimes$	Kronecker product
$\boxtimes$	Khatri-Rao product

2.2 Tensor Decomposition

Canonical Polyadic Decomposition (CPD) [10] on sparse tensors decomposes a sparse tensor X into a sum of dense matrices for each mode which best approximates the original sparse tensor. For example, given a sparse tensor $X \in \mathbb{R}^{\square_0 \times \dots \times \square_{\square-1}}$, our goal is to express it as $X \approx \sum_{r=0}^{\square-1} \alpha_r \cdot \mathbf{a}_r \circ \mathbf{b}_r \circ \mathbf{c}_r$, where $\square \in \mathbb{Z}^+$ refers to the rank of X , which is defined as the smallest sum of rank-one tensors required to generate X , $\mathbf{a}_r \in \mathbb{R}^{\square_0}$, $\mathbf{b}_r \in \mathbb{R}^{\square_1}$, and $\mathbf{c}_r \in \mathbb{R}^{\square_2}$ for $r = 0, \dots, \square-1$. For illustration purposes, we assume the number of modes of tensors to be three in the rest of the section. The components of the above summation can be expressed as input factor matrices, i.e., $\mathbf{A} = [\mathbf{a}_0, \dots, \mathbf{a}_{\square-1}]$ and likewise for $\mathbf{B}$ and $\mathbf{C}$. It is often useful to constrain the components to unit length, factoring the weights into the vector $\alpha = [\alpha_0, \dots, \alpha_{\square-1}] \in \mathbb{R}^{\square}$, which allows to concisely express the model as $X \approx [(\square \mathbf{A}, \mathbf{B}, \mathbf{C})] = \sum_{r=0}^{\square-1} \alpha_r \cdot \mathbf{a}_r \circ \mathbf{b}_r \circ \mathbf{c}_r$.

Since the problem is non-convex and has no closed-form solution, existing methods for this optimization problem rely on iterative schemes. The alternating least squares algorithm for computing the CP decomposition (CP-ALS) is the most popular method due to its simplicity and efficiency. Algorithm 1 shows a common formulation of CP-ALS for 3-mode tensors. It consists of several iterations of the sparse Matricized Tensor-Times Khatri-Rao Product (spMTTKRP) operation on each mode.

As illustrated in Algorithm 1, executing spMTTKRP in each mode is the most expensive operation of CP-ALS. spMTTKRP involves the mode- $\square$ matricization $X_{(\square)}$ and the Khatri-Rao product [5] — given two matrices $\mathbf{B} \in \mathbb{R}^{\square_1 \times \square}$ and $\mathbf{C} \in \mathbb{R}^{\square_2 \times \square}$, their Khatri-Rao product $\mathbf{B} \boxtimes \mathbf{C} = [\mathbf{b}_1 \boxtimes \mathbf{c}_1 \quad \mathbf{b}_2 \boxtimes \mathbf{c}_2 \dots \mathbf{b}_{\square} \boxtimes \mathbf{c}_{\square}]$. spMTTKRP can be expressed as

$$\text{spMTTKRP}(X_{(\square)}, \mathbf{B}, \mathbf{C}) = X_{(\square)} (\mathbf{B} \boxtimes \mathbf{C}).$$

Algorithm 1: CP-ALS for the 3-mode tensors

- 1 Input: A tensor $X \in \mathbb{R}^{\square_0 \times \square_1 \times \square_2}$, the rank $\square \in \mathbb{Z}^+$
- 2 Output: CP decomposition $[[\square \mathbf{A}, \mathbf{B}, \mathbf{C}]]$, $\square \in \mathbb{R}^{\square}$, $\mathbf{A} \in \mathbb{R}^{\square_0 \times \square}$, $\mathbf{B} \in \mathbb{R}^{\square_1 \times \square}$, $\mathbf{C} \in \mathbb{R}^{\square_2 \times \square}$
- 3 while stopping criterion not met do
- 4 $\mathbf{A} \leftarrow \text{spMTTKRP}(X_{(0)}, \mathbf{B}, \mathbf{C})$
- 5 $\mathbf{B} \leftarrow \text{spMTTKRP}(X_{(1)}^1, \mathbf{A}, \mathbf{C})$
- 6 $\mathbf{C} \leftarrow \text{spMTTKRP}(X_{(2)}^2, \mathbf{A}, \mathbf{B})$
- 7 Normalize $\mathbf{A}, \mathbf{B}, \mathbf{C}$ and store the norms as $\square$

3 RELATED WORK

Srivastava et al. [21] propose a custom CGRA fabric to accelerate sparse computations, including spMTTKRP. The authors use a mode-specific tensor format to compute spMTTKRP. In our work, we propose a mode-agnostic tensor format that only requires one additional tensor copy regardless of the number of modes of the tensor.

Nisa et al. [17] optimize MTTKRP on GPUs. They propose a tensor slicing technique for the load balancing between GPU warps. We develop a simple load balancing scheme that fairly shares the total workload between ACCELS (see Section 4.6).

There are several CPU-based MTTKRP acceleration algorithms proposed in the literature. J. Li et al. propose HiCOO [12], a block-based format that compresses the sparse tensor. Helal et al. propose ALTO [5], a space-filling curve-based tensor ordering method that can efficiently encode spaces with irregular shapes. ALTO requires the least amount of external memory to store tensors. Unlike prior formats, we perform on-the-fly memory layout remapping to reduce external memory communication time.

4 ACCELERATOR DESIGN

4.1 Target Platform

We develop a hardware accelerator for a data center FPGA device directly connected to external DRAM memory.

4.2 Hypergraph Representation of Tensors

The hypergraph representation of a tensor has been used to describe the spMTTKRP operation in the literature [8, 13]. To describe our proposed tensor format based on the hypergraph, we will briefly introduce the hypergraph representation of a tensor in this section. For a given tensor X with $\square$ modes, we build a hypergraph $\mathcal{H} = (\mathcal{V}, \mathcal{E})$ with the vertex

set $\mathcal{V}$ and the hyperedge set $\mathcal{E}$ as follows: $\mathcal{V} = \{i_0, i_1, i_2, \dots, i_{\square-1}\}$, where $i_{\square}$ is the set of all the tensor indices in mode $\square$ and contains

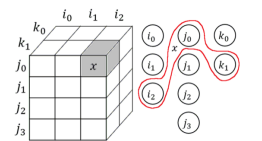


Figure 1: Hypergraph of a sparse tensor

hyperedges that represent the nonzero tensor elements in X . For a 3-mode sparse tensor $X \in \mathbb{R}^{I_1 \times I_2 \times I_3}$ with N nonzero tensor elements, its hypergraph $\mathcal{H} = (\mathcal{V}, \mathcal{E})$ consists of $|\mathcal{V}| = I_1 + I_2 + I_3$ vertices and $|\mathcal{E}| = N$ hyperedges. A hyperedge $X(i, j, k)$ connects the three vertices i , j , and k which correspond to the row indices of the factor matrices. Figure 1 shows an example of the hypergraph for a sparse tensor.

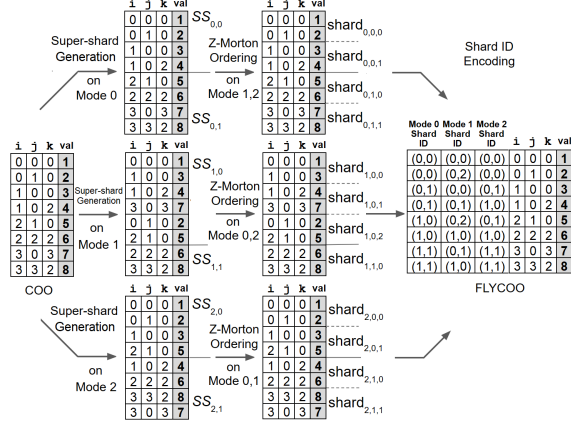


Figure 2: Example FLYCOO Format Generation

4.3 Tensor Format

During tensor decomposition, spMTTKRP is computed along each mode, one mode after the other (see Algorithm 1). For an n -mode tensor, when computing spMTTKRP for mode o we refer to mode o as the output mode and its corresponding factor matrix as the output factor matrix. Meanwhile, the rest of the modes becomes input modes, and their factor matrices become input factor matrices.

For each mode, the FLYCOO format assigns each nonzero tensor element to a tensor partition. Then embed the partition ids to each tensor element. Figure 2 illustrates the tensor format generation process for an example tensor with 3 modes and 8 elements. For a given output mode, FLYCOO divides the tensor into multiple partitions with an equal number of output mode indices. We call these partitions super-shards. Each PE processes the super-shards one by one. The number of intermediate values generated while processing each super-shard is proportional to the number of output mode indices in the super-shard. It ensures the intermediate values are combined in the FPGA internal memory to generate the output factor matrices. Despite generating many intermediate values similar to other tensor formats, our method enables combining them into output factor matrix elements before communicating the results to the FPGA external memory.

Due to the sparsity of input tensors, each super-shard contains a different number of nonzero tensor elements. We further divide each super-shard into shards, where the shards have the same number of nonzero tensor elements residing inside. Having the same size shards throughout the execution enables streaming memory access to the input tensor stored in FPGA external memory. It also leads to a static load balancing scheme described in Section 4.7 that fairly distributes the workload among PEs in all the output mode computations. Before partitioning the super-shards into shards, each super-shard is ordered based on Z-Morton order [12] using input mode indices of each nonzero element (see Figure 2). As a result, the nonzero tensor elements with the same input mode indices

reside within the same shard. Z-Morton order recursively partitions multidimensional data into one dimension while preserving the locality of the data [12]. We adopt the Z-Morton order that lays out the elements along a recursive $\square$ -shaped curve. The Z-Morton vector of each nonzero element is computed from the indices of each input mode by interleaving their binary coordinate values. For each nonzero tensor element, the input factor matrices are accessed based on their indices of the input modes. Hence, the proposed ordering improves the data locality in each shard while accessing input factor matrices.

4.3.1 Tensor Format Definition. Following the hypergraph representation described in Section 4.2, we can define the tensor format as follows. Consider a FPGA internal memory with enough space to store N rows of the output factor matrix. For each output mode o we partition the vertex set $\mathcal{V}_o$ which represent the indices of the input tensor in mode o into equal-size vertex sets $\mathcal{V}_{o,0}, \mathcal{V}_{o,1}, \dots, \mathcal{V}_{o, \frac{N}{n_o}-1}$, where $\mathcal{V}_{o,i} = \{i\}$. Here, $|\mathcal{V}_o|$ refers to the size of vertex set $\mathcal{V}_o$. For $i = 0, 1, \dots, (\frac{N}{n_o} - 1)$, each vertex set $\mathcal{V}_{o,i}$ of size n_o is defined as a subset of vertex set $\mathcal{V}_o$. We call each vertex set $\mathcal{V}_{o,i}$ as an interval. Then, we collect all the hyperedges incident on the vertices (i.e., nonzero tensor elements) in $\mathcal{V}_{o,i}$ together as a super-shard, denoted by $SS_{o,i}$. Since real-world sparse tensors have high variance in the distribution of nonzero tensor elements, each super-shard contains a different number of hyperedges. This leads to load imbalance during the spMTTKRP computation. To address this, we further divide each super-shard into equal-sized sets called shards. Each super-shard $SS_{o,i}$ is further divided into $\frac{|SS_{o,i}|}{n_o}$ shards to fit in the FPGA buffers of size n_o . Here, we denote the i -th shard in $SS_{o,i}$ as $\mathcal{S}_{o,i,j}$. The total number of shards for mode o is $\sum_{i=0}^{\frac{N}{n_o}-1} \frac{|SS_{o,i}|}{n_o} \approx \frac{N}{n_o}$ for a tensor with N nonzero tensor elements.

FLYCOO format maps each nonzero element in the tensor to a shard in each mode. A tensor of size N with n modes in the FLYCOO format is a sequence $\mathcal{Q}_0, \dots, \mathcal{Q}_{n-1}$, where each element $\mathcal{Q}_i$ is a tuple $(\mathcal{Q}_{i,0}, \mathcal{Q}_{i,1}, \dots, \mathcal{Q}_{i, \frac{N}{n_i}-1})$. $\mathcal{Q}_{i,j} = (\mathcal{Q}_{i,j,0}, \dots, \mathcal{Q}_{i,j, n_i-1})$ is a shard ID vector where each shard ID corresponds to a mode of the tensor. Here, $\mathcal{Q}_{i,j} = (\mathcal{Q}_{i,j,0}, \dots, \mathcal{Q}_{i,j, n_i-1})$ if and only if $\mathcal{Q}_{i,j,0} \dots \mathcal{Q}_{i,j, n_i-1} = \mathcal{Q}_{i,j}$. This is used to locate the shards where each nonzero tensor element belongs in each mode. $\mathcal{Q}_i = (\mathcal{Q}_{i,0}, \dots, \mathcal{Q}_{i, \frac{N}{n_i}-1})$ is the original indices of the nonzero tensor element in each dimension. $\mathcal{Q}_{i,j}$ is the value of the nonzero tensor elements of the tensor at $\mathcal{Q}_i$. Following the notation used in Section 4.3.1, a single nonzero element in the FLYCOO format requires approximately $n \times \log_2 \frac{N}{n_o} + \sum_{i=0}^{n-1} \log_2 |\mathcal{Q}_i| + \mathcal{Q}_{float}$ bits, where $\mathcal{Q}_{float}$ is the number of bits needed to store the floating-point value of the nonzero tensor element. We encode $|\mathcal{Q}_i| \approx n \times \log_2 \frac{N}{n_o}$, $|\mathcal{Q}_i| = \sum_{j=0}^{n_i-1} \log_2 |\mathcal{Q}_{i,j}|$, and $|\mathcal{Q}_{float}| = \mathcal{Q}_{float}$.

The FLYCOO format generation is a preprocessing task. Even though the remapping process is executed in real-time, the FPGA external memory spaces (i.e., memory layout) for each shard can be statically decided during the preprocessing time (see Section 4.7).

While FLYCOO and HiCOO [12] adopt similar tensor ordering strategies, such as Z-Morton ordering during the format generation, they have significant differences, including: (1) Intermediate value communication: According to the FLYCOO tensor format introduced in Section 4.3.1, each super-shard contains all the nonzero

tensor elements of an output factor matrix interval. Since each interval can be fit in the FPGA internal memory, each interval can be computed without communicating intermediate values generated during the computation to the FPGA external memory. However, HiCOO requires intermediate values to be communicated to the FPGA external memory. Since spMTTKRP is a memory-bound operation and FLYCOO avoids communicating intermediate values to the FPGA external memory, FLYCOO significantly reduces the overall execution time. (2) Tensor partitioning scheme and nonzero element distribution: HiCOO partitions contain an equal number of tensor indices for each mode, while the FLYCOO shards have the same number of nonzero tensor elements. The approach used by FLYCOO leads to load-balanced computation in a multi-PE accelerator, as discussed in Section 4.6. HiCOO partitions vary in size because nonzero elements are not distributed evenly across tensor indices. Hence, HiCOO leads to non-uniform partitioning of the tensor. (3) Applying Z-Morton ordering: FLYCOO applies Z-Morton ordering for each super-shard independent of each other along the indices of the input mode. HiCOO applies Z-Morton ordering once considering all the modes. FLYCOO enables better data reuse of input factor matrices while computing spMTTKRP with the use of the accelerator hardware (see Section 4.5).

4.4 Parallel Algorithm

We perform mode-by-mode super-shard computation. Since each super-shard can be executed independently by construction (see Section 4.3), the order of super-shard computation does not affect the outputs. Further, multiple super-shards can be executed in parallel.

Algorithm 2 shows the parallel algorithm for a super-shard for a given mode. The functions Load and Store correspond to loading and storing data from the FPGA external memory. The parallel algorithm consists of (1) spMTTKRP Computation and (2) Data Remapping for the next mode. These 2 stages are executed concurrently.

In Algorithm 2, all the factor matrices are accessed in row-major order. Hence, the factor matrices are stored in the FPGA external memory in row-major order.

At the beginning of Algorithm 2, a super-shard gets assigned to a PE for execution. In the spMTTKRP Computation, the shards that belongs to the same super-shard are loaded into an internal buffer one by one (Algorithm 2: line 7). After a shard is loaded, the tensor elements inside the shard are assigned for execution. First, the coordinates of the modes are extracted from the tensor element in FLYCOO format (Algorithm 2: line 10). Then the corresponding rows (based on input mode indices of the tensor element) of the input factor matrices are loaded into the PE from the FPGA external memory. After, the element-wise operations between the tensor element and the rows of the input factor matrices are performed inside the PEs (Algorithm 2: line 15-20). The PE maintains an internal memory buffer to store the intermediate values of the computation. Keeping this internal buffer size proportional to interval size, $|I|$ (see Section 4.3.1) ensures that all the intermediate values can be updated only using the internal memory of FPGA. After a super-shard is completely processed, the generated output interval is stored in the FPGA external memory (Algorithm 2: line 26).

After computing spMTTKRP for mode $\square$ the accelerator should compute spMTTKRP for the subsequent mode (i.e., mode $(\square + 1)$

mod $\square$). Nonzero tensor elements should be ordered according to the output mode to support our proposed parallel algorithm. Therefore, the tensor should be remapped according to the shard IDs of mode $(\square + 1)$ mod $\square$ to compute the spMTTKRP of the upcoming mode. Hence, as the current mode runs, the tensor is remapped in parallel according to the shard IDs of the upcoming mode, allowing sequential execution of all the modes.

Algorithm 2: Computation of a super-shard in mode $\square$

```

1 Input: Input factor matrices-set  $Y = \{\square_0, \square_1, \dots, \square_{\square-1}\} \setminus \{\square\}$ , 2
super-shard in mode  $\square$   $\square_{\square} = \{\text{shard}_{\square\square_0}, \dots, \text{shard}_{\square\square_{\square-1}}\}$ 
Result: Store interval  $\square$  from output factor matrix  $\square_{\square}$ ,
remap and store each tensor element  $\square_{\square}$ 
3 Compute & On-the-fly Remap( $\square$   $\square_{\square}$ ):
4 Initialize  $\square$  as a zero matrix
5  $\square \leftarrow (\square + 1) \bmod \square$  // upcoming mode
6 for each shard  $\square_{\square}$  in  $\square_{\square}$  do
7   Buffer $\square_{\square} \leftarrow \text{Load}(\text{shard}_{\square\square_{\square}})$ 
8   for each element  $\square_{\square} = \langle \square_0, \square_1, \dots, \square_{\square-1} \rangle$  in Buffer $\square_{\square}$ 
9      $\square_{\square} \leftarrow \square_{\square}$ 
10     $\square_{\square} = (\square_0, \dots, \square_{\square-1})$ 
11     $\square_{\square} \leftarrow \square_{\square}$ 
12     $\square_{\square} \leftarrow \square_{\square}$ 
13    //  $\ell$  is a vector of size  $R$ . Each element initialized to 1
14     $\ell \leftarrow 1$ 
15    for each input mode  $\square_{\square} \in \{0, \dots, \square-1\} \setminus \{\square\}$  do
16       $\square_{\square} \leftarrow \text{Load}(\text{row } \square_{\square} \text{ from } \square^{\text{th}} \text{ factor matrix})$ 
17    for each rank  $\square$  in  $\square$  parallel do
18       $\ell(\square) \leftarrow \ell(\square) \times \square_{\square}(\square)$ 
19    for each rank  $\square$  in  $\square$  parallel do
20       $\square(\square, \square) \leftarrow \square(\square, \square) + \square_{\square}(\square) \times \ell(\square)$ 
21    // Inside Remap_Cache_Buffer
22    shard_collector( $\square_{\square}$ )  $\leftarrow \text{Update}(\text{shard_collector}(\square_{\square}), \square_{\square})$ 
23    if Remap_Cache_Buffer full then
24      for each collected shard_collector( $\square_{\square}$ ) in
25        Remap_Buffer do
26          Store(append shard_collector( $\square_{\square}$ ) to  $\square_{\square} \square_{\square}$ )
26 Bufferoutput  $\leftarrow \square$ 
27 Store(Bufferoutput to  $\square_{\square}$  factor matrix)
28 for each collected shard_collector( $\square_{\square}$ ) in
29   Remap_Cache_Buffer do
30     Store(append shard_collector( $\square_{\square}$ ) to  $\square_{\square} \square_{\square}$ )
```

The on-the-fly remapping for the upcoming mode is performed as follows: First, the shard ID of the subsequent mode to be executed is extracted from each tensor element. Then, the tensor elements with the same shard ID of the subsequent mode are collected together inside the Remap_Cache_Buffer using the update function (Algorithm 2: line 22). Remap_Cache_Buffers are not large enough to store all shards; therefore, as the Remap_Cache_Buffer becomes full, we store partially aggregated shards (i.e., shard collectors) in

the corresponding shard locations in the FPGA external memory (Algorithm 2: lines 23-25 & 28-29).

The total computations per mode in the proposed Algorithm 2 is $\prod_{i=1}^M |\Omega_i| \times \prod_{i=1}^M |\Omega_i|$. Here, the factor $\prod_{i=1}^M |\Omega_i|$ comes from $\prod_{i=1}^M |\Omega_i| - 1$ multiplications and the addition we perform per nonzero tensor element. The accelerator loads all the nonzero tensor elements and stores the factor matrix of the output mode. Meanwhile, it also stores the remapped tensor optimized for the next mode. Hence it requires a total of $2 \times |\Omega_1| + \prod_{i=1}^M |\Omega_i| \times \prod_{i=1}^M |\Omega_i|$ memory transfers. Here, $\prod_{i=1}^M |\Omega_i|$ represents the length of the output dimension ($\prod_{i=1}^M |\Omega_i| = \{|\Omega_1|, |\Omega_2|, \dots, |\Omega_M|\}$). For factor matrices of rank $\prod_{i=1}^M |\Omega_i|$ with no data reuse, the total factor matrix elements transferred per mode is $(\prod_{i=1}^M |\Omega_i| - 1) \times |\Omega_1| \times \prod_{i=1}^M |\Omega_i|$. Hence, the total amount of data transferred from the FPGA external memory per mode is $2 \times |\Omega_1| + (\prod_{i=1}^M |\Omega_i| - 1) \times |\Omega_1| \times \prod_{i=1}^M |\Omega_i|$.

4.5 FPGA Design

Algorithm 3 is designed for multi-SLR [25] datacenter FPGAs. These FPGAs contain a large number of DSPs and multiple DRAM memories to provide high compute power and memory bandwidth.

Algorithm 3: FLYPAR: FLYCOO-based Parallel algorithm

```

1 Input: Input factor matrices-set  $Y = \{\hat{\Omega}_0, \hat{\Omega}_1, \dots, \hat{\Omega}_{M-1}\}$ , 2
super-shards of mode 0,  $\{H_0\} = \{\hat{\Omega}_0 \hat{\Omega}_1 : \hat{\Omega}_M\}$ 
3 Output: Updated factor matrices-set  $\hat{\Omega} = \{\hat{\Omega}_0, \hat{\Omega}_1, \dots, \hat{\Omega}_{M-1}\}$  4
for each mode  $\square = 0, \dots, M - 1$  do
5   while  $H_\square \neq \emptyset$  do
6     for each ACCEL $_\square$  parallel do
7       if ACCEL $_\square$  is idle then
8         Fetch super-shard  $\square \hat{\Omega}_\square$  in DRAM $_\square$  to
          ACCEL $_\square$ 
9          $H_\square \leftarrow H_\square \setminus \{\square \hat{\Omega}_\square\}$ 
10         $(H_{(\square+1)} \bmod \square, \hat{\Omega}_\square) \leftarrow \text{Compute \&}$ 
          On-the-fly Remap( $\square \hat{\Omega}_\square$ ) // Algorithm 2
11   Wait(All ACCELS are idle)
```

The proposed accelerator contains multiple custom hardware units called ACCELS. For a FPGA with $\prod_{i=1}^M$ DRAMs, we employ $\prod_{i=1}^M$ ACCELS where each ACCEL is directly connected to a DRAM via memory controller as shown in Figure 3 (a). ACCELS communicate with each other using ACCEL routers in a ring interconnection. Inter-SLR switches in the FPGA [24] are used to connect neighboring ACCELS in different Super Logic Regions (SLRs).

Our design has the following features: (1) custom multi-pipeline PEs to support concurrent element-wise tensor multiply and add operations, (2) internal buffers in each pipeline to ensure all the intermediate values remain in the FPGA internal memory while processing a super-shard, (3) use Direct Memory Accesses (DMAs) to load shards and store output factor matrix rows as bulk transfers while optimally using the FPGA external memory bandwidth, (4) multi-cache subsystem to exploit locality while loading input factor matrices, and (5) tensor remapping module to support on-the-fly remapping.

Each ACCEL is assigned a super-shard during the execution time in which the super-shard resides in its directly connected DRAM as described in FLYPAR: Parallel spMTTKRP Accelerator Algorithm based on FLYCOO (Algorithm 3). The $\prod_{i=1}^M$ ACCELS process $\prod_{i=1}^M$ super-shards in parallel following the parallel algorithm in Section 4.4.

All the ACCELS synchronize at the end of the computations in each output mode. Our workload distribution (see Section 4.6) ensures PE idle time is minimized. As shown in Figure 3 (b), each ACCEL has a Processing Engine (PE) and a memory controller. Each PE executes the element-wise computation on the tensor elements. The memory controller manages the data flow between the PEs and their external memory.

4.5.1 Processing Engine (PE). PE is the compute unit inside an ACCEL. It can concurrently process $\prod_{i=1}^M$ input tensor elements in each clock cycle using $\prod_{i=1}^M (\prod_{i=1}^M \geq 1)$ processing pipelines as shown in Figure 3 (c). A PE processes a shard with the following 3 steps: (1) a pipeline reads an element of the shard from the tensor buffer, (2) pipeline extracts the coordinates of the tensor element and requests corresponding input factor matrices from the caches in the memory controller, (3) When all the requested input factor matrix rows become available inside the pipeline, it performs element-wise compute as indicated in Algorithm 2 and stores the corresponding partial value in its matrix buffer, (4) after computing all the elements in a super-shard, PE computes the final output factor matrix rows by adding the partial sums in the partial matrix buffers in all the pipelines with the same output mode indices using PE's adder tree. For a PE with $\prod_{i=1}^M$ pipelines, the adder tree has $\log_2(\prod_{i=1}^M)$ stages with $\prod_{i=1}^M (\prod_{i=1}^M + 1)/2$ total adders, (5) outputs the output factor matrix rows to the output matrix buffer inside the memory controller. This process generates a interval of the output factor matrix that corresponds to the input super-shard. At the end of the process, it is transferred to the FPGA external memory as a bulk transfer using DMA inside the memory controller.

4.5.2 Memory Controller. Figure 3 (d) shows the details of the memory controller (MC). It consists of a cache subsystem, Direct Memory Accesses (DMA) module, a Tensor Remapper Module (TRM), an ACCEL router, and an external memory interface. The cache subsystem exploits the spatial and temporal locality of the input factor matrix accesses enabled by the proposed FLYCOO format. Each cache is a set-associative cache with the Least Recently Used (LRU) cache-line replacement policy. Multiple input factor matrices can share a cache. The number of caches inside a MC, the number of input factor matrices that share a single cache, and cache size are chosen based on available FPGA resources. DMAs can efficiently access sequential data as bulk transfers. DMAs load shards to the tensor buffer and store the output factor matrix from the output matrix buffer as bulk transfers. By keeping the size of a shard as same as the tensor buffer size and the interval size equal to the output matrix buffer size, we avoid data overflows inside the DMAs. The DMA uses double buffers to overlap the communication time of the shards from the FPGA external memory with the compute time of the PEs. On-the-fly remapping is supported by the tensor remapping module (TRM) in the MC. When the PE requests a tensor element from the tensor buffer, a copy of the same tensor element is passed to the TRM using a shared bus, as shown in Figure 3 (d). TRM extracts the shard ID of the upcoming output mode of the spMTTKRP in the CPD decomposition iteration (see Algorithm 1). The TRM collects the tensor elements with the same shard ID until it fills a complete row of the TRM buffer and transfers it to the FPGA external memory.

ACCEL router maintains the super-shard scheduling while controlling the data flow between PE, DRAM, and neighboring ACCELS.

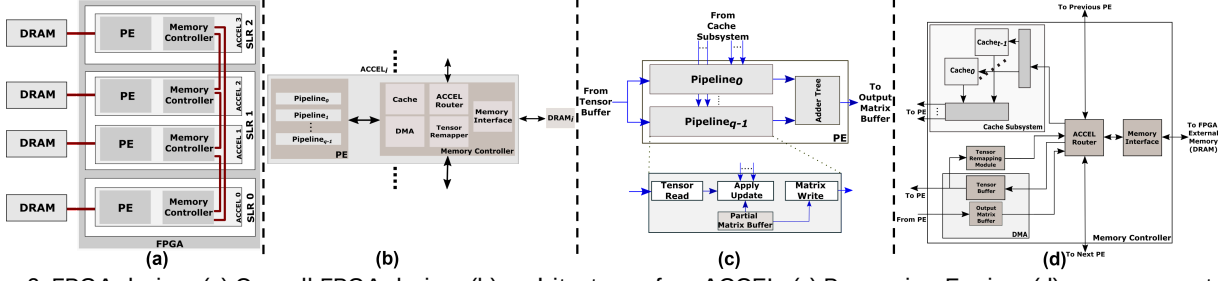


Figure 3: FPGA design: (a) Overall FPGA design, (b) architecture of an ACCEL, (c) Processing Engine, (d) memory controller

ACCEL router resides inside the memory controller of each ACCEL. Each ACCEL Router is connected to neighboring ACCEL routers as a ring (see Figure 3 (b)). Each ACCEL router communicates data as data packets between its neighboring ACCEL routers in the ring. It uses large data buses between ACCELS while maximally using inter SLR [24] routing resources to transfer multiple data packets simultaneously. The accelerator maintains a virtual address space to keep track of the memory addresses of each tensor shards of all the modes and the intervals of all the factor matrices. The accelerator also maintains memory pointers to identify the location to be read or stored in a shard to support the proposed tensor remapping. The virtual addresses are generated during the tensor generation process and stored in the FPGA external memory as initial meta-data.

4.5.3 ACCEL Router. Each ACCEL Router performs Algorithm 4 that, (1) loads shards from directly connected external DRAM, (2) stores the output factor matrix rows depending on the PE requests, (3) loads the input factor matrix rows depending on the PE requests from different ACCELS, and (4) stores the remapped tensor elements in the corresponding shard location. The destination location of the remapped tensor element can be in any of the external DRAMs connected to the FPGA.

The data router loads each shard from its directly connected DRAM and forwards it to Tensor Buffer inside DMA (Algorithm 4 line 2-8). Additionally, memory access requests are received from Cache Subsystem (Algorithm 4 line 9-10) and Tensor Remapping Module (Algorithm 4 line 11-12) to the ACCEL router. In addition, data packets are forwarded from the previously connected ACCEL router (Algorithm 4 line 13-14). For the input factor matrix row read requests, if the factor matrix is in the directly connected DRAM, it is loaded from the DRAM and forwarded either to the previous ACCEL router in the ring or to the cache subsystem. Otherwise, the request is forwarded to the next ACCEL router in the ring (Algorithm 4 line 15-26). When a tensor remapping request is reached to the ACCEL, if the corresponding shard of the nonzero tensor element is located inside the directly connected DRAM, it is stored in the DRAM. Otherwise, remap request with the tensor element is forwarded to the next ACCEL router in the ring (Algorithm 4 line 27-32).

We use memory interface IPs [7, 26] to maintain the low-level signals (e.g., refresh, and pre-charge) between DRAMs and FPGA.

4.6 Load balancing

Since the accelerator performs spMTTKRP computation for a single output mode at a time, we consider load balancing the total computation mode by mode. In a given mode, the total number of computations corresponding to a super-shard is proportional to the number of nonzero tensor elements in the super-shard. Since each super-shard is further partitioned into shards with same number

Algorithm 4: Data Routing Algorithm for ACCEL_i

```

1 Routing (ACCELi,  $\square$ ): // For output mode  $\square$ 
2 if PE is Idle then
3   SS = Get Active super-shard ID
4   shard = Next shard (SS)
5   if all shards of SS processed then
6     interval data = Get output data from PE
7     Store DRAMi(Get Interval (SS), interval data)
8     SS = Get Next super-shard ID
9 if Factor Matrix Row  $\square$  is requested by a Cache Miss then
10  {type, src_accel, w, info} ← {fm_read, i, n, k}
11 if else remapped tensor y available from Tensor Remapper
    then
12  {type, src_accel, w, info} ← {remap_store, i, n, y}
13 if else data_packet from from ACCELi-1 then
14  {type, src_accel, w, info} ← data_packet
15 if type = fm_read then
16  I = Extract Interval (info, w)
17  if I is in DRAMi then
18    value = Read DRAMi(w, info)
19    if src_ACCEL = ACCELi then
20      Forward Cache(value, w, info)
21    else Forward ACCELi-1 ({fm_write, src_accel, n, w,
22      info})
23  else Forward ACCELi+1 ({fm_read, src_accel, n, w,
24    info})
25 if type = remap_store then
26   $\square$  = Get Next Mode shard ID (y)
27  if  $\square$  is in DRAMi then
28    Store DRAMi(info, S)
29  else Forward ACCELi+1 ({remap, src_accel, n, w, info})

```

of tensor elements, the total number of computations also proportional to the number of shards in a super-shard. Each super-shard contains a different number of shards depending on the sparsity of the tensor. We use a greedy approach to distribute the super-shards and perform the spMTTKRP computation. For each mode, the proposed method evenly distributes the total workload among the ACCELS. Suppose a FLYCOO tensor of size $\square$ is partitioned into super-shards $\{\square \square : \square \square\}$ for a mode $\square$. Let $\square_{i,j}$ be the number of

shards in the super-shard $\square_{\ell}$. We reorder the indices of super-shards so that $\square_{\ell} \geq \square_{\ell'}$ if $\ell < \ell'$, so that they are sorted in descending order of number of shards. Each super-shard $\square_{\ell}$ is iteratively assigned to the ACCEL that is currently the least heavily loaded (i.e., with the least number of shards). We perform the above operation for all the modes of the tensor.

For a given output mode, let $\square_{\max}$ be the largest number of shards assigned to a single ACCEL and $\square_{\max}$ be the value of $\square_{\max}$ in an optimal shard distribution among the ACCELS. Then our proposed greedy approach above guarantees that $\square_{\max} \leq 4/3 \cdot \square_{\max}$ [4]. The proof of this approximation guarantee is shown by induction in the following.

Theorem 4.1. For a given mode ℓ based on our greedy approach, the number of shards assigned to each ACCEL is at most $4/3 \cdot \square_{\max}$.

Proof. Assume the theorem is true for at most $\ell - 1$ super-shards. Suppose there are ℓ super-shards, and the ℓ -th super-shard is the last one that is assigned to the ACCEL ℓ containing the most shards. Consider two cases: (i) If $\ell < \ell$ then since the ℓ -th super-shard is not assigned to ℓ removing it doesn't affect the result. Therefore, the theorem follows immediately by the induction hypothesis on the first $\ell - 1$ super-shards. (ii) If $\ell = \ell$ then $\square_{\max} \leq \square_{\max} + \square_{\ell}$. If $\square_{\ell} \leq \square_{\max}/3$ then the theorem holds immediately. Suppose $\square_{\ell} > \square_{\max}/3$, then $\square_{\max} < 3\square_{\ell}$, so each ACCEL has been assigned either one or two super-shards, which is in fact an optimal assignment: exchanging any two super-shards in different ACCELS will increase the largest number of shards in an ACCEL. So $\square_{\max} = \square_{\max}$ in this case. $\square$

4.7 Tensor Remapping

As described in Section 4.4, tensor elements are remapped according to the shard IDs of the upcoming mode. Theorem 4.2 shows that it only requires $2 \times |\square|$ FPGA external memory to perform on-the-fly tensor remapping. Note that the mode-specific formats require the number of tensor copies proportional to the number of modes of an input tensor. Theorem 4.3 shows the tensor layout generated through the proposed tensor remapping technique is load balanced across all ACCELS, and the tensor data is locally available to each ACCEL in its directly connected external DRAM memory during each output mode computing time.

The accelerator requires the external memory address pointers for each shard to identify the destination memory address of nonzero elements while tensor remapping. The memory address pointers for the initial position of each shard are computed during the tensor format generation time. These memory address pointers are updated as the shards are getting filled during the runtime. During the remapping, the accelerator only requires to keep track of the address pointers corresponding to a single mode. Also, following the routing algorithm discussed in Section 4.5.2, each ACCEL only requires to keep track of the shards mapped into its directly connected DRAM.

Theorem 4.2. The total FPGA external memory required to store an input tensor $\square$ in the FLYCOO format for our proposed algorithm FLYPAR is $2 \times |\square|$, where $|\square|$ is the size of the tensor, independent of the number of modes.

Proof. Consider a FPGA external memory space of size $\square = 2 \times |\square|$, divided to $\square[0]$ and $\square[1]$. If the current mode (ℓ) being executed is even (i.e., $\ell = 2, 4, 6, \dots$), then it uses $\square[0]$ to keep the tensor copy ordered according to mode ℓ shard IDs. $\square[1]$ is used

to keep the tensor copy ordered according to shard IDs of mode $(\ell + 1) \bmod \square$ during the on-the-fly remapping. Similarly, if the current output mode being executed is odd (i.e., $\ell = 1, 3, 5, \dots$), it uses space $\square[1]$ to keep the tensor copy ordered according to mode ℓ shard IDs. $\square[0]$ is used to keep the tensor copy ordered according to mode $(\ell + 1) \bmod \square$. As a result, we only need $|\square| = 2 \times |\square|$ space to store the tensor to support the algorithm FLYPAR. $\square$

Theorem 4.3. Algorithm FLYPAR performs spMTTKRP for all the modes of an input tensor such that the computation is load balanced across all ACCELS and the tensor data is available for each ACCEL in its directly connected external DRAM memory for all the output mode computations.

Proof. In the FLYCOO format, each tensor element contains its shard IDs for all the modes (see Section 4.3). For an input tensor with $\square$ modes, during the computation of mode $\ell - 1$, the on-the-fly-remapping remaps each nonzero tensor element according to the shard IDs of mode ℓ ($0 \leq \ell < \square$). The data routing algorithm (Algorithm 4.5.3) allows this by remapping each tensor element to the DRAM directly connected to the PE. Hence, at the beginning of the spMTTKRP computation for an output mode ℓ (Algorithm 2), all the nonzero elements in the tensor have been assigned to the shards of mode ℓ ($\square_{\ell}$). Therefore, the tensor data is available for each ACCEL in its directly connected external DRAM memory for all the output mode computations.

The load balancing approach in Section 4.6 proves that the computation is load balanced if each tensor element is in its assigned shard for an output mode ℓ . According to the on-the-fly remapping, at the beginning of the computation, all nonzero elements in the tensor are in the shards of mode ℓ . Since we perform the overall computation mode-by-mode, the overall computation is load balanced for all the modes. $\square$

Theorem 4.4 proves that the remapping cost introduced in FLYPAR is insignificant compared to total FPGA external memory accesses while accessing input factor matrices in each mode.

Theorem 4.4. In any mode, without data reuse, the ratio of the total amount of data transferred in remapping to that for accessing the input factor matrices is $O\left(\frac{1}{(\square-1)\square}\right)$, where $\square$ is the number of modes and $\square$ is the rank of the factor matrices.

Proof. For a given mode, since every nonzero element is remapped for spMTTKRP in the upcoming mode, the total amount of data transferred per mode for remapping is equal to the size of the total number of nonzero elements ($|\square|$). With no data reuse, the total amount of data transferred per mode for the input factor matrices is $(\square - 1) \times |\square| \times \square$ (see Section 4.4). Therefore, the ratio of the amount of data transferred for the remapping to that for accessing input factor matrices is $O\left(\frac{1}{(\square-1)\square}\right)$. $\square$

4.8 Super-shard Scheduling

After distributing super-shards among ACCELS in Section 4.6, we further optimize the memory access cost between distinct super-shards assigned to each ACCEL in every mode by statically scheduling these super-shards. The objective of the super-shard scheduling is to reuse cached rows of input factor matrices across multiple super-shards as much as possible. The first step to achieving this is to create a weighted directed complete graph $\Gamma = (\Gamma_{\square}, \Gamma_{\square}, \Gamma_{\square})$ whose vertex set $\Gamma_{\square}$ consists of the super-shards, which are assigned to the target ACCEL. Here, the weight $\Gamma_{\square} \rightarrow \Gamma_{\square}$ of each edge

Γ represents the number of rows of factor matrices that any two super-shards have in common. After constructing the complete graph Γ , our goal is to find a maximum weighted Hamiltonian path π in Γ . The order of vertices in the Hamiltonian path π represents the scheduling of the super-shards assigned to each ACCEL. It leads to efficient cache utilization by maximizing the reusability of the common input factor matrices among super-shards after loading them into the caches. Since the maximum weighted Hamiltonian path problem is NP-complete [1], we use a greedy heuristic by adding one edge at a time, as summarized in Algorithm 5. The schedule computed as a pre-processing step for each input tensor.

Algorithm 5: super-shard Scheduling

```

1 Input: A weighted directed complete graph  $\Gamma = (\Gamma_v, \Gamma_e, \Gamma_w)$ 
   consisting of the super-shards
2 Output: A Hamiltonian path  $\pi$  representing the scheduling
   of the super-shards assigned to each ACCEL
3  $\pi \leftarrow \emptyset$ 
4 Sort  $\Gamma_v$  in descending order of  $\Gamma_w$ 
5 for each edge  $e = (u, v) \in \Gamma_e$  do
6   if  $\deg^+(u) = 0$  &  $\deg^-(v) = 0$  & no cycle in  $\pi \cup \{e\}$  then
7      $\pi \leftarrow \pi \cup \{e\}$ 
8 return  $\pi$ 

```

5 EVALUATION

5.1 Experimental Setup

5.1.1 FPGA Platform. We implement our hardware design on Xilinx Alveo U250 Data Center Accelerator Card [24] using Verilog HDL. This Alveo Card consists of 4 super logic regions (SLRs) [25]. The SLRs connected to DRAM modules through memory interface IPs [26]. Simulation, synthesis, and place and route are performed using Xilinx Vivado Design Suite 2020.2 [27].

5.1.2 Datasets. We use the sparse tensors from real-world applications shown in Table 2. All the tensors are from the Formidable Repository of Open Sparse Tensors and Tools (FROSTT) dataset [19]. The selected datasets have tensors with different shapes, sizes, and sparsities.

5.1.3 FLYCOO Format Generation Time. Since we implement the preprocessing step using Python libraries, our preprocessing algorithm is substantially slower than C/C++ implementations. There are no meaningful comparisons of the execution time between the pre-processing and processing steps with C/C++-based implementations in related works. Hence in this work, we do not show the implications of preprocessing costs which are executed offline.

5.1.4 Performance Model Simulator. Performance Model Simulator (PMS) is used for optimizing the accelerator configuration. PMS can estimate the number of clock cycles spent on computing spMTTKRP for a given input tensor. PMS models each hardware module in the accelerator design at the cycle level. The PMS can further estimate the total FPGA internal memory requirement and the DSP usage for a given accelerator configuration.

Our objective is to use PMS to (1) identify the best set of hardware parameters of the accelerator to obtain the least average execution time for a given collection of datasets (see Section 5.1.5), and (2) evaluate the impact of various hardware modules in the proposed design and algorithmic optimizations due to the versatility of PMS.

Obtaining module-wise performance results is effortless with PMS compared with the actual hardware implementation.

The system designer has to provide the following inputs to the PMS: (1) resources of target FPGA (i.e., total DSPs, BRAMs, and URAMs and data width of memory interface), (2) design parameters (i.e., number of pipelines per ACCEL, DMA buffer sizes, number of caches, number of cache lines, associativity of a cache, and number of factor matrices shared by a cache), (3) algorithm parameters (i.e., rank of the factor matrices, size of a shard, and size of an interval), and (4) input tensor parameters (i.e., number of modes, and length (dimension size) of each mode).

Table 2: Characteristics of the sparse tensors

Tensor	Shape	#NNZs	Density
NELL-1	$2.9 \times 2.1 \times 25.5$	143.6	9.1×10^{-13}
NELL-2	$12.1 \times 9.2 \times 28.8$	76.9	2.4×10^{-05}
PATENTS	$46 \times 239.2 \times 239.2$	3.6	1.4×10^{-03}
LBNL	$1.6 \times 4.2 \times 1.6 \times 4.2 \times 868.1$	1.7	4.2×10^{-14}
DELICIOUS	$532.9 \times 17.3 \times 2.5 \times 1.4$	140.1	4.3×10^{-15}

The overall execution time (in cycles) reported by PMS is validated using FPGA run-time results. Figure 5 compares the total execution time of each tensor on FPGA and the PMS. The behavior of DRAM technology and packet routing on ACCEL routers can not be accurately simulated using a software model. Hence PMS shows below 10% error compared with the actual hardware.

We performed an extensive parameter search to select the configuration of the accelerator. The goal is to identify the hardware parameters with the least average execution time for all the datasets in Table 2. spMTTKRP is executed per single iteration for each dataset.

5.1.5 Optimizing Accelerator Configuration. The execution time depends on the total number of clock cycles and FPGA operating frequency (MHz). We use PMS to estimate the number of clock cycles for computing spMTTKRP. We determine the average number of FPGA cycles (ns) by taking the PMS results for all the targeted tensors. PMS can also estimate the total FPGA internal memory requirement and the DSP usage for a given accelerator configuration.

We can choose an accelerator configuration that provides the least ns and is subject to the resource constraints of the FPGA using PMS. First, we focus on identifying the shard and interval sizes that are suitable for the target FPGA. These sizes directly correlate with the accelerator buffer sizes. We use PMS to identify the shard and interval sizes that fit the target FPGA and take the minimum number of FPGA clock cycles. Using PMS, we select 10 suitable configurations with a minimum of ns and fit the target FPGA resources. Then we use Place and Route (P&R) using Xilinx Vivado to obtain the ns of each selected accelerator configuration. Finally, we pick the configuration with the least ns configuration. Table 3 shows the selected configuration for our targeted dataset and FPGA. We observe that using more FPGA resources leads to lower FPGA operating frequency which reduces the overall performance of the design.

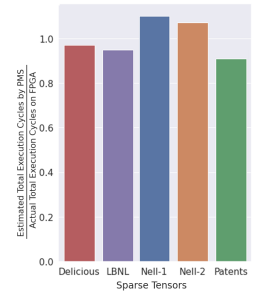


Figure 5: PMS accuracy

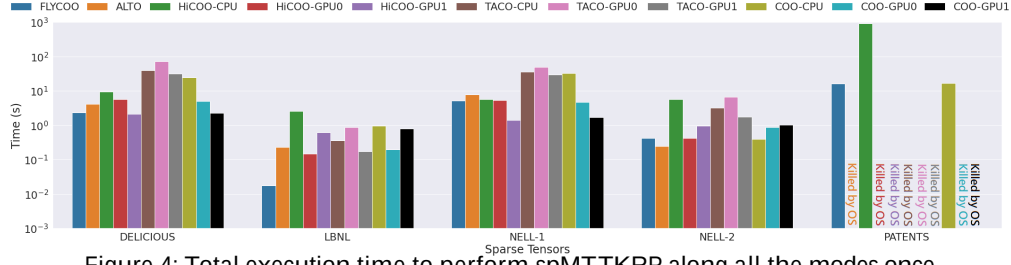


Figure 4: Total execution time to perform spMTTKRP along all the modes once

Table 3: Module Parameter Configuration

Module	Parameter	Configuration
PE	No. of pipelines	16
	Partial Matrix Buffer size	16 KB
Cache subsystem	No. of caches	2
	Set-associativity	4
	No. of cachelines	4096
	cache line width	64 B
DMA	Tensor Buffer	64 KB
	Output Matrix Buffer	16 KB
	Tensor Remapping Buffer	64 KB

Table 4: Resource utilization of the selected design

LUT	FF	BRAM	URAM	DSP
37.23%	24.52%	20.84%	20.52%	25%

5.1.6 Baselines. The baseline experiments are conducted on an Intel Xeon Gold 5120 CPU, an NVIDIA RTX 3090 GPU, and an NVIDIA RTX A6000 GPU. Platform specifications are summarized in Table 5.

Table 5: Specifications of the platforms

Platform	CPU Intel Xeon Gold 5120	GPU0 NVIDIA RTX 3090	GPU1 NVIDIA RTX A6000	FPGA Xilinx Alveo U250
Technology	Intel 14 nm	Samsung 8 nm	Samsung 8 nm	TSMC 16 nm
Frequency	2.20 GHz	1695 MHz	1410 MHz	230 MHz
Peak Device Performance	14.9 GFLOPS	35.6 TFLOPS	38.71 TFLOPS	0.6 TFLOPS
Device Internal Memory	19.25 MB L3 Cache	6 MB L2 Cache	6 MB L2 Cache	54 MB
External Memory Bandwidth	107.3 GB/s	936.2 GB/s	768 GB/s	77 GB/s

We evaluate our work against mode-specific CSF format[9], mode-specific COO format[11], mode-agnostic HiCOO format[12] on both CPU and GPU platforms. We also evaluate the ALTO format [5] on the CPU. For the COO format, we use the library ParTI[11]. HiCOO is the required input parameter block size for its partitioning scheme, similar to the shard size in our proposed work. Similar to literature [5, 12], we use a block size of 128 to evaluate the performance of HiCOO. We ran OpenMP-enabled HiCOO CPU implementation using all the CPU threads. We evaluate the performance of mode-specific CSF formats using TACO[9]. We optimize the TACO code to our target CPU and GPU platforms using their command-line tool[9].

5.2 Overall Performance

Following the baselines[5, 12], we set the tensor rank (κ) as 16. The interval size (ℓ) is set to the same size as the output matrix buffer size, which is determined during optimizing the accelerator. We also keep the shard size equal to the tensor buffer size. Our experiments are conducted on the actual FPGA hardware. The tensor format generation is a one-time preprocessing step. As in the baselines[9, 12], we do not include the tensor generation time in the overall execution time.

Figure 4 displays the total execution time while computing spMTTKRP along all the modes. Our work is the only implementation that delivers consistent performance across all the modes of the datasets. Our work achieves an average of 8.8 $\times$ and 3.8 $\times$ speedup compared with the baseline CPU and GPU implementations. Table 6 summarizes the average speedup of each baseline and our work compared to COO-CPU[11].

We observe that LBNL generates a large number of intermediate values. Therefore, baselines that hold intermediate data in the external memory have to load and store intermediate data several times while computing a single row of the output factor matrix. Since our approach avoids partial outputs being stored in the external memory, it reduces the communication time between the FPGA and its external memory. LBNL also has limited data reuse while accessing the rows of input factor matrices. Our cache system exploits the data reuse of the input factor matrices of LBNL better than the other datasets. This leads to a significant execution time reduction for the LBNL dataset compared with the rest of the benchmarks. When executing PATENT, the OS kills ALTO, TACO CPU, and all GPU implementations because the data generated during compute time does not fit in the device external memory. The selected devices did not have sufficient external memory to hold the intermediate values generated during the execution time of PATENT for our baselines.

5.3 Impact of Tensor Remapping

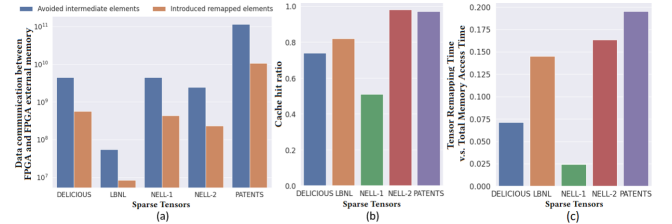


Figure 6: (a) Number of intermediate elements avoided being transferred to FPGA external memory by tensor remapping vs. total tensor elements remapped, (b) Cache-hit ratio while accessing rows of input factor matrices, (c) Tensor remapping time over total memory access time

The proposed parallel algorithm, FLYPAR (Algorithm 3), avoids intermediate values being transferred to the FPGA external memory at the additional cost of tensor remapping. Note that it introduces additional data transfers during on-the-fly tensor remapping. Figure 6 (a) shows a comparison between the amount of remapped tensor elements transferred to the FPGA external memory and the amount of data transfers (to the FPGA external memory) avoided by combining the intermediate values inside the FPGA internal memory for every super-shard. Our results indicate that FLYCOO

Table 6: Comparison of baselines on their targeted platforms and this work

Tensor Formats	FLYCOO FPGA	ALTO CPU	HiCOO CPU	HiCOO GPU0	HiCOO GPU1	TACO CPU	TACO GPU0	TACO GPU1	COO CPU	COO GPU0	COO GPU1
Speedup (over COO-CPU)	18.0	3.9	2.2	4.4	9.2	1.0	0.5	1.90	1.0	4.2	7.8
Device Peak Bandwidth (GB/s)	77	107.3	107.3	936.2	768	107.3	936.2	768	107.3	936.2	768
Device Peak Performance (TFLOPS)	0.60	0.0149	0.0149	35.6	38.71	0.0149	35.6	38.71	0.0149	35.6	38.71

with FLYPAR reduces the FPGA external memory traffic by 9.2× on average. Figure 6 (b) shows the cache hit rate while accessing input factor matrices. It confirms that FLYCOO significantly increases data reuse. Figure 6 (c) compares the total tensor remapping time with the time spent on the rest of the memory accesses. The results show that even though the input factor matrices are cached in the FPGA, the tensor remapping cost is still insignificant compared with the rest of the memory access cost.

5.4 Impact of Memory Controller

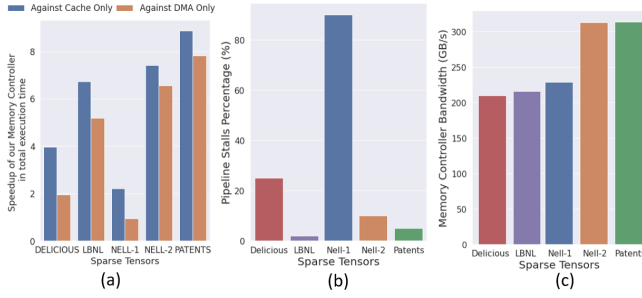


Figure 7: (a) Improvement of proposed MC in execution time, (b) Number of pipeline stalls wrt. total number of pipeline requests, (c) Sustained MC bandwidth

We compare our proposed memory controller (MC[cache + DMA]) with 2 other alternative memory controller designs, namely cache-only (MC[cache only]) and DMA-only (MC[DMA only]) memory controllers. As the name implies, MC[cache only] uses only caches. The DMA in our proposed MC[cache + DMA] is replaced by a cache of the same size. Similarly, MC[DMA only] replaces the caches in the original design with a DMA of the same size. A comparison of the total execution time for each dataset is shown in Figure 7 (a). The datasets with higher data locality while loading input factor matrices show significant execution time reduction with MC[cache + DMA] compared with MC[DMA only]. Replacing DMAs with caches hinders bulk data transfers when loading, storing, and remapping shards. As a result, MC[cache only] takes longer to perform the above memory operations than MC[cache + DMA]. With MC[cache + DMA], total execution time is improved by 4.5x and 5.8x compared with MC[cache only] and MC[DMA only].

Figure 7 (b) shows the total number of pipeline stalls in the system compared with the total number of operations executed by all the pipelines. Due to irregular external memory access to the input factor matrices, NELL-1 has a significant percentage of pipeline stalls. The rest of the datasets has less than 25% pipeline stalls during computing spMTTKRP on all modes.

Figure 7 (c) shows the sustained memory controller bandwidth while executing spMTTKRP. The memory controller bandwidth is defined as the total amount of data communicated between PEs and memory controllers over the total spMTTKRP execution time for each input tensor. Our memory controller achieves sustained

memory controller bandwidth that is more than twice the peak DRAM bandwidth.

5.5 Impact of Rank of the Factor Matrices

In this section, we evaluate the impact of the rank ($\square$) of the factor matrices on total execution time. We alter $\square$ to 8, 16, and 32 as they are the most common sizes used in tensor decomposition applications [2, 15, 18, 22]. Results from PMS (see Section 5.1.5) suggest that we can use the same accelerator parameters listed in Table 3 by adjusting the interval size of the tensor format.

Figure 8 shows

the total execution time for different datasets over $\square = 8, 16$ and 32. With increasing $\square$, the number of elements in a row in output interval and input factor matrices increases. It leads to

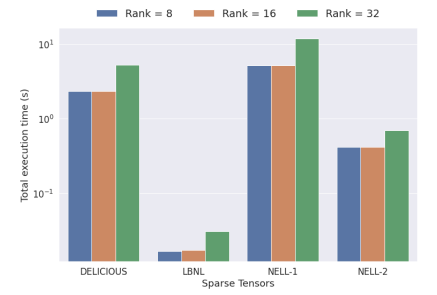


Figure 8: Impact of rank (R)

the following 2 important observations: (1) as $\square$ increases, the caches need multiple cache-line requests to load a single row of factor matrices, and (2) as $\square$ decreases, multiple neighboring rows of factor matrices can fit in a single cache line. In the case of $\square = 8$, the contribution of each subsequent factor matrix row to the computation is very low even if multiple neighboring factor matrices are loaded in one cache line load. Hence we can not observe a significant reduction in total execution time. When $\square = 16$, a single cache line request loads a single row of the input factor matrix. On the other hand, when $\square = 32$, a single input factor matrix row occupies multiple cache lines. Hence a cache miss on a single input factor matrix row results in multiple cache line misses.

6 CONCLUSION AND FUTURE WORK

We proposed the FLYCOO format to reduce the total memory access time on FPGA. The on-the-fly tensor remapping supported by FLYCOO avoids storing intermediate values in the external memory. Using FLYCOO and the proposed FPGA accelerator, we outperform existing benchmarks on a variety of real-world sparse tensors.

In the future, we plan to parallelize the preprocessing algorithm for the tensor format generation. The algorithmic optimizations discussed in Section 4 can be adapted to general purpose hardware such as CPU and GPU. Using our proposed algorithmic optimizations, CPU and GPU can accelerate spMTTKRP by efficiently using their limited internal cache and external memory bandwidth.

ACKNOWLEDGMENTS

This work is supported by the National Science Foundation (NSF) under grants OAC-2209563, CNS-2009057 and in part by DEVCOM Army Research Lab (ARL) under ARL-USC collaborative grant DIRA-ECI:DEC21-CI-037.

REFERENCES

- [1] Alfred V Aho and John E Hopcroft. 1974. The design and analysis of computer algorithms. (1974).
- [2] Zhiyu Cheng, Baopu Li, Yanwen Fan, and Yingze Bao. 2020. A novel rank selection scheme in tensor ring decomposition based on reinforcement learning for deep neural networks. In *ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 3292–3296.
- [3] Sofia Fernandes, Hadi Fanaee-T, and João Gama. 2020. Tensor decomposition for analysing time-evolving social networks: An overview. *Artificial Intelligence Review* (2020), 1–26.
- [4] Ronald L. Graham. 1969. Bounds on multiprocessing timing anomalies. *SIAM Journal on Applied Mathematics* 17, 2 (1969), 416–429.
- [5] Ahmed E. Helal, Jan Laukemann, Fabio Checconi, Jesmin Jahan Tithi, Teresa Ranadive, Fabrizio Petrini, and Jeewhan Choi. 2021. ALTO: Adaptive Linearized Storage of Sparse Tensors. In *Proceedings of the ACM International Conference on Supercomputing (Virtual Event, USA) (ICS '21)*. Association for Computing Machinery, New York, NY, USA, 404–416. <https://doi.org/10.1145/3447818.3461703>
- [6] David Hong, Tamara G. Kolda, and Jed A. Duersch. 2020. Generalized Canonical Polyadic Tensor Decomposition. *SIAM Rev.* 62, 1 (2020), 133–163. <https://doi.org/10.1137/18M1203626> arXiv:<https://doi.org/10.1137/18M1203626>
- [7] Intel. 2022. External Memory Interfaces Intel Arria 10 FPGA IP User Guide. <https://www.intel.com/content/dam/www/programmable/us/en/pdfs/literature/ug/ug-20115.pdf>. Online; accessed 29 January 2022.
- [8] Oguz Kaya and Bora Uçar. 2015. Scalable sparse tensor decompositions in distributed memory systems. In *SC'15: Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*. IEEE, 1–11.
- [9] Fredrik Kjolstad, Stephen Chou, David Lugato, Shoaib Kamil, and Saman Amarasinghe. 2017. Taco: A tool to generate tensor algebra kernels. In *2017 32nd IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE, 943–948.
- [10] Tamara G Kolda and Brett W Bader. 2009. Tensor decompositions and applications. *SIAM review* 51, 3 (2009), 455–500.
- [11] Jiajia Li, Yuchen Ma, and Richard Vuduc. 2018. ParTI! : A Parallel Tensor Infrastructure for multicore CPUs and GPUs. <http://parti-project.org> Last updated: Jan 2020.
- [12] Jiajia Li, Jimeng Sun, and Richard Vuduc. 2018. HiCOO: Hierarchical Storage of Sparse Tensors. In *SC18: International Conference for High Performance Computing, Networking, Storage and Analysis*. 238–252. <https://doi.org/10.1109/SC.2018.00022>
- [13] Jiajia Li, Bora Uçar, Ümit V Çatalyürek, Jimeng Sun, Kevin Barker, and Richard Vuduc. 2019. Efficient and effective sparse tensor reordering. In *Proceedings of the ACM International Conference on Supercomputing*. 227–237.
- [14] Jiajia Li, Bora Uçar, Ümit V. Çatalyürek, Jimeng Sun, Kevin Barker, and Richard Vuduc. 2019. Efficient and Effective Sparse Tensor Reordering. In *Proceedings of the ACM International Conference on Supercomputing (Phoenix, Arizona) (ICS '19)*. Association for Computing Machinery, New York, NY, USA, 227–237. <https://doi.org/10.1145/3330345.3330366>
- [15] Marco Mondelli and Andrea Montanari. 2019. On the connection between learning two-layer neural networks and tensor decomposition. In *The 22nd International Conference on Artificial Intelligence and Statistics*. PMLR, 1051–1060.
- [16] Israt Nisa, Jiajia Li, Aravind Sukumaran-Rajam, Prasant Singh Rawat, Sriram Krishnamoorthy, and P. Sadayappan. 2019. An Efficient Mixed-Mode Representation of Sparse Tensors. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis (Denver, Colorado) (SC '19)*. Association for Computing Machinery, New York, NY, USA, Article 49, 25 pages. <https://doi.org/10.1145/3295500.3356216>
- [17] Israt Nisa, Jiajia Li, Aravind Sukumaran-Rajam, Richard Vuduc, and P. Sadayappan. 2019. Load-Balanced Sparse MTTKRP on GPUs. In *2019 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*. 123–133. <https://doi.org/10.1109/IPDPS.2019.00023>
- [18] Nicholas D. Sidiropoulos, Lieven De Lathauwer, Xiao Fu, Kejun Huang, Evangelos E. Papalexakis, and Christos Faloutsos. 2017. Tensor Decomposition for Signal Processing and Machine Learning. *IEEE Transactions on Signal Processing* 65, 13 (2017), 3551–3582. <https://doi.org/10.1109/TSP.2017.2690524>
- [19] Shaden Smith, Jee W. Choi, Jiajia Li, Richard Vuduc, Jongsoo Park, Xing Liu, and George Karypis. 2017. FROST: The Formidable Repository of Open Sparse Tensors and Tools. <http://frostt.io/>
- [20] Shaden Smith, Niranjan Ravindran, Nicholas D. Sidiropoulos, and George Karypis. 2015. SPLATT: Efficient and Parallel Sparse Tensor-Matrix Multiplication. In *2015 IEEE International Parallel and Distributed Processing Symposium*. 61–70. <https://doi.org/10.1109/IPDPS.2015.27>
- [21] Nitisri Srivastava, Hanchen Jin, Shaden Smith, Hongbo Rong, David Albonese, and Zhiru Zhang. 2020. Tensaurus: A Versatile Accelerator for Mixed Sparse-Dense Tensor Computations. In *2020 IEEE International Symposium on High Performance Computer Architecture (HPCA)*. 689–702. <https://doi.org/10.1109/HPCA47549.2020.00062>
- [22] Fuxi Wen, Hing Cheung So, and Henk Wymeersch. 2020. Tensor decomposition-based beamspace esprit algorithm for multidimensional harmonic retrieval. In *ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 4572–4576.
- [23] Sasindu Wijeratne, Rajgopal Kannan, and Viktor Prasanna. 2021. Reconfigurable Low-latency Memory System for Sparse Matricized Tensor Times Khatri-Rao Product on FPGA. In *2021 IEEE High Performance Extreme Computing Conference (HPEC)*. 1–7. <https://doi.org/10.1109/HPEC49654.2021.9622851>
- [24] Xilinx. 2022. Alveo U250 Data Center Accelerator Card. <https://www.xilinx.com/products/boards-and-kits/alveo/u250.html>. Online; accessed 29 January 2022.
- [25] Xilinx. 2022. UltraFast Design Methodology Guide for Xilinx FPGAs and SoCs. <https://docs.xilinx.com/r/2021.2-English/ug949-vivado-design-methodology/>. Online; accessed 1 August 2022.
- [26] Xilinx. 2022. UltraScale Architecture-Based FPGAs Memory IP v1.4. https://www.xilinx.com/support/documentation/ip_documentation/ultrascale_memory_ip/v1_4/pg150-ultrascale-memory-ip.pdf. Online; accessed 29 January 2022.
- [27] Xilinx. 2022. Vivado Design Suite User Guide. <https://docs.xilinx.com/v/u/2018.2-English/ug910-vivado-getting-started>. Online; accessed 29 January 2022.