

Tucker-SVD-FFT-Accelerated Solution of Volume Integral Equation for Magneto-Quasi-Static Analysis of Voxelized Geometries

Xiaofan Jia

School of Electrical and Electronic
Engineering
Nanyang Technological University
Singapore
xiaofan002@e.ntu.edu.sg

Yang Liu

Applied Mathematics & Computational
Research Division
Lawrence Berkeley National Laboratory
Berkeley, USA
liuyangzhuan@lbl.gov

Mingyu Wang

Xpeedic Co., LTD.
Shanghai, China
mingyu.wang@xpeedic.com

Chao-Fu Wang

National University of Singapore
Singapore
cfwang@nus.edu.sg

Abdulkadir C. Yucel

School of Electrical and Electronic
Engineering
Nanyang Technological University
Singapore
acyucel@ntu.edu.sg

Abstract—A Tucker decomposition-singular value decomposition (SVD)-fast Fourier transform (FFT)-based acceleration scheme, called Tucker-SVD-FFT, is proposed for rapidly solving volume integral equation (VIE) for magneto-quasi-static (MQS) analysis of voxelized geometries. The proposed acceleration scheme uses Tucker, SVD, and FFT acceleration techniques to perform matrix-vector multiplications related to far, nearby, and self interactions of the voxel groups, respectively. The resulting Tucker-SVD-FFT-accelerated VIE solver is applicable to multiscale and non-uniformly discretized voxelized geometries, unlike its pure FFT-accelerated counterpart. Preliminary results demonstrate that the Tucker compression yields at least *one order* of magnitude better compression compared to SVD compression, while the proposed acceleration scheme outperforms the FFT and SVD acceleration schemes for MQS analysis of a uniformly voxelized and well-separated structure. While the scheme's application to the MQS analysis is demonstrated here, its adoption to the full-wave analysis is underway.

Keywords—Magneto-quasi-static (MQS), matrix-vector multiplications, tensor decompositions, Tucker decomposition, volume integral equation (VIE)

I. INTRODUCTION

Volume integral equation (VIE) solvers have been widely used for magneto-quasi-static (MQS) analysis for applications in low-frequency bioelectromagnetic analysis, chip design, and geophysical exploration. To accelerate these solvers, matrix-vector multiplications (MVMs) arising during the iterative solution of the VIEs are often accelerated by fast Fourier transforms (FFTs) for voxelized geometries [1-4]. Nonetheless, the computational time and memory requirements of FFT-accelerated solvers scale with the total number of voxels in the computational domain. To this end, these solvers lose their efficiency when applied to scenarios with well-separated structures. Furthermore, the FFT method's uniform grid requirement restricts its application to uniform voxel-based discretization and makes its application infeasible to multiscale voxelized structures (i.e., structures discretized by voxels with different sizes). For performing fast analysis of multiscale and/or well-separated voxelized structures, matrix decomposition technique [5-7] can be used to accelerate MVMs. However, these techniques fail to effectively exploit the low-rank nature of tensors naturally arising due to the voxelized setting. Recently, tensor

decompositions have proven their effectiveness in compressing low-rank tensors used in various integral equation solvers. Especially, Tucker decomposition significantly reduced the memory requirement and setup time of static and quasi-static solvers [2-4] and alleviated computational time and/or memory requirements of stages in fast multipole method-FFT (FMM-FFT) accelerated full-wave solvers [8, 9]. However, to date, no study has been performed to use the Tucker decompositions for compressing the low-rank blocks in the system matrices and performing fast MVMs in Tucker-compressed format.

This study proposes a Tucker-SVD-FFT accelerated VIE solver for fast MQS analysis of the voxelized structures. In the proposed scheme, the computational domain bounding the structure is divided into the boxes/groups, enclosing the voxels. During the iterative solution of VIE, FFT and SVD are used to accelerate the MVMs involving the self-interactions of voxel groups and nearby interactions between (neighbouring) groups, respectively. Furthermore, Tucker decomposition is used to accelerate the MVMs involving the interactions between far groups. To do that, Tucker decomposition compresses low-rank six-dimensional (6-D) tensors. Compressed tensors are then used to perform fast tensor-vector multiplications. Numerical results show that Tucker decomposition introduces at least one order of magnitude reduction in the memory required to store the far interactions between voxel groups. Moreover, the Tucker-SVD-FFT-accelerated VIE solver requires 2.13x less memory and 3.38x less computational time compared to the FFT acceleration technique for MQS analysis of a well-separated voxelized structure. These savings, compared to FFT technique, increase linearly with increasing separation between structures.

II. FORMULATION

This section briefly explains the Tucker-SVD-FFT acceleration scheme. Let conductor(s) with conductivity σ reside in free space and the structure be discretized by voxels with edge length Δx . Furthermore, let the computational domain split into small boxes/groups enclosing voxels. During the iterative solution of the VIE, MVMs involving the groups' self-interactions are expedited by FFT [1]. The blocks storing the interactions between basis functions in adjacent groups are compressed via SVD and their corresponding MVMs are

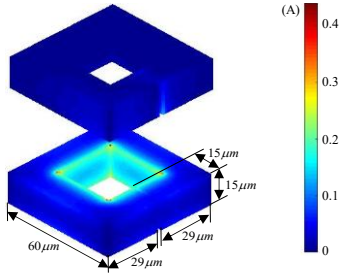


Fig. 1. Two parallel identical square coils with $R_x = R_y = R_z = 30 \mu\text{m}$.

performed in SVD-compressed format [5]. For well-separated far groups, the interactions of all basis functions are computed and stored in 6-D tensors. Given a pair of far groups, the number of voxels along x -, y -, and z -directions are N_1 , N_2 , and N_3 for the source group, and N_4 , N_5 , and N_6 for the observation group. The interactions between these groups are stored in a tensor $\mathcal{S}$ with dimensions $N_1 \times \dots \times N_6$ and compressed by Tucker decomposition as

$$\mathcal{S} = \mathcal{C}_T \times_1 \bar{\mathbf{U}}_T^1 \times_2 \bar{\mathbf{U}}_T^2 \times_3 \bar{\mathbf{U}}_T^3 \times_4 \bar{\mathbf{U}}_T^4 \times_5 \bar{\mathbf{U}}_T^5 \times_6 \bar{\mathbf{U}}_T^6. \quad (1)$$

Here, $\mathcal{C}_T$ is the 6-D core tensor with dimensions $r_1 \times \dots \times r_6$, $\bar{\mathbf{U}}_T^i$ denotes a factor matrix with dimensions $N_i \times r_i$, and $\times_i$ represents the mode- i product, $i = 1, \dots, 6$. As N_i is much larger than the multilinear rank r_i , the Tucker-compressed representation of $\mathcal{S}$ requires much less memory than $\mathcal{S}$. While iteratively solving the VIE, the product between $\mathcal{S}$ and current coefficient vector is performed in Tucker-compressed format. Procedures outlining the fast tensor-vector multiplication and efficient computation of Tucker decomposition will be explained during the presentation.

III. NUMERICAL RESULTS

The efficiency of Tucker-SVD-FFT-accelerated VIE solver is examined for inductance extraction of two parallel identical square coils with $\sigma = 5.8 \times 10^7 \text{ S/m}$ at 300 MHz [Fig. 1]. The structure is discretized by voxels with $\Delta x = 1 \mu\text{m}$, grouped into boxes with size of $15 \times 15 \times 15$. The tolerance for Tucker decomposition and SVD is set to 10^{-6} . The distances between the centers of two coils along the x -, y -, and z -directions are identical, i.e., $R_x = R_y = R_z$, and increased from $30 \mu\text{m}$ to $240 \mu\text{m}$. For inductance extraction of this scenario with increasing R_x , the proposed Tucker-SVD-FFT scheme is compared with FFT and SVD schemes. FFT-acceleration is the original routine in VoxHenry [1],

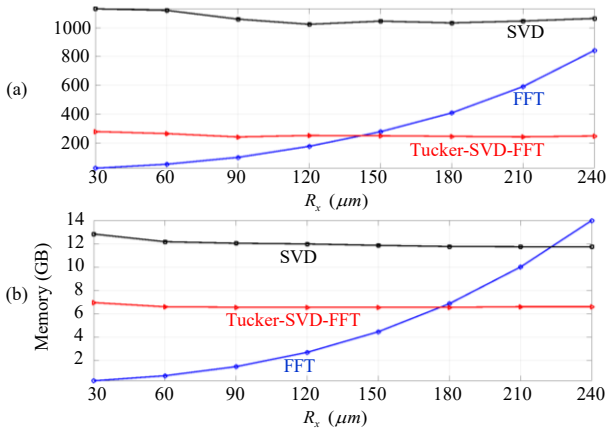


Fig. 2. (a) CPU time and (b) memory requirements of FFT, SVD, and Tucker-SVD-FFT schemes w.r.t. the distance between coil centers.

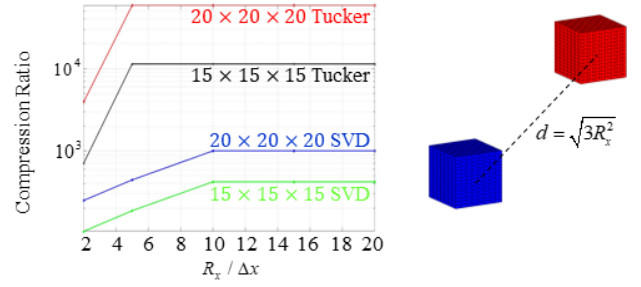


Fig. 3. Compression ratio obtained via SVD and Tucker decomposition for far interactions when the group sizes are $15 \times 15 \times 15$ and $20 \times 20 \times 20$.

while SVD scheme leverages the conventional MVM and SVD-accelerated MVM schemes for self-interactions of groups and nearby/far interactions of groups, respectively. Extracted parameters using all acceleration schemes agree well with each other and will be shared during the presentation. Fig. 2 illustrates that the Tucker-SVD-FFT scheme yields significant CPU time and memory savings. For the analysis at $R_x = 240 \mu\text{m}$, CPU time/memory requirements of FFT, SVD, and Tucker-SVD-FFT schemes for iterative solutions are 842.2 s/14.0 GB, 1064.2 s/11.7 GB, and 249.2 s/6.59 GB, respectively. To this end, the proposed scheme requires 3.38x and 4.27x less CPU time during iterative solution and 2.13x and 1.78x less memory compared to FFT and SVD. The memory savings mainly result from the high compression ratio achieved by Tucker decomposition for far interactions, which reaches over one order of magnitude larger than that obtained by SVD, as shown in Fig. 3.

IV. CONCLUSION

A Tucker-SVD-FFT-based acceleration scheme is proposed for efficient MQS analysis of voxelized structures. Numerical results show that the proposed scheme outperforms well-established and widely-used FFT and SVD acceleration schemes for analyzing well-separated voxelized structures.

REFERENCES

- [1] A. C. Yucel, I. P. Georgakis, A. G. Polimeridis, H. Bağcı, and J. K. White, "VoxHenry: FFT-accelerated inductance extraction for voxelized geometries," *IEEE Trans. Microw. Theory Tech.*, vol. 66, no. 4, pp. 1723-1735, 2018.
- [2] M. Wang, C. Qian, J. K. White, and A. C. Yucel, "VoxCap: FFT-accelerated and Tucker-enhanced capacitance extraction simulator for voxelized structures," *IEEE Trans. Microw. Theory Techn.*, vol. 68, no. 12, pp. 5154-5168, Dec. 2020.
- [3] M. Wang, Y. Liu, P. Ghysels, and A. C. Yucel, "VoxImp: Impedance extraction simulator for voxelized structures," *IEEE Trans. Comput.-Aided Design Integr. Circuits Syst.*, early access, 2022.
- [4] M. Wang, C. Qian, E. D. Lorenzo, L. J. Gomez, V. Okhmatovski, and A. C. Yucel, "SuperVoxHenry: Tucker-enhanced and FFT-accelerated inductance extraction for voxelized superconducting structures," *IEEE Trans. Appl. Supercond.*, vol. 31, no. 7, pp. 1-11, 2021.
- [5] S. Kapur and D. E. Long, "IES³: A fast integral equation solver for efficient 3-dimensional extraction," in *Proc. ICCAD*, 1997, pp. 448-455.
- [6] K. Zhao, M. N. Vouvakis, and J.-F. Lee, "The adaptive cross approximation algorithm for accelerated method of moments computations of EMC problems," *IEEE Trans. Electromagn. Compat.*, vol. 47, no. 4, pp. 763-773, 2005.
- [7] S. B. Sayed, Y. Liu, L. J. Gomez, and A. C. Yucel, "A butterfly-accelerated volume integral equation solver for broad permittivity and large-scale electromagnetic analysis," *IEEE Trans. Antennas Propag.*, vol. 70, no. 5, pp. 3549-3559, 2021.
- [8] C. Qian, M. Wang, and A. C. Yucel, "Compression of far-fields in the fast multipole method via Tucker decomposition," *IEEE Trans. Antennas Propag.*, vol. 69, no. 10, pp. 6660-6668, 2021.

- [9] C. Qian and A. C. Yucel, "On the compression of translation operator tensors in FMM-FFT-accelerated SIE simulators via tensor decompositions," *IEEE Trans. Antennas Propag.*, vol. 69, no. 6, pp. 3359-3370, 2020.