

SparseLNR: Accelerating Sparse Tensor Computations Using Loop Nest Restructuring

Adhitha Dias

Electrical and Computer Engineering
Purdue University
West Lafayette, IN, USA
kadhitha@purdue.edu

Charitha Saumya

Electrical and Computer Engineering
Purdue University
West Lafayette, IN, USA
cgusthin@purdue.edu

Kirshanthan Sundararajah

Electrical and Computer Engineering
Purdue University
West Lafayette, IN, USA
ksundar@purdue.edu

Milind Kulkarni

Electrical and Computer Engineering
Purdue University
West Lafayette, IN, USA
milind@purdue.edu

Abstract

Sparse tensor algebra computations have become important in many real-world applications like machine learning, scientific simulations, and data mining. Hence, automated code generation and performance optimizations for tensor algebra kernels are paramount. Recent advancements such as the Tensor Algebra Compiler (TACO) greatly generalize and automate the code generation for tensor algebra expressions. However, the code generated by TACO for many important tensor computations remains suboptimal due to the absence of a scheduling directive to support transformations such as *distribution/fusion*.

This paper extends TACO's scheduling space to support kernel distribution/loop fusion in order to reduce asymptotic time complexity and improve locality of complex tensor algebra computations. We develop an intermediate representation (IR) for tensor operations called *branched iteration graph* which specifies breakdown of the computation into smaller ones (kernel distribution) and then fuse (loop fusion) outermost dimensions of the loop nests, while the innermost dimensions are distributed, to increase data locality. We describe exchanges of intermediate results between space iteration spaces, transformation in the IR, and its programmatic invocation. Finally, we show that the transformation can be used to optimize sparse tensor kernels. Our results show that this new transformation significantly improves the performance of several real-world tensor algebra computations compared to TACO-generated code.

CCS Concepts: • Software and its engineering → Source code generation; Domain specific languages.

Keywords: Sparse Tensor Algebra, Loop Transformations, Kernel Distribution, Loop Fusion

ACM Reference Format:

Adhitha Dias, Kirshanthan Sundararajah, Charitha Saumya, and Milind Kulkarni. 2022. SparseLNR: Accelerating Sparse Tensor Computations Using Loop Nest Restructuring. In *2022 International Conference on Supercomputing (ICS '22)*, June 28–30, 2022, Virtual Event, USA. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3524059.3532386>

1 Introduction

Sparse tensor algebra is used in many machine learning domains such as graph neural networks [13, 14, 27]. Tensors are a generalization of matrices and are typically represented using n-dimensional arrays. However, when used to represent large graph-like structures, representing a tensor with a dense array is wasteful, as most values in the tensor are zero. In such cases, programmers use *compressed* representations of these *sparse tensors*.

The problem of compiler optimizations for sparse codes is well known [5, 17, 18, 34, 35], and there are several challenges that compilers face: (1) tensor computations have to deal with specific data formats; (2) load imbalance can arise due to irregular structure; and (3) data locality issues arise due to the sparsity of the tensors. TACO provides a compiler for automatically generating kernels for dense and sparse tensor algebra operations [17]. The tensor application is expressed in terms of three languages: a tensor algebra language for expressing the computation (Section 2.1), a data representation language for specifying how sparse tensors are compressed, and a scheduling language that specifies the schedule of the computation (Section 2.3).

The scheduling language provides the ability to define different schedules for computations depending on tensors' dimensionality and sparsity patterns, because one schedule



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

ICS '22, June 28–30, 2022, Virtual Event, USA

© 2022 Copyright held by the owner/author(s).

ACM ISBN 978-1-4503-9281-5/22/06.

<https://doi.org/10.1145/3524059.3532386>

may not fit all data formats and datasets. This allows the separation of *algorithmic* specification from the *scheduling* details of the computation. Once both are specified, code can be generated to implement the desired algorithm and schedule.

One important consequence of TACO's code generation is that the asymptotic complexity of the kernels grows with the number of index variables in the tensor index notation [2]. For example, the complexity of $A_{ij} = \sum_k B_{ij} \cdot C_{ik} \cdot D_{jk}$ ¹ is $O(\text{nnz}(B_{IJ})K)^2$, where B is sparse. If this example is extended with an additional computation, as in $A_{il} = \sum_{kj} B_{ij} \cdot C_{ik} \cdot D_{jk} \cdot E_{jl}$, then the complexity is $O(\text{nnz}(B_{IJ})KL)$ ³—and this complexity increases with each additional index variable. Hence, with increasing terms in the tensor expression, the asymptotic complexity of the resulting code blows up.

Interestingly, this asymptotic blowup is a consequence of doing multiple tensor operations in a single kernel. The computation could instead be expressed as two separate kernels, with the result of the first computation stored in a temporary tensor: $T_{ij} = \sum_k B_{ij} \cdot C_{ik} \cdot D_{jk}$; $A_{il} = \sum_j T_{ij} \cdot E_{jl}$. This computation has a complexity of $O(\text{nnz}(B_{IJ})(K + L))$. However, writing complex computations as *separate* TACO expressions has two downsides. First, it is no longer possible to apply schedule transformations, such as outer-loop parallelization, across the entire computation. Second, if the computations require large temporaries, materializing them results in performance degradation due to exhaustion of the last-level cache.

The correct schedule looks like neither the single-kernel approach nor the separate-kernels approach. Instead, it performs a single outer loop over the i and j indices and then performs the inner loop of the first kernel, stores the results in a temporary, then uses those results in the inner loop of the second kernel. This approach has an asymptotic complexity of $O(\text{nnz}(B_{IJ})(K + L))$, comparable to the separate kernel approach, but because the temporary is only live within the inner loops, it is much smaller and hence can fit in cache. Moreover, the overall computation is a single loop nest, allowing for the outer loops to be parallelized, tiled, etc.

The above schedule transformation is analogous to ones in *dense* tensor contraction that combine loop distribution and fusion to create imperfectly-nested loops [4]. But it is less clear how to use this technique on sparse loops for several reasons: (i) analysis is harder, because of the sparse tensor accesses and non-affine bounds, as polyhedral techniques do not work due to the use of dynamic array bounds in loops; (ii) producing good schedules is harder because performance can degrade by forcing a sparse tensor to be processed using dense iteration; and (iii) code generation is harder, as you

need to deal with storage format-specific iteration machinery. For example, a sparse matrix and dense matrix multiplication (SpMM) may be performed with a sparse matrix of Compressed Sparse Row format (CSR), Coordinate format (COO), etc. [8]. Hence, the compiler needs to tackle format-specific access patterns to generate code for SpMM for different storage formats.

Our insight for tackling the complex scheduling transformations needed to avoid asymptotic blowup while preserving locality, is to use dense temporaries and introduce *Sparse Loop Nest Restructuring (SparseLNR)*⁴ for tensor computations. Crucially, these transformations can co-exist with TACO's other scheduling primitives [30].

This paper introduces a new representation called *branched iteration graphs* that support imperfect nesting of sparse iteration. Given this representation, our compiler can restructure sparse tensor computations to remove the asymptotic blowup in sparse tensor algebra code generation while delivering good locality. Our specific contributions are;

Branched iteration graph for tensor multiplications We generalize the iteration graph intermediate representation (IR) of TACO to support imperfectly nested loop structures.

Branch IR transformation We design a sparse tensor transformation that transforms iteration graphs to express fusion and distribution.

New scheduling primitives We introduce a new scheduling primitive that lets programmers integrate fusion and distribution into TACO schedules.

For several real-world tensor algebra computations (Described in Section 6.2) on various datasets (Shown in Table 1), using our new representation and transformations, we show that SparseLNR can achieve 1.23–1997x (single-thread) and 0.86–1263x (multi-thread) speedup over baseline TACO schedules, and 0.27–3.21x (single-thread) and 0.51–3.16x (multi-thread) speedup over TACO schedules of manually separated computations.

2 Background

This section provides the necessary background to understand sparse tensor algebra computations and different ways to schedule those computations.

2.1 Tensor Index Notation

Tensor index notation is a high-level representation used for describing tensor algebra expressions [17]. Throughout the paper we will be using both the *standard notation* and *tensor index notation* to denote tensor operations. For instance, the tensor computation $A_{ik} = \sum_j B_{ij} C_{jk}$ written in standard notation is equivalent to $A(i, k) = B(i, j) * C(j, k)$, written in tensor index notation.⁵ Here, all the tensors are matrices

¹Highlighted tensors denote sparse tensors.

² $\text{nnz}(B_{IJ})$ denotes the nonzero values of the sparse tensor B bounded by the hierarchical accesses i and j .

³ K and L denote the number of iterations or the dimensionality of k and l dimensions respectively.

⁴<https://github.com/adhithadidas/SparseLNR>

⁵This computation is classic matrix-matrix multiply.

and indices i, j , and k are used to iterate over matrices A, B , and C . In this computation, index j must be iterated over the intersection of second dimension coordinates of B and first dimension coordinates of C , whereas index i and k must be iterated over the first and second dimension coordinates of B and C respectively.

2.2 Iteration Graph

We first summarize TACO's iteration graph representation, which Kjolstad *et al.* describes in great detail [17]. When computing the tensor expression $A_{ij} = \sum_k B_{ij} C_{ik} D_{jk}$, coordinates (i, j) of B , coordinates (i, k) of C , and (j, k) of D need to be iterated. An iterator on indices (i, j, k) can iterate through all the coordinates of B, C , and D and store the results in A . TACO represents the iteration space of a tensor expression using an iteration graph, an intermediate representation that defines tensor access patterns of indices.

Figure 1 shows a few examples of iteration graphs. For example, a tensor expression $A_{ij} = \sum_k B_{ij} C_{ik} D_{jk}$ results in an iteration graph as shown in Figure 1a such that the indices lay in i, j, k order. Here, the order of j and k is not strict if C and D are dense. Figures 1b and 1c are the iteration graphs of tensor expressions $A_{ik} = \sum_{kl} B_{ikl} C_{lj} D_{kj}$ and $y_i = \sum_{jk} B_{ij} C_{jk} v_k$ respectively.

Nodes in the iteration graph represent indices of tensor index notation. In other words, the iteration graph is a directed graph of these indices. These indices of the graph are topologically sorted such that it imposes sparse iteration constraints (*i.e.*, constraints that define the sparse tensor access patterns of indices due to lack of random access in general). Each index in the iteration graph can be expressed as a loop to iterate through a tensor. Therefore, a given tensor multiplication can be computed using nested loops, where each loop corresponds to an index variable in the iteration graph.

Definition 2.1. An iteration graph is a directed graph $G = (V, P)$ where $V = v_1, v_2, \dots, v_n$ defines the set of index variables in the tensor index notation, and $P = p_1, p_2, \dots, p_n$ defines the set of tensor paths, a tensor path is a tuple of index variables associated with a particular tensor variable.

2.3 Scheduling Primitives

A tensor expression can have multiple valid schedules of computation as there are different valid orders of iterating through indices and multiple parallelization strategies. Kjolstad *et al.* [17] and Senanayake *et al.* [30] have introduced scheduling primitives for tensor computations, with which the user can describe schedules to execute a given tensor computation. The scheduling primitives in TACO are the *split* directive to split a loop into two loops for tiling, *collapse* directive to collapse doubly nested loops into a single loop for balancing load among threads, *reorder* directive⁶

⁶Also referred to as *permute* directive in the literature.

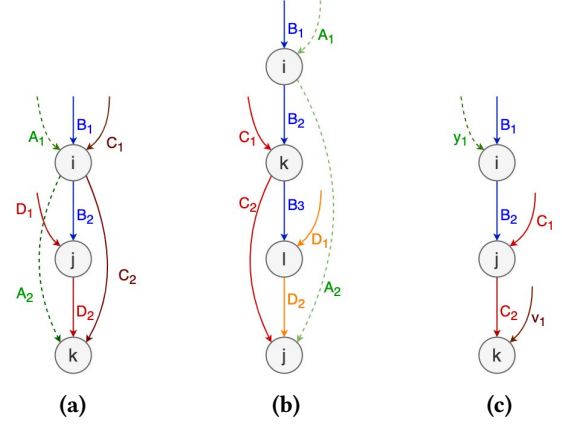


Figure 1. Iteration graphs (a) SDDMM kernel $A_{ij} = \sum_k B_{ij} C_{ik} D_{jk}$ (b) Khatri-Rao product (MTTKRP) kernel $A_{ik} = \sum_{kl} B_{ikl} C_{lj} D_{kj}$ (c) Sparse matrix vector multiplication (SpMV) kernel preceded by another SpMV kernel $y_i = \sum_{jk} B_{ij} (C_{jk} v_k)$

to reorder loops, *unroll* directive to perform loop unrolling, *parallelize* directive to parallelize loops with OpenMP-based multithreaded execution (for outer loops) or vectorized execution (for inner loops). Furthermore, Kjolstad *et al.* [16] added *precompute* scheduling directive to use intermediate dense workspaces to remove sparse accesses when storing data values to output tensors.

3 Overview

There are a number of factors taken into account when deciding whether to apply transformations across kernels. If the working sets are small, running the kernels separately with good schedules defined on each individual kernel maybe faster than a fused kernel. But if the working sets are large resulting in large temporaries that do not fit in caches, it is better to fuse two kernels and try to maximize the data reuse by using the results produced by the first kernel and execute part of the second kernel without waiting for the completion of the first kernel.

3.1 Motivating Example

Consider the computation, $A = \text{Sparse } B \odot (CD) \cdot E$ that is used in graph embedding and graph neural networks [27, 36]. The Hadamard product, or element-wise product, is denoted by $\odot$ and matrix multiplication is denoted by $\cdot$. We can perform the above computation in the following order with fine-grained smaller tensor operations. $T = \text{gemm}(C, D)$, $\text{Sparse } U = \text{spelmm}(\text{Sparse } B, T)$, $A = \text{spmm}(\text{Sparse } U, E)$. Here, *gemm* stands for the generalized matrix multiplication, *spelmm* stands for sparse element-wise multiplication, and *spmm* stands for sparse matrix multiplication. Materialization of these intermediate tensors leads to multiple issues:

1. Dense matrix multiplication results in redundant calculations and unnecessary increase in asymptotic complexity, because later it is sampled by the Sparse B matrix.

```

1  int32_t jY = 0;
2  for (int32_t i = 0; i < C1_dimension; i++) {
3    for (int32_t jB = B2_pos[i]; jB < B2_pos[(i + 1)]; jB++) {
4      int32_t j = B2_crd[jB];
5      double tkY_val = 0.0;
6      for (int32_t k = 0; k < D2_dimension; k++) {
7        tkY_val += B_vals[jB] * C_vals[i,k] * D_vals[j,k];
8      }
9      Y_vals[jY] = tkY_val;
10     jY++;
11   }
12 }

```

$$(a) Y_{ij} = \sum_k B_{ij} C_{jk} D_{kj}$$

```

1  for (int32_t i = 0; i < C1_dimension; i++) {
2    for (int32_t jB = B2_pos[i]; jB < B2_pos[(i + 1)]; jB++) {
3      int32_t j = B2_crd[jB];
4      for (int32_t l = 0; l < E2_dimension; l++) {
5        double tkA = 0.0;
6        for (int32_t k = 0; k < D2_dimension; k++) {
7          tkA += B_vals[jB] * C_vals[i,k] * D_vals[j,k] * E_vals[j,l];
8        }
9        A_vals[i,l] = A_vals[i,l] + tkA;
10     }
11   }
12 }

```

$$(c) A_{il} = \sum_{jk} B_{ij} C_{jk} D_{kj} E_{jl}$$

```

1  for (int32_t i = 0; i < Y1_dimension; i++) {
2    for (int32_t jY = Y2_pos[i]; jY < Y2_pos[(i + 1)]; jY++) {
3      int32_t j = Y2_crd[jY];
4      for (int32_t l = 0; l < E2_dimension; l++) {
5        A_vals[i,l] = A_vals[i,l] + Y_vals[jY] * E_vals[j,l];
6      }
7    }
8 }

```

$$(b) A_{il} = \sum_j Y_{ij} E_{jl}$$

```

1  for (int32_t i = 0; i < C1_dimension; i++) {
2    for (int32_t jB = B2_pos[i]; jB < B2_pos[(i + 1)]; jB++) {
3      int32_t j = B2_crd[jB];
4      double Y_val = 0.0;
5      for (int32_t k = 0; k < D2_dimension; k++) {
6        Y_val += B_vals[jB] * C_vals[i,k] * D_vals[j,k];
7      }
8      for (int32_t l = 0; l < E2_dimension; l++) {
9        A_vals[i,l] = A_vals[i,l] + Y_val * E_vals[j,l];
10     }
11   }
12 }

```

$$(d) A_{il} = \sum_j (\sum_k B_{ij} C_{jk} D_{kj}) E_{jl}$$

Figure 2. Different schedules of executing $A_{il} = \sum B_{ij} \cdot C_{jk} \cdot D_{kj} \cdot E_{jl}$. The code snippet 2b executed immediately after the code snippet 2a computes the same result as fused operations explained in the code snippets 2c and 2d. Here, the code snippet 2c has a perfectly nested loop structure while the code 2d describes a nested loop structure for the same computation.

- Values are produced long before they are consumed, which may cause them to be evicted from caches.
- Having intermediate tensors is justifiable if intermediate results are needed for some other computation, nevertheless a single kernel maybe needed for faster operation.

Introducing *kernel fusion* to tensor computations can reduce these issues [27]. In this section, we discuss different schedules for performing the computation $A = \text{Sparse } B \cdot (CD) * E$, and motivate the need for supporting loop fusion for sparse tensor computations.

First, we discuss the opportunities for distribution in the running example using a fused kernel with high asymptotic complexity (Section 3.1.1). Next, we discuss opportunities for fusion when the computation is split into two smaller kernels (Section 3.1.2). Finally, in Section 3.1.3 we discuss how we can exploit these different scenarios to construct a distributed (versus the fused kernel in Section 3.1.1) and then fused (as compared to the kernel in Section 3.1.2) implementation.

3.1.1 Asymptotic expensive fused kernel. The computation $A_{il} = \sum B_{ij} \cdot C_{ik} \cdot D_{jk} \cdot E_{jl}$ can be fully realized using a nested loop iterator defined by all indices i, j, k , and l . The generalized way of producing kernels for a tensor multiplication of this kind in TACO is by generating an iteration graph (see Section 2.2). Since the iteration graph contains all the indices in a linear tree pattern, TACO generates a kernel

as in Figure 2c, with time complexity of $O(nnz(B_{IJ})KL)$ due to the quadruple linearly nested loops (lines 1–6).

3.1.2 Asymptotically inexpensive distributed kernels.

However, the computation $A_{il} = \sum_{kj} B_{ij} \cdot C_{ik} \cdot D_{jk} \cdot E_{jl}$ can be performed by evaluating two smaller kernels: sampled dense-dense matrix multiplication (SDDMM): $Y_{ij} = \sum_k B_{ij} \cdot C_{ik} \cdot D_{jk}$ followed by SpMM: $A_{il} = \sum_j Y_{ij} \cdot E_{jl}$. As these separate kernels are triply nested loops (lines 2–6 in Figure 2a and 1–3 in Figure 2b), they have lower asymptotic complexity.

Here, the Hadamard product, in SDDMM, results in Y_{ij} matrix's sparse structure to be same as B_{ij} . Therefore, the asymptotic complexity of performing two tensor computations with an intermediary matrix Y_{ij} is $O(nnz(B_{IJ})(K+L))$. These separate kernels can be realized through loop distribution of the kernel from Section 3.1.1.

Although we achieve a lower asymptotic complexity, we are using an intermediary tensor to pass values between SDDMM and SpMM, Hence, we miss the opportunity to exploit the temporal locality of the operation. The tensor contraction computed using linearly nested loops in Section 3.1.1. is expensive because of the high degree of nesting in the computation and the redundant duplicate computations, but may still be good for memory-constrained systems because the computation does not require any memory for storing intermediate results.

Using a temporary tensor to hold the result of the SDDMM operation is acceptable as long as the dimensionality of the index variables i and j , and the density of the temporary tensor, are small. The code generation algorithm in TACO is limited to generating sequential code when the output tensor is of sparse format, (see jY variable in Figure 2a). The kernel is sequential because the data format used to store the results of the computation limits random accesses. Here, the output of SDDMM operation is sparse (and the output from SpMM is dense) in which case we cannot parallelize the outermost loop of the SDDMM operation in separate kernel execution whereas the kernel in Figure 2c can be parallelized because the output of the combined kernel is dense. This is another valid reason to prefer the single kernel implementation despite its high asymptotic complexity.

3.1.3 Fused kernel with low asymptotic complexity.

Since both the kernels $Y_{ij} = \sum B_{ij} \cdot C_{ik} \cdot D_{jk}$ in Figure 2a and $A_{il} = \sum Y_{ij} \cdot E_{jl}$ in Figure 2b have the same access patterns in their two outer-most loops, we can fuse them as shown in Figure 2d, removing the use of the intermediary tensor to pass the values between the two separate kernels as explained in Section 3.1.2. This execution has a time complexity of $O(nnz(B_{IJ})(K + L))$, and at the same time removes the usage of a large tensor temporary by using an imperfectly nested loop structure (Lines 1–2,5 and 8 in Figure 2d).

Note that this partially-fused kernel provides the best of both worlds. Like the separate kernel approach, it has low asymptotic complexity. Like the fused kernel approach, it has good locality (since the temporaries only need to store data from the inner loops, their sizes much smaller and the reuse distances are reduced). Furthermore, because the outer loops of the partially fused are shared between both computations, and there is no longer a loop-carried dependence for SDDMM, the overall kernel can be parallelized in the same way as the kernel of Figure 2c.

3.2 Our approach: SparseLNR

While the schedule of computation in Figure 2d provides both good asymptotic complexity and good locality, no existing system can automatically generate this type of schedule when generating code for sparse computations. TACO only handles “linear” iteration graphs that yield perfectly-nested loops, and hence cannot handle the partially-fused, imperfectly nested loop structure needed by our example. On the other hand, prior work on distribution and fusion for tensor computations [4], can support this type of code structure only for operations on dense tensors.

SparseLNR provides mechanisms for generating the code in Figure 2d from a high level representation of the computation as well as scheduling directives that inform the structure of the code. We introduce several components to perform this code generation and Section 4 discuss them in detail.

1. We introduce a new representation called a *branched iteration graph* that allows the representation of partially-fused iteration structures, where some loops in a loop nest are common between computations and others are separate. Hence, this graph represents imperfect nesting. We carefully place constraints on these graphs to ensure that the requirements of nested iteration over sparse structures are met. The branched iteration graph is described in more detail in Section 4.2.
2. We introduce new scheduling primitives for loop distribution and fusion that allow programmers to *generate* the branched iteration graph by applying scheduling transformations to linear TACO iteration graph. We describe the primitives and describe how they systematically transform a branched iteration graph in Section 4.3.1.
3. We adapt TACO’s code generation strategies to the branched iteration graph, allowing SparseLNR to generate sparse iteration code for tensor kernels that have had our distribution and fusion transformations applied to them. We discuss code generation in Section 4.4.

4 Detailed Design

This section describes the key components of SparseLNR. Section 4.1 describes SparseLNR’s new branched iteration graph representation. Section 4.2 shows how partial fusion is represented through iteration graph transformations. Section 4.3.1 explains how scheduling directives can guide partial fusion while still composing with TACO’s existing scheduling language. Finally, Section 4.4 explains how SparseLNR generates code.

4.1 Representation

SparseLNR uses a *branched iteration graph* to represent sparse tensor algebra kernels, which is an extension to the concrete index notation described in [16]. A branched iteration graph can be understood as an iteration graph with branches in index access patterns. By transforming the linear index tree iteration graph generated by TACO to a branched iteration graph in the context of tensor multiplication, we try to remove the asymptotic complexity that arises from perfectly linearly nested loops in dense/sparse iterations.

Definition 4.1. A branched iteration graph is a directed graph $G = (V, G_p, G_c, P)$, where V is a set of *unbranched* indices, organized as a sequence starting from the root of the iteration graph that then has two children graphs, G_p (producer) and G_c (consumer), that define the two branches of G , where G_p and G_c themselves are branched iteration graphs, such that there is a *dependence edge* from G_p to G_c and a *boundary* between V and (G_c, G_p) . The *dependence edge* tracks the common set of indices in G_p and G_c . $P = p_1, p_2, \dots, p_n$ defines the set of tensor paths, a tuple of indices associated with a particular tensor variable.

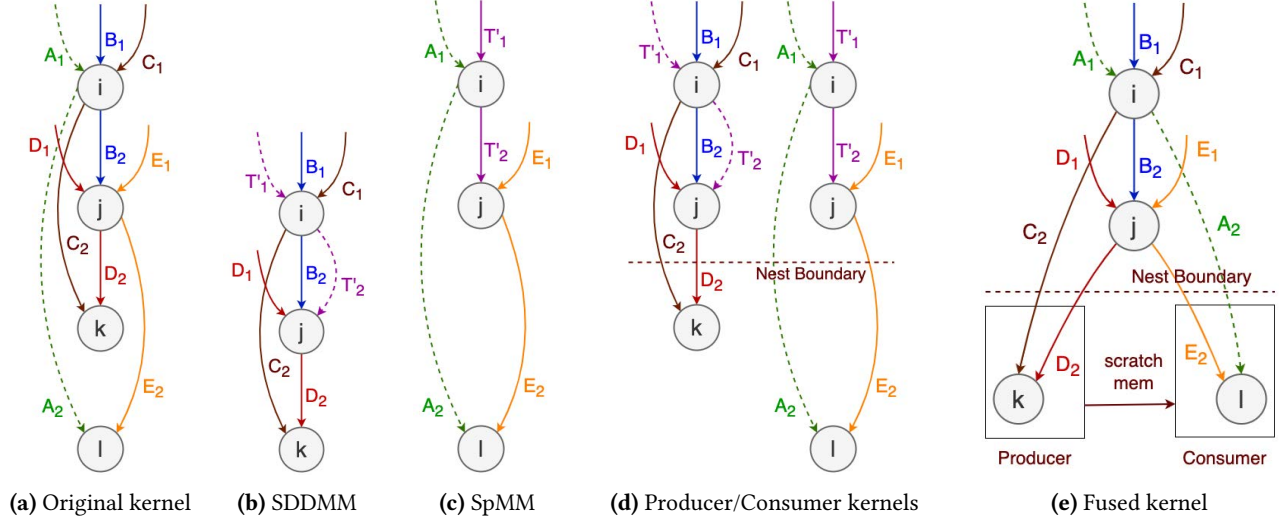


Figure 3. loopfuse transformation performed on $A_{il} = \sum_{jk} B_{ij} C_{ik} D_{jk} E_{jl}$

Intuitively, where a TACO iteration graph corresponds to a perfectly-nested loop where the order of the vertices in the graph corresponds to the nesting order of the loops, a branched iteration graph represents an imperfectly nested loop. V corresponds to the common outer loops, just as in a TACO graph, while G_p and G_c correspond to the inner loop nests (which can themselves be imperfectly nested). For example, in Figure 3e, V refers to the set of indices $\{i, j\}$, and G_c , G_p refer to the boxes *Producer* and *Consumer*, respectively.

4.2 Branched Iteration Graph Transformation

In Section 3.1 we saw how we could perform loop fusion or distribution for a sparse tensor algebra computation. We recognize this pattern in index traversal and exploit it to generate the branched iteration graph. We name this pattern recognition algorithm *fusion after distribution* because it proceeds in two steps as described in Algorithm 1: (i) distributing the perfectly-nested indices in the iteration graph, and then (ii) fusing the common indices.

Topologically sorted iteration graph. The iteration graph in Figure 3a relates to the index expression $A(i, l) = B(i, j) * C(i, k) * D(j, k) * E(j, l)$, where B is sparse. We denote this kernel as $\langle \text{SDDMM}, \text{SpMM} \rangle$. The indices here are topologically ordered such that the ordering of the indices are constrained by the sparsity patterns of the sparse tensors. The ordering $i \rightarrow j \rightarrow k \rightarrow l$ would be consistent with the access patterns of all the tensors $i \rightarrow l$ in A , $i \rightarrow j$ in B , $i \rightarrow k$ in C , $j \rightarrow k$ in D , and $j \rightarrow l$ in E . However, there should be a hard ordering imposed on i and j index variables because j cannot be accessed without accessing i first. The index access patterns of the tensor access variables are marked in the graph as paths. A_1 denotes the first access dimension of the A and A_2 denotes the second access dimension of the A tensor. The fusion algorithm requires the iteration graph in Figure 3a and the tensor index notation expression

Algorithm 1 Loop fusion after distribution

Input: Topologically Ordered Iteration Graph $G_I = (I_G, P)$
Input: Index Expression $Expr$: $A_{out} = A_1 * A_2 * \dots * A_n$
Input: Bool *recursive*
Output: Branched Iteration Graph G'_I

```

1: fusible = isFusible( $G_I$ )
2: if !fusible then return  $G_I$ 
3:  $P_{T'} = P_{A_{out}} - P_{A_n} \triangleright$  index path for temporary tensor  $T'$ 
4:  $G_{I-Producer}, ProducerExpr_{temp} := T'(P_{T'}) = Expr \setminus A_n$ 
5:  $G_{I-Consumer}, ProducerExpr_{temp} := A_{out} = T'(P_{T'}) * A_n$ 
6: if recursive then
7:    $G_{I-Producer} = recursiveCall(G_{I-Producer},$ 
      $\hookrightarrow ProducerExpr_{temp}, recursive)$ 
8:  $List_{I-Producer} = GetIndices(G_{I-Producer})$ 
9:  $List_{I-Consumer} = GetIndices(G_{I-Consumer})$ 
10: Define:  $I_{sharable} = \emptyset$ 
11: for Each  $i \in I_G$  do
12:   if  $i \in List_{I-Producer}$  and  $i \in List_{I-Consumer}$  then
13:      $I_{sharable} = I_{sharable} \cup i$ 
14:   else break;
15: Define:  $I_{fusible} = \emptyset$ 
16: for  $i \leftarrow 1$  to  $N$  do
17:   if  $i \notin I_{sharable}$  and  $i \in List_{I-Producer}$  and
      $\hookrightarrow i \in List_{I-Consumer}$  then
18:      $I_{fusible} = I_{fusible} \cup i$ 
19:   else break;
20: Define:  $T(P_{I_{fusible}})$ 
21:  $ProducerExpr := T(P_{I_{fusible}}) = T'(P_{T'})$ 
22:  $ConsumerExpr := A_{out} = T(P_{I_{fusible}}) * A_n$ 
23: return  $GraphRewrite(G_I, I_{sharable},$ 
      $\hookrightarrow ProducerExpr, ConsumerExpr)$ 

```

$A(i, l) = B(i, j) * C(i, k) * D(j, k) * E(j, l)$. We identify the iteration graph as fusible if there are indices that are only

present in the last tensor and the output tensor in the tensor expression (line 1 of the Algorithm 1).

Distribution into two kernels. The description of the tensor kernel above captures all the information of performing kernel executions SDDMM: $T'(i, j) = B(i, j) * C(i, k) * D(j, k)$ and SpMM: $A(i, l) = T'(i, j) * E(j, l)$ sequentially. (Notice that separation of kernels requires a temporary matrix T') Therefore, we can recover the separate 2 smaller kernels that would yield the same result given the larger tensor expression. We denote the first kernel as the producer and the second kernel as the consumer. To find these separate smaller kernels, we need to remove the last tensor $E(j, l)$ from the original expression. Line 8 of the Algorithm 1 creates the producer index expression and iteration graph for the tensor computation performed first (SDDMM in our running example) by removing the last tensor from the original expression, and then line 9 of the Algorithm 1 creates the consumer index expression and iteration graph for the tensor computation that is performed second (SpMM in our running example) by adding it back to the producer's expression. These 2 separate kernels would have iteration graphs shown in Figures 3b and 3c respectively. We perform this recovery of the two separate operations in order to identify the fusible and shared indices between two separate tensor operations as we will further explain in a next paragraph.

Fusing common loops. Once we have the iteration graphs for the separate kernels we reason about them together (See 3d). We reason that both the sparse iterations need to iterate through the space using index variables i and j . Also, iteration space defined by the index k is iterated only by the SDDMM operation, and the iteration space defined by the index l is only iterated by the SpMM operation. But those iterations over index k and l need to happen one after the other. The producer-consumer dependence must be satisfied such that the values consumed by the consumer must have been produced by the producer before its use. The values shared between the producer and consumer can be stored in an intermediate scratch memory. Furthermore, the comparison of the two graphs, the producer graph and the consumer graph, helps identify the indices that can and cannot be shared among the iterations.

The producer graph in Figure 3b and the consumer graph in Figure 3c have a common prefix defined by some indices in their iteration graphs. We run a prefix match to identify the shared indices by the two kernels (lines 8–14 of the Algorithm 1), in Figure 3d. We see that both i and j indices are shared, and the other variables are not shared. The addition of indices i and j to the set of sharable indices is described in lines 10–14 of the Algorithm 1. We identify this point as a *nest boundary* in the iteration graph 3d, and denote the indices above the *nest boundary* as fusible. The final output of executing the *fusion after distribution algorithm* is a

branched iteration graph. Therefore, if the algorithm is applied recursively (see the kernel <SDDMM, SpMM, GEMM> in the benchmark Section 6.2) on the producer (lines 6–7 of the Algorithm 1), our algorithm can still match the prefix even if the producer graph is already branched.

Materializing temporary variables. The next step of the Algorithm is to identify the indices that cannot be fused as outermost loops but are common to the producer and the consumer. In Figure 3d we see that there are no common variables below the *nest boundary*. The variables that are below the *nest boundary* line and common to both the producer and consumer define the dimensions of the temporary variable that is shared between them. For the case of <SDDMM, SpMM> described in Figure 3d, since no indices are common below the *nest boundary* line, we can define the temporary as a scalar. However, for the same case of <SDDMM, SpMM> described in Section 4.3.1, where transpose of D is used to define the computation, we can see that index j is a common index below the *nest boundary* line. Therefore, the algorithm defines a temporary vector bounded by the size of the index j . Lines 15–19 of the Algorithm 1 explain how we perform the identification of the common indices below the *nest boundary*, and line 20 defines this temporary variable.

Rewrite the iteration graph. After we find the fusible indices, shared indices and define the temporary variable, we define the producer expression and consumer expression using the temporary variable that is shared between the producer and the consumer (lines 21, 22 of the Algorithm 1). Then, we rewrite the iteration graph to model this behavior with the temporary variable, the producer and the consumer (see Figure 3d) which would eventually generate the code shown in Figure 2d for our running example.

4.3 Scheduling

In this section we describe, (1) the invocation of scheduling transformation and (2) the impact it has on the space of possible schedules.

4.3.1 Scheduling Directive. SparseLNR introduces a new scheduling directive to TACO. The user can call the `loopfuse` scheduling transformation as shown in Figure 4b with other scheduling directives. Here, **1** refers to applying the algorithm once. By passing **2** or a higher number, the algorithm can be applied recursively.

Sometimes it is necessary to combine `loopfuse` with other TACO scheduling directives. Hence, it is important that our new directive compose with the existing scheduling language. For example, applying Algorithm 1 to the tensor expression $A(i, l) = B(i, j) * C(i, k) * D(k, j) * E(j, l)$ would not yield the code in Figure 2d by default because now the access pattern of the D matrix is different since we are using the transpose of D for this example. This difference results in a

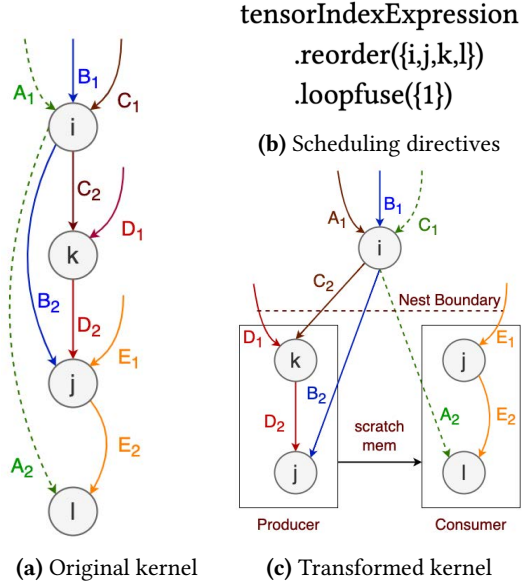


Figure 4. The loopfuse transformation performed on $A_{il} = \sum_{jk} B_{ij}C_{ik}D_{kj}E_{jl}$.

different iteration graph as shown in Figure 4a, because now the iteration graph needs to preserve the ordering of $i \rightarrow j$ for B, $i \rightarrow k$ for D, $k \rightarrow j$ for D, $j \rightarrow l$ for E, and $i \rightarrow l$ for A, with a hard ordering of $i \rightarrow j$ because B is a sparse matrix. Applying the *fusion after distribution* algorithm would result in an iteration graph as depicted in Figure 4c.

However, since D is dense, there is no hard constraint on the ordering of indices k and j . Therefore, to arrive at the code in Figure 2d, a loop reordering can be performed before the loopfuse scheduling directive (See Figure 4b).

The *nest boundary* between the branching point in the iteration graph constrains loop reordering between the *nest boundary*. But loop reordering can be still allowed in indices within *nest boundaries*.

4.3.2 Scheduling Space. In our current implementation, if two kernels have n common outer loops, we fuse them all. However, if loop reordering is possible, and we choose only certain loops to be fused, then there are 2^n possible schedules to start with, given that there are n number of fusible loops (i.e. n number of common iterators in the two kernels). This is an upper bound without considering any constraints of sparse access patterns. Reordering of inner loops can be performed after fusion, and other scheduling directives (split, parallelization, etc.) can be applied separately, giving more scheduling opportunities with imperfect nesting. This scheduling space is obviously very large, so smart strategies for searching that space is a promising avenue for future work.

4.4 Code Generation

We carefully redesigned intermediate representation (IR) in TACO to support the branched iteration graph and manage temporaries such that code generation backend does not require any changes. We rewrite the graph loop structure with

where statements defining a producer-consumer relationship. This placement of temporaries for the producer-consumer relationship and the change of iteration graph explained in Section 4.2 preserves all the attributes that are necessary for TACO code generation backend.

In TACO each index in the iteration graph is converted to one or more loops to iterate through dense loops or co-iterate over the levels of sparse data formats. An iteration lattice [17] is used to co-iterate through the intersections of the sparse dimensions which results in a single for-loop, single while-loop or multiple while-loops.

5 Implementation

We implement the branch iteration graph transformation described in Section 4 on top of the TACO [17] intermediate representation (IR). Furthermore, we introduce a new scheduling directive to separate it from the algorithmic language and to provide the scheduling language with more opportunities to generate more (performant) schedules.

We change the iteration graph [17] and use the concrete index notation [16] to introduce intermediate temporaries that are shared between the producer and the consumer. We implement a *nest boundary* between the fused loops and shared index loops to constrain performing loop reordering transformations between them. In our running example, the user cannot interchange loops with an outer level, once the distribution operation is performed.

This new transformation can be used in the context of tensor multiplication. Hence, it does not generalize to tensor expressions with tensor additions. We limit the number of tensors and index variables removed from the index expression, to identify the producer and consumer graphs, per iteration to one. We believe that the algorithm could be generalized to support fusion of indices shared between multiple tensors which would be able to support high order tensors and complex tensor contractions.

6 Evaluation

We compare SparseLNR to two other techniques:

TACO Original. Given a large combined index expression containing multiple smaller index expressions, the code generated by TACO has a perfectly nested loop structure with at least one loop per each index variable in the index expression. We refer to this version as *TACO Original*.

TACO Separate. In some cases, the asymptotic complexity of *TACO Original* can be reduced by manually separating a larger index expression into multiple smaller index expressions by using temporary tensors to store the intermediate results. We refer to this version as *TACO Separate*. When there are multiple ways to break down the computation into smaller kernels, we evaluate all those combinations and report the best execution time.

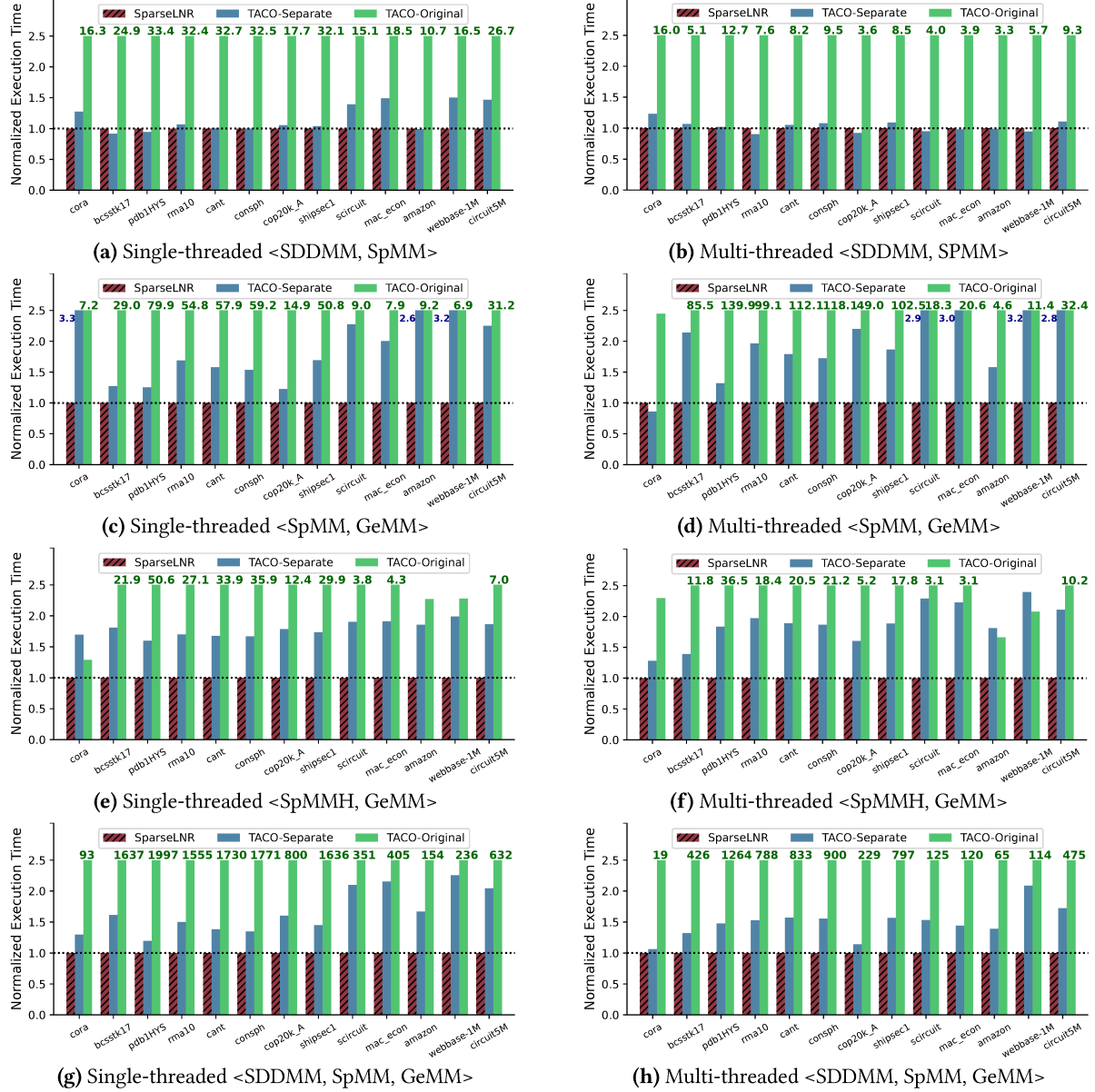


Figure 5. Performance Comparison with TACO for benchmarks with 2-D matrices.

6.1 Experimental Setup

All experiments run on a single socket 64-Core AMD Ryzen Threadripper 3990X at 2.2 GHz, with 32KB L1 data cache, 512KB shared L2 cache, and 16MB shared L3 cache. We compile the code using GCC 7.5.0 with `-O3 -ffast-math`. We use `-fopenmp` for parallel versions with *OpenMP version 4.5*. All parallel versions use 64 threads which is the number of physical cores available in the machine.

Datasets. We use sparse tensors from four sources: SuiteSparse [12]; Network Repository [28]; Formidable Repository of Open Sparse Tensors and Tools (FROSTT) [31]; and the 1998 DARPA Intrusion Detection Evaluation [15]. Dense tensors in kernels are randomly generated. Table 1 gives the details of the sparse tensors.

6.2 Benchmarks

<SDDMM, SpMM>. SDDMM computation followed by the SpMM operation, $A_{il} = \sum B_{ij} \cdot C_{ik} \cdot D_{jk} \cdot E_{jl}$. This operation is used in graph neural networks [27]. We set the inner dimensions k and l to `<64, 64>`. Fusion of SDDMM with SpMM results in a scalar intermediate to share the results between the fused loops as shown in Figure 2d.

<SpMMH, GeMM>. SpMMH here is pre-multiplying the Hadamard product of two dense matrices by a sparse matrix. The combined kernel we evaluate is $A_{il} = \sum B_{ik} \cdot C_{kj} \cdot D_{kj} \cdot E_{jl}$ with inner dimensions j and l set to `<128, 128>`.

<SpMM, GeMM>. SpMM kernel followed by another GeMM kernel, $A_{il} = \sum B_{ij} \cdot C_{jk} \cdot D_{kl}$. The inner dimensions k, m

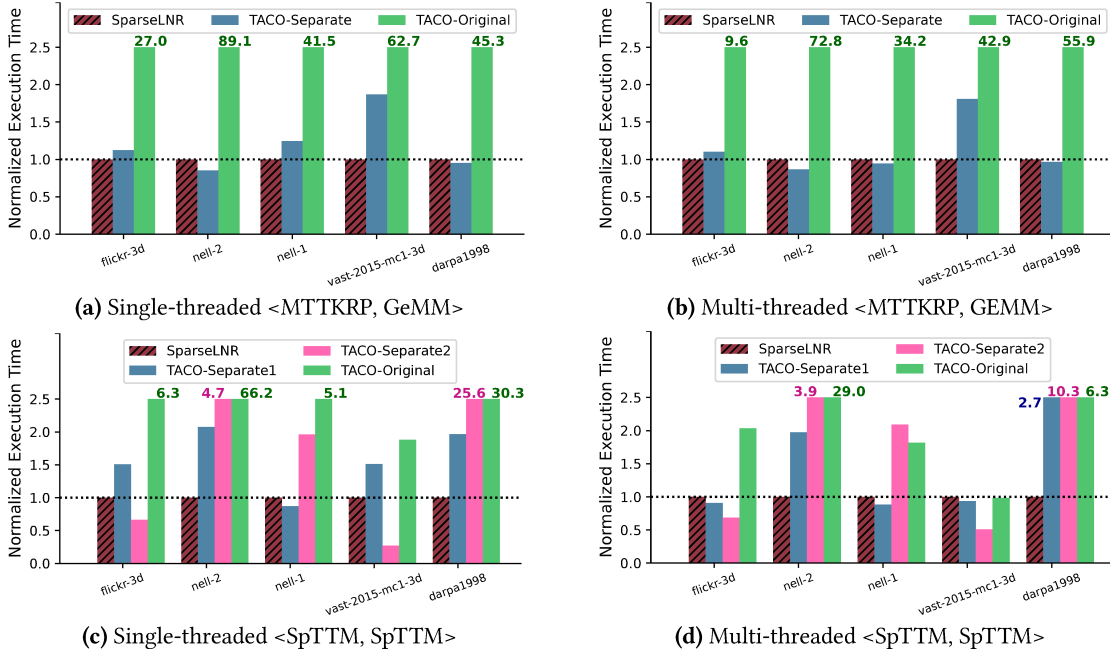


Figure 6. Performance Comparison with TACO for benchmarks with 3-D tensors.

Table 1. Test tensors used in the evaluation from various matrix and tensor collections mentioned in the Section 6.1

Tensor	Dimensions	Non-zeros
cora	$2.7K \times 2.7K$	5.4K
bcsstk17	$11K \times 11K$	429K
pdf1HYS	$36K \times 36K$	4.34M
rma10	$47K \times 47K$	2.37M
cant	$62K \times 62K$	4.01M
consph	$83K \times 83K$	6.01M
cop20k_A	$12K \times 12K$	2.62M
shipsec1	$140K \times 140K$	7.81M
scircuit	$171K \times 171K$	959K
mac_econ_fwd500	$207K \times 207K$	1.27M
amazon	$334K \times 334K$	1.85M
webbase-1M	$1.00M \times 1.00M$	3.11M
circuit5M	$5.56M \times 5.56M$	59.52M
flickr-3d	$320K \times 2.82M \times 1.60M$	112.89M
nell-2	$12K \times 9K \times 288K$	76.88M
nell-1	$2.9M \times 2.14M \times 25.5M$	143.6M
vast-2015-mc1-3d	$165K \times 11K \times 2$	26.02M
darpa1998	$22K \times 22K \times 23.7M$	28.42M

are set to $\langle 128, 64 \rangle$. This calculation is commonly used for updating the hidden state in GNNs [37].

<SDDMM, SpMM, GEMM>. We combine two of the prior kernels to show the recursive applicability of the algorithm. $A_{im} = \sum B_{ij} \cdot C_{ik} \cdot D_{jk} \cdot F_{jl} \cdot W_{lm}$. We could relate this execution to performing SDDMM operation to get the attention values along the edges of a graph, multiplying the feature matrix of the graph with a weight matrix to get the new feature set of the graph and then doing a neighbor sum of the graph. The inner dimensions k, l and m are set to $\langle 64, 64, 64 \rangle$.

<MTTKRP, GEMM>. Khatri-Rao product (MTTKRP) followed by GEMM operation, $A_{im} = \sum B_{ikl} \cdot C_{lj} \cdot D_{kj} \cdot E_{jm}$. MTTKRP kernel is used in various sparse computations like signal processing and computer vision [7]. We set the inner dimensions j and m to $\langle 32, 64 \rangle$.

<SpTTM, SpTTM>. Sparse Tensor Times Matrix (SpTTM) operation followed by another SpTTM operation, $A_{ijm} = \sum B_{ijk} \cdot C_{kl} \cdot D_{lm}$. SpTTM is a computational kernel used in data analytics and data mining applications such as the popular Tucker decomposition [25]. The inner dimensions l and m are set to $\langle 32, 64 \rangle$.

Sparse Formats. SpMM, SDDMM kernels use standard compressed sparse row (CSR) format for their sparse matrices whereas SpTTM, MTTKRP kernels use compressed sparse fiber (CSF) format.

For the SDDMM, SpMM, MTTKRP kernels in *TACO separate* we use the versions provided in Senanayake *et al.* [30]. For the rest of kernels we evaluate multiple schedules and select the best performing one. TACO does not generate multi-threaded code when the output tensor is sparse. Prior work has evaluated against single-threaded code in such situations [16, 34]. Following the strategy of Senanayake *et al.* [30], we modified the TACO generated code manually to add multithreading.

In general, the speedups we see compared to the TACO original comes from the reduction in asymptotic complexity while the speedups we see compared to the TACO separate comes from the reduction in cache reuse distances by removing large tensors used to store intermediate results.

We see speedups for our approach of 0.90–1.23x compared TACO Separate and 3.31–16.05 compared to TACO Original in <SDDMM, SpMM> kernel’s multi-threaded execution. For single-threaded execution, we get 0.91–1.50x compared to TACO separate and 10.75–33.39x compared to TACO original. In multithreaded execution the fused kernel performs better on circuit5M, shipsec1, consph, pdf1HYS, and cant,

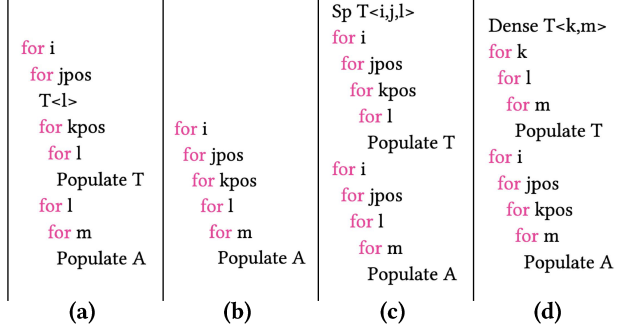


Figure 7. Basic loop structure of different schedules for <SpTTM, SpTTM> kernel.

the matrices with most non-zeros, from the tested datasets 1 compared against the separate execution.

We observe speedups of 1.60–1.99x against TACO separate and 1.29–50.55x against TACO original in single-threaded execution for <SpMMH, GEMM> kernel. For the same kernel in multi-threaded execution, we observed speedups of 1.28–2.40x against TACO separate and 1.66–36.50x against TACO original. For the <SpMM, GEMM> kernel in single-threaded execution, speedups of 1.23–3.27x, and 6.91–79.86x are observed for TACO original and TACO separate, respectively. For the same kernel in multi-threaded execution, we observed speedups of 0.86–3.16x, and 2.44–139x for TACO original and TACO separate, respectively.

We see substantial speedups in <SDDMM, SpMM, GEMM> due to the kernel presenting two opportunities for fusion: <SDDMM, SpMM> and <SpMM, GEMM>. We see speedups ranging from 1.20–2.26x for single-threaded execution against TACO separate, 93–1997x over single-threaded TACO original, 1.06–2.09x for multithreaded TACO separate and 19–1263x for multithreaded TACO original.

We see that our approach under-performed in datasets nell-2 and darpa1998 against the TACO separate for the <MT-KRP, GEMM> benchmark. In these datasets, the first dimensions of the tensors are bounded by 12092, and 22476. Therefore, the intermediate matrix in the TACO separate execution has sizes that fit within the last level cache. Furthermore, executing kernels separately sometimes offer more opportunities to optimize smaller kernels individually. Due to these reasons, there may be datasets that the separate kernel execution performs better than the fused version. But we see considerable speedups versus TACO Separate when the intermediate tensors are large.

For <SpTTM, SpTTM>, there are two different schedules due to different associativity choices. We see varying results based on which choice is made, so we report both of them as *separate1* and *separate2* (see Figures 6c and 6d). Figure 7 shows the basic loop structure for different versions of <SpTTM, SpTTM>. The fused version has asymptotic complexity of $O(nnz(B_{IJK})L + nnz(B_{IJL})LM)$. The TACO original version $A_{ijm} = \sum B_{ijk} \cdot C_{kl} \cdot D_{lm}$ has asymptotic complexity

of $O(nnz(B_{IJK})LM)$. But *separate1* ($T_{ijl} = \sum B_{ijk} \cdot C_{kl}$ followed by $A_{ijl} = \sum T_{ijl} \cdot C_{kl}$) and *separate2* ($T_{km} = \sum C_{kl} \cdot D_{lm}$ followed by $A_{ijm} = \sum B_{ijk} \cdot T_{km}$) versions have complexities of $O(nnz(B_{IJK})L + nnz(B_{IJL})M)$ and $O(KLM + nnz(B_{IJK})M)$, respectively.

When k is small (for instance, dimension k of the dataset *vast2015-mc1-3d* in Table 1 is bounded by 2), the asymptotic complexity of our approach is comparable to TACO’s baseline approach, and the size of the working set is small. Therefore, there are no cache misses for the TACO separate approach; the re-associated schedule hence has the best performance. In the other datasets, where k is large (for instance, the dimension k of the dataset *darpa1998* in Table 1 is bounded by 28.42M), these effects vanish, and our approach is considerably faster than any competing version. We note that SparseLNR’s representation could support re-association scheduling directives, that would allow it to use the better schedules of TACO separate, but we leave that for future work.

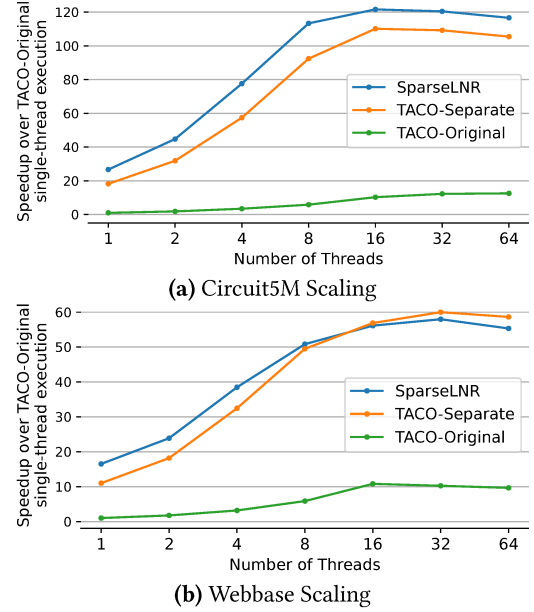
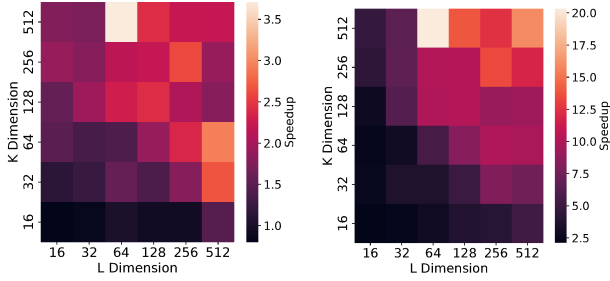


Figure 8. Speedup change with respect to the number of threads for <SDDMM, SpMM> benchmark.

The scaling results for webbase and circuit5M datasets, on <SDDMM, SpMM> benchmark are shown in Figure 8; SparseLNR delivers comparable scaling to the best alternative approach. We observed similar scaling for other benchmarks and datasets.

6.3 Case Study: Performance with different inner dimensions

We chose the inner dimensions of the benchmarks explained in Section 6.2 arbitrarily. In this case study we consider change of performance wrt. varying inner dimension sizes (k and l) in the benchmark <SpMM, GEMM> ($A_{il} = \sum B_{ijk} \cdot C_{jk} \cdot D_{kl}$) in Section 6.2. Here, the dimensions i and j are



(a) TACO-Sep. vs. SparseLNR (b) TACO-Orig. vs. SparseLNR

Figure 9. Speedup variation wrt. inner dimension sizes (k and l) on the benchmark $\langle \text{SpMM}, \text{GEMM} \rangle$. The Figures (a) and (b) correspond to multithreaded execution of SparseLNR against TACO-Separate and TACO-Original, respectively.

determined by size of the graphs read into the sparse matrix B_{ij} and cannot be arbitrarily changed. However, the dimensions k and l can change in size since the matrices C_{jk} and D_{kl} are dense. Usually, in GNN literature, these dimensions correspond to feature sizes in hidden layers.

Performing the *loop fusion after distribution* in the benchmark $\langle \text{SpMM}, \text{GEMM} \rangle$ results in a temporary vector of the length of the size of dimension k , and k and l dimensions completely define the size of the matrix D_{kl} . When the size of the dimension l is small, D_{kl} can completely fit in higher level caches for the k values considered. Therefore, for smaller l values, the speedup increases with the size of k (See Figure 9). Because with increasing k , the temporary tensor gets larger and it decreases the performance of TACO-Separate. We see this behavior for l dimension sizes of 16–128 in Figure 9a. But with increasing k , when the size of the l dimension gets larger, the sizes of D_{kl} and temporary vector get larger. As a result, they keep getting evicted from the higher level caches in SparseLNR. Therefore, as shown in Figure 9a, the peak performance in columns 256 and 512 of the dimension l occurs not when k is 512, but when size of k takes values in the range 64–256. Increasing sizes of k and l results in higher time complexity for TACO-Original. Hence, speedup of SparseLNR increases with the sizes of k and l in the Figure 9b.

7 Related Work

Code generation for tensor algebra has been extensively researched. This area of research can be primarily divided into two subareas — sparse and dense. First, we discuss the related work on code generation and optimization techniques for sparse tensor algebra and then move to the ones for dense.

7.1 Sparse Tensor Algebra

Automated sparse code generation [5, 6, 17, 34] is a heavily-researched topic. Even though these methods are highly effective, they lack fine-grained optimizations like ours that applies across kernels to reduce the time-complexity of the

computation. Ahrens *et al.* [2] proposed splitting large tensor expressions into smaller kernels to minimize the time-complexity. The Sparse Polyhedral Framework [21, 32, 33] employs an inspector-executor strategy to transform the data layout and schedule of sparse computations to achieve locality and parallelism. Kurt *et al.* [20] improved SpMM and SDDMM kernels by optimizing their tile sizes using a sparsity signature. However, these methods do not consider loop nest restructuring transformations to improve data locality across kernels.

Athena [23], Sparta [24], and HiParTi [22] are techniques which provide highly optimized kernels for sparse tensor operations and contraction sequences that shows significant performance improvements. Kernel fusion has been used in FusedMM [27] to accelerate SDDMM and SpMM used in graph neural network applications. Their transformation is structurally analogous to SparseLNR, but is specific to graph embeddings, and further performs kernel-specific optimizations. None of these prior techniques handle arbitrary sparse tensor expressions supporting a variety of input formats.

7.2 Dense Tensor Algebra

Optimizations for computations over dense tensors have been well-studied for decades. Numerous loop optimizations for dense tensor contractions [3, 4, 9–11, 19, 29] for CPUs and tensor contractions for GPUs [1, 26] have been proposed that exhibits superior performance. However, these transformations are not directly applicable to sparse tensor algebra since there data access restrictions for sparse tensors and the non-affine nature of loop nests.

8 Conclusion

We presented SparseLNR, a loop restructuring framework for sparse tensor algebra programs. SparseLNR improves the performance of sparse computations by reducing time complexity and enhancing data locality. SparseLNR enables kernel distribution and loop fusion and achieves significant performance improvements for real-world benchmarks. The new scheduling transformations introduced by SparseLNR expands the scheduling space of sparse tensor applications and facilitates fine-grained tuning.

Acknowledgments

We appreciate the feedback from the anonymous reviewers for their suggestions and comments that helped to improve this paper. We would also like to thank Fredrik Kjolstad for the valuable discussions we had regarding the SparseLNR transformation. This work was supported in part by the National Science Foundation awards CCF-1908504 and CCF-1919197. Any opinions, findings, and conclusions or recommendations expressed in this paper are those of the authors and do not necessarily reflect the views of the National Science Foundation.

References

- [1] A. Abdelfattah, M. Baboulin, V. Dobrev, J. Dongarra, C. Earl, J. Falcou, A. Haidar, I. Karlin, Tz. Kolev, I. Masliah, and S. Tomov. 2016. High-performance Tensor Contractions for GPUs. *Procedia Computer Science* 80 (2016), 108–118. <https://doi.org/10.1016/j.procs.2016.05.302>
- [2] Peter Ahrens, Fredrik Kjolstad, and Saman Amarasinghe. 2021. *An Asymptotic Cost Model for Autoscheduling Sparse Tensor Programs*. arXiv:2111.14947
- [3] A. Allam, J. Ramanujam, G. Baumgartner, and P. Sadayappan. 2006. Memory minimization for tensor contractions using integer linear programming. In *Proceedings 20th IEEE International Parallel Distributed Processing Symposium*. 8 pp.–. <https://doi.org/10.1109/IPDPS.2006.1639717>
- [4] Alina Bibireata, Sandhya Krishnan, Gerald Baumgartner, Daniel Cociorva, Chi-chung Lam, P. Sadayappan, J. Ramanujam, David E. Bernholdt, and Venketesh Choppella. 2004. Memory-Constrained Data Locality Optimization. In *16th International Workshop on Languages and Compilers for Parallel Computing (LCPC'03)*. Springer-Verlag, 93–108.
- [5] Aart J. C. Bik and Harry A. G. Wijshoff. 1993. Compilation Techniques for Sparse Matrix Computations. In *Proceedings of the 7th International Conference on Supercomputing (Tokyo, Japan) (ICS '93)*. Association for Computing Machinery, New York, NY, USA, 416–424. <https://doi.org/10.1145/165939.166023>
- [6] Aart J. C. Bik and Harry A. G. Wijshoff. 1993. On Automatic Data Structure Selection and Code Generation for Sparse Computations. In *Proceedings of the 6th International Workshop on Languages and Compilers for Parallel Computing*. Springer-Verlag, Berlin, Heidelberg, 57–75.
- [7] Jee Choi, Xing Liu, Shaden Smith, and Tyler Simon. 2018. Blocking Optimization Techniques for Sparse Tensor Computation. In *2018 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*. 568–577. <https://doi.org/10.1109/IPDPS.2018.00066>
- [8] Stephen Chou, Fredrik Kjolstad, and Saman Amarasinghe. 2018. Format Abstraction for Sparse Tensor Algebra Compilers. *Proc. ACM Program. Lang.* 2, OOPSLA, Article 123 (oct 2018), 30 pages. <https://doi.org/10.1145/3276493>
- [9] Daniel Cociorva, Gerald Baumgartner, Chi-Chung Lam, P. Sadayappan, J. Ramanujam, Marcel Nooijen, David E. Bernholdt, and Robert Harrison. 2002. Space-Time Trade-off Optimization for a Class of Electronic Structure Calculations. In *Proceedings of the ACM SIGPLAN 2002 Conference on Programming Language Design and Implementation (Berlin, Germany) (PLDI '02)*. Association for Computing Machinery, New York, NY, USA, 177–186. <https://doi.org/10.1145/512529.512551>
- [10] D. Cociorva, Xiaoyang Gao, S. Krishnan, G. Baumgartner, Chi-Chung Lam, P. Sadayappan, and J. Ramanujam. 2003. Global communication optimization for tensor contraction expressions under memory constraints. In *Proceedings International Parallel and Distributed Processing Symposium*. 8 pp.–. <https://doi.org/10.1109/IPDPS.2003.1213121>
- [11] D. Cociorva, J. W. Wilkins, C. Lam, G. Baumgartner, J. Ramanujam, and P. Sadayappan. 2001. Loop Optimization for a Class of Memory-Constrained Computations. In *Proceedings of the 15th International Conference on Supercomputing (Sorrento, Italy) (ICS '01)*. Association for Computing Machinery, New York, NY, USA, 103–113. <https://doi.org/10.1145/377792.377814>
- [12] Timothy A. Davis and Yifan Hu. 2011. The University of Florida Sparse Matrix Collection. *ACM Trans. Math. Softw.* 38, 1, Article 1 (dec 2011), 25 pages. <https://doi.org/10.1145/2049662.2049663>
- [13] William L. Hamilton, Rex Ying, and Jure Leskovec. 2018. Inductive Representation Learning on Large Graphs. arXiv:1706.02216 [cs.SI]
- [14] Yuwei Hu, Zihao Ye, Minjie Wang, Jiali Yu, Da Zheng, Mu Li, Zheng Zhang, Zhiru Zhang, and Yida Wang. 2020. FeatGraph: A Flexible and Efficient Backend for Graph Neural Network Systems. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis (Atlanta, Georgia) (SC '20)*. IEEE Press, Article 71, 13 pages.
- [15] Inah Jeon, Evangelos E. Papalexakis, U Kang, and Christos Faloutsos. 2015. HaTen2: Billion-scale tensor decompositions. In *2015 IEEE 31st International Conference on Data Engineering*. 1047–1058. <https://doi.org/10.1109/ICDE.2015.7113355>
- [16] Fredrik Kjolstad, Peter Ahrens, Shoaib Kamil, and Saman Amarasinghe. 2019. Tensor Algebra Compilation with Workspaces. (2019), 180–192. <http://dl.acm.org/citation.cfm?id=3314872.3314894>
- [17] Fredrik Kjolstad, Shoaib Kamil, Stephen Chou, David Lugato, and Saman Amarasinghe. 2017. The Tensor Algebra Compiler. *Proc. ACM Program. Lang.* 1, OOPSLA, Article 77 (Oct. 2017), 29 pages. <https://doi.org/10.1145/3133901>
- [18] Vladimir Kotlyar, Keshav Pingali, and Paul Stodghill. 1997. A Relational Approach to the Compilation of Sparse Matrix Programs. In *Proceedings of the Third International Euro-Par Conference on Parallel Processing (Euro-Par '97)*. Springer-Verlag, Berlin, Heidelberg, 318–327.
- [19] Sandhya Krishnan, Sriram Krishnamoorthy, Gerald Baumgartner, Daniel Cociorva, Chi-Chung Lam, P. Sadayappan, J. Ramanujam, David E. Bernholdt, and Venkatesh Choppella. 2003. Data Locality Optimization for Synthesis of Efficient Out-of-Core Algorithms. In *HiPC*.
- [20] Süreyya Emre Kurt, Aravind Sukumaran-Rajam, Fabrice Rastello, and P. Sadayappan. 2020. Efficient Tiled Sparse Matrix Multiplication through Matrix Signatures. In *SC20: International Conference for High Performance Computing, Networking, Storage and Analysis*. 1–14. <https://doi.org/10.1109/SC41405.2020.00091>
- [21] Alan LaMille and Michelle Mills Strout. 2010. Enabling Code Generation within the Sparse Polyhedral Framework.
- [22] Jiajia Li. [n.d.]. A Hierarchical Parallel Tensor Infrastructure (HiParTI). <https://github.com/pnrl/HiParTI> [Accessed 02-Feb-2022].
- [23] Jiawen Liu, Dong Li, Roberto Gioiosa, and Jiajia Li. 2021. Athena: High-Performance Sparse Tensor Contraction Sequence on Heterogeneous Memory. In *Proceedings of the ACM International Conference on Supercomputing (Virtual Event, USA) (ICS '21)*. Association for Computing Machinery, New York, NY, USA, 190–202. <https://doi.org/10.1145/3447818.3460355>
- [24] Jiawen Liu, Jie Ren, Roberto Gioiosa, Dong Li, and Jiajia Li. 2021. *Sparta: High-Performance, Element-Wise Sparse Tensor Contraction on Heterogeneous Memory*. Association for Computing Machinery, New York, NY, USA, 318–333. <https://doi.org/10.1145/3437801.3441581>
- [25] Yuchen Ma, Jiajia Li, Xiaolong Wu, Chenggang Yan, Jimeng Sun, and Richard Vuduc. 2019. Optimizing sparse tensor times matrix on GPUs. *J. Parallel and Distrib. Comput.* 129 (2019), 99–109. <https://doi.org/10.1016/j.jpdc.2018.07.018>
- [26] Thomas Nelson, Axel Rivera, Prasanna Balaprakash, Mary Hall, Paul D. Hovland, Elizabeth Jessup, and Boyana Norris. 2015. Generating Efficient Tensor Contractions for GPUs. In *2015 44th International Conference on Parallel Processing*. 969–978. <https://doi.org/10.1109/ICPP.2015.106>
- [27] Md. Khaledur Rahman, Majedul Haque Sujon, and Ariful Azad. 2020. *FusedMM: A Unified SDDMM-SpMM Kernel for Graph Embedding and Graph Neural Networks*. arXiv:2011.06391
- [28] Ryan A. Rossi and Nesreen K. Ahmed. 2016. An Interactive Data Repository with Visual Analytics. *SIGKDD Explor.* 17, 2 (2016), 37–41. <http://networkrepository.com>
- [29] S.K. Sahoo, S. Krishnamoorthy, R. Panuganti, and P. Sadayappan. 2005. Integrated Loop Optimizations for Data Locality Enhancement of Tensor Contraction Expressions. In *SC '05: Proceedings of the 2005 ACM/IEEE Conference on Supercomputing*. 13–13. <https://doi.org/10.1109/SC.2005.35>

- [30] Ryan Senanayake, Changwan Hong, Ziheng Wang, Amalee Wilson, Stephen Chou, Shoaib Kamil, Saman Amarasinghe, and Fredrik Kjolstad. 2020. A Sparse Iteration Space Transformation Framework for Sparse Tensor Algebra. *Proc. ACM Program. Lang.* 4, OOPSLA, Article 158 (Nov. 2020), 30 pages. <https://doi.org/10.1145/3428226>
- [31] Shaden Smith, Jee W. Choi, Jiajia Li, Richard Vuduc, Jongsoo Park, Xing Liu, and George Karypis. 2017. *FROSTT: The Formidable Repository of Open Sparse Tensors and Tools*. New York, NY, USA. <https://doi.org/10.1145/2049662.2049663>
- [32] Michelle Mills Strout, Mary Hall, and Catherine Olschanowsky. 2018. The Sparse Polyhedral Framework: Composing Compiler-Generated Inspector-Executor Code. *Proc. IEEE* 106, 11 (2018), 1921–1934. <https://doi.org/10.1109/JPROC.2018.2857721>
- [33] Michelle Mills Strout, Alan LaMille, Larry Carter, Jeanne Ferrante, Barbara Kreaseck, and Catherine Olschanowsky. 2016. An Approach for Code Generation in the Sparse Polyhedral Framework. *Parallel Comput.* 53, C (apr 2016), 32–57. <https://doi.org/10.1016/j.parc.2016.02.004>
- [34] Ruiqin Tian, Luanzheng Guo, Jiajia Li, Bin Ren, and Gokcen Kestor. 2021. A High Performance Sparse Tensor Algebra Compiler in MLIR. In *2021 IEEE/ACM 7th Workshop on the LLVM Compiler Infrastructure in HPC (LLVM-HPC)*. 27–38. <https://doi.org/10.1109/LLVMHPC54804.2021.00009>
- [35] Anand Venkat, Mary Hall, and Michelle Strout. 2015. Loop and Data Transformations for Sparse Matrix Code. In *Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation (Portland, OR, USA) (PLDI '15)*. Association for Computing Machinery, New York, NY, USA, 521–532. <https://doi.org/10.1145/2737924.2738003>
- [36] Mengjia Xu. 2020. Understanding graph embedding methods and their applications. *CoRR* abs/2012.08019 (2020). arXiv:2012.08019 <https://arxiv.org/abs/2012.08019>
- [37] Jie Zhou, Ganqu Cui, Shengding Hu, Zhengyan Zhang, Cheng Yang, Zhiyuan Liu, Lifeng Wang, Changcheng Li, and Maosong Sun. 2021. Graph Neural Networks: A Review of Methods and Applications. arXiv:1812.08434 [cs.LG]