



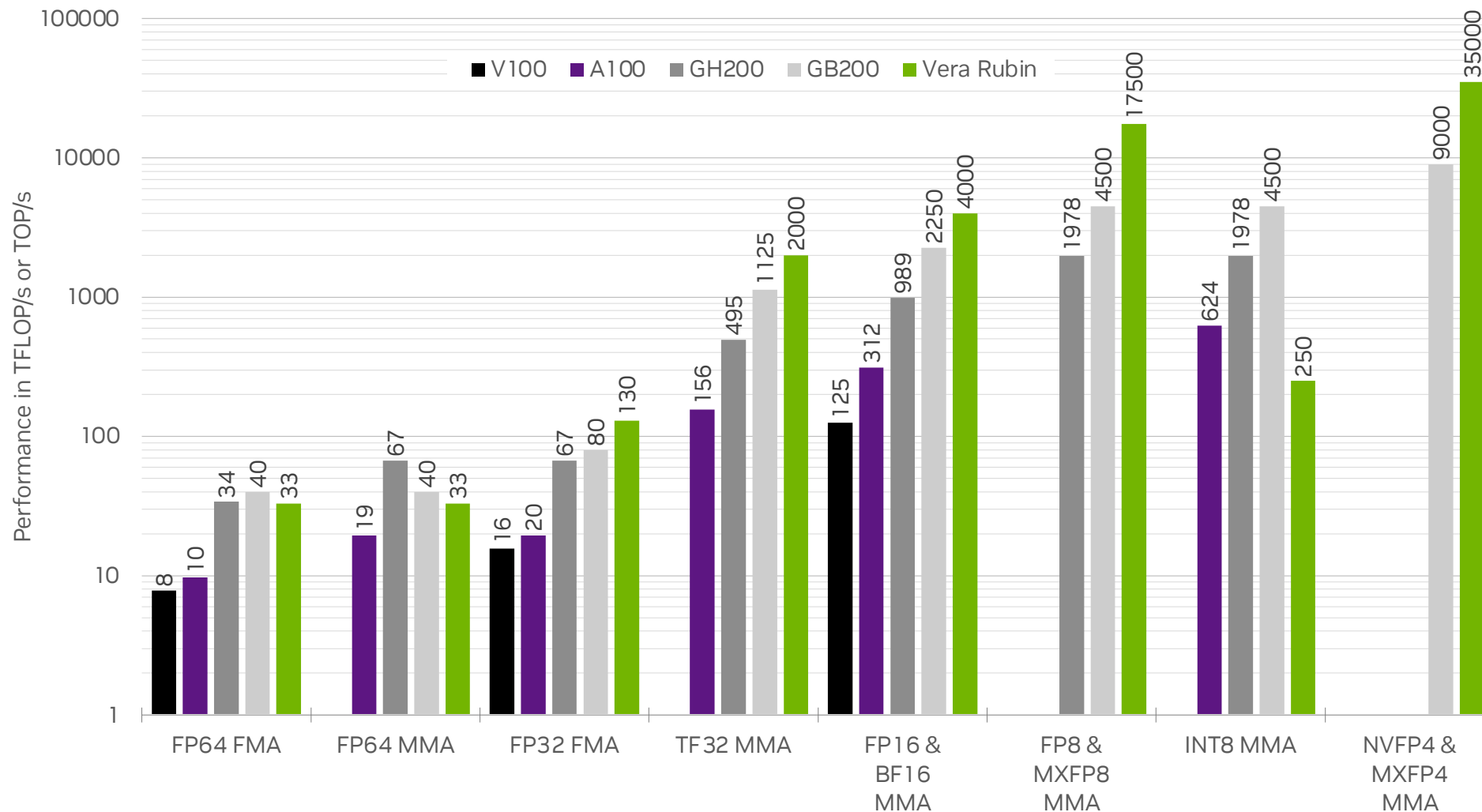
Dispelling Many of the Perceived Challenges in Adopting the Ozaki Scheme in Scientific Computing

Harun Bayraktar, Sr. Director of Engineering – Math Libraries
IPDPS 2026 MxP4S Workshop | May 25, 2026 | New Orleans, LA



Evolution of Peak Performance

Hardware Specs



arXiv:2306.11975v4 [cs.DC] 30 Mar 2024

10

Rio Yokota
rioyokota@gsic.titech.ac.jp
Tokyo Institute of Technology
Tokyo, Japan

Deep learning hardware achieves high throughput and low power consumption by reducing computing precision and specializing in matrix multiplication. For machine learning inference, fixed-point value computation is commonplace, where the input and output tensors are quantized to integers. However, some deep learning accelerators are now equipped with fast integer matrix multiplication units (IMMU). It is of significant interest to find a way to harness these IMMU to improve the performance of H3C applications while maintaining accuracy. We focus on the OpenAI schime, which computes a high-dimensional matrix multiplication by using lower-precision integer units. We first analyze the accuracy of the OpenAI schime by using IMMU. The experiment using integer TensorCores shows that we can compute double-precision matrix multiplication faster than cuBLAS and an existing OpenAI schime implementation on FP16 TensorCores on NVIDIA consumer GPUs. Furthermore, we propose a new integer TensorCores architecture that achieves up to 4.33x while maintaining the FP16 accuracy.

- Computing methodologies → Parallel algorithms; • Theory of computation → Numeric approximation algorithms.

matrix multiplication, fixed-point arithmetic, floating-point arithmetic. Tensor Cores

Hiroyuki Ootomo, Katsuhisa Ozaki, and Rio Yokota. 2018. DGEMM on Integer Matrix Multiplication Unit. In *Proceedings of Make sure to enter the correct conference title from your rights confirmation email (Conference acronym 'XX)*. ACM, New York, NY, USA, 12 pages. <https://doi.org/XXXXXXX.XXXXXXX>

The development of processors for machine learning is progressing rapidly, and researchers are exploring their potential for use in other HPC applications [1]. These processors often use low/mixed-precision floating-point, or low-bit-length integer computing units, such as NVIDIA Tensor Cores [26] and Google TPUs [16–18], to achieve high throughput. While the training of deep learning models typically requires floating-point types like FP16, Bfloat16, FP32

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies on the first page. Copyrights for components of this work must be honored. Abstracting with credit is permitted to post on servers or to redistribute to lists, requires fee. Request permissions from permissions.acm.org

ACM Conference on XX, June 03-05, 2018, Waco, TX
© 2018 Association for Computing Machinery
ACM ISBN 978-1-4503-XXXXX-X/18/06... \$15.00
<http://doi.org/10.1145/XXXXXX.XXXXXX>



FP32, the inference involves quantizing model parameters and input/output values to use lower-bit-length integer formats. This reduces data size, processor circuit area, and power consumption [14, 16]. As a result, high-performance processors like NVIDIA TensorRT [17] and Intel MKL-DNN [18] have been developed, which are equipped with low-bit-length integer matrix multiplication units (MMU). For instance, NVIDIA GPUs provide a DP4A instruction that can efficiently compute the inner product of two length-4 8-bit integer (INT8) vectors, and the result is in a 32-bit integer (INT32) format. In addition, NVIDIA TensorRT [17] uses the multiplication of INT8 matrices with INT32 accumulation from the Turing architecture and the multiplication of 4-bit integer (INT4) matrices from the Ampere architecture. These integer Tensor Cores are 4 times faster than the floating-point Tensor Cores. Other processors, such as Google TPU v1 [18], Intel AMX-INT8 [5], and Gropq TSP2 [2], also support the multiplication of INT8 matrices with INT32 accumulation. In light of these advancements, our research aims to develop a high-performance inference framework for deep learning models on high-performance computing (HPC) applications.

Recent research has explored the use of machine learning processors for HPC applications, in particular low- and mixed-precision floating-point matrix multiplication units (FMMU). For instance, Haider et al. used the FP16 Tensor Core to solve FP64 linear equations with an iterative refinement method [10]. Finkelstein et al. used Tensor Cores to accelerate quantum circuit simulations, achieving quantum response calculations with sufficient accuracy [7]. Our previous study has utilized FP16 Tensor Cores to improve the throughput of quantum circuit simulation by emulating single-precision matrix multiplication using an error correction method and avoiding rounding inside Tensor Cores without loss of accuracy [27, 28]. However, despite the promising results obtained by these studies, few have explored the potential of integer matrix multiplications. In this paper, we explore the potential of integer matrix multiplication by leveraging integer matrix multiplication units, for instance, INT8 Tensor Cores in NVIDIA GPUs, for HPC applications.

Even before the advent of deep learning and specialized hardware for it, researchers had been exploring ways to perform high-precision matrix multiplication on lower-precision computing units. One such approach is the double-double and quad-double precision scheme, where a single value is represented as the sum of two and four FP64 values, respectively, and arithmetic operations are performed using a sequence of FP64 operations [13]. GEMM and other double-double precision have been evaluated on the AMD Cypress GPUs [24, 24]. Another approach is the split-matrix scheme, which also splits a value into multiple pieces [29, 30]. In this scheme, input matrices are split into matrix slices, and the resulting matrix is obtained by multiplying these slices. The key idea is to reduce the multiplication of these slices. The ob-

Journal Title
XX(X):1-13
©The Author(s) 2016
Reprints and permission:
sagepub.com/journalsPermissions.nav
DOI: 10.1177/TcBeAssigned
www.sagepub.com/

arXiv:2504.08009v3 [cs.MS] 27 Apr 2025

matrix multiplication, floating-point arithmetic, matrix engines, high-precision emulation

This paper surveys numerical methods for emulating high-precision matrix multiplication. When the accuracy of the computed results is insufficient, multiple-precision arithmetic provides a viable alternative (Higham 2018). For such computations, high-precision datatypes, such as the built-in `float128_t` in the C++23 standard, and multiple-precision arithmetic libraries, such as MPFR (Fousse et al. 2025), offer robust and reliable functionality. In the domain of linear algebra, `MPPLA` (Fousse et al. 2025) is available for high-precision port when it is unnecessary to extend floating-point numbers and only pseudo-extension, *multiple-complex* an efficient solution. Such methods quad-word arithmetic (Hida et al. 2000) as well as triple-word arithmetic (A

When the computational task is limited to matrix multiplication, the Ozaki scheme (Ozaki et al. 2012, 2013) is recognized as a highly reliable method. It achieves high accuracy by leveraging standard floating-point operations and high performance by exploiting optimized routines, e.g., General Matrix-Matrix Multiplication (GEMM), in Basic Linear Algebra Subroutines (BLAS).

In recent years, increasing attention has been paid to matrix engines optimized for low-precision arithmetic. From the perspective of power efficiency, mixed-precision computation utilizing low-precision formats has also



Sciences, Shibaura Institute of Technol-
 onal Science, Japan
 u, Minuma-ku, Saitama-shi, Saitama 337-
 o.jp



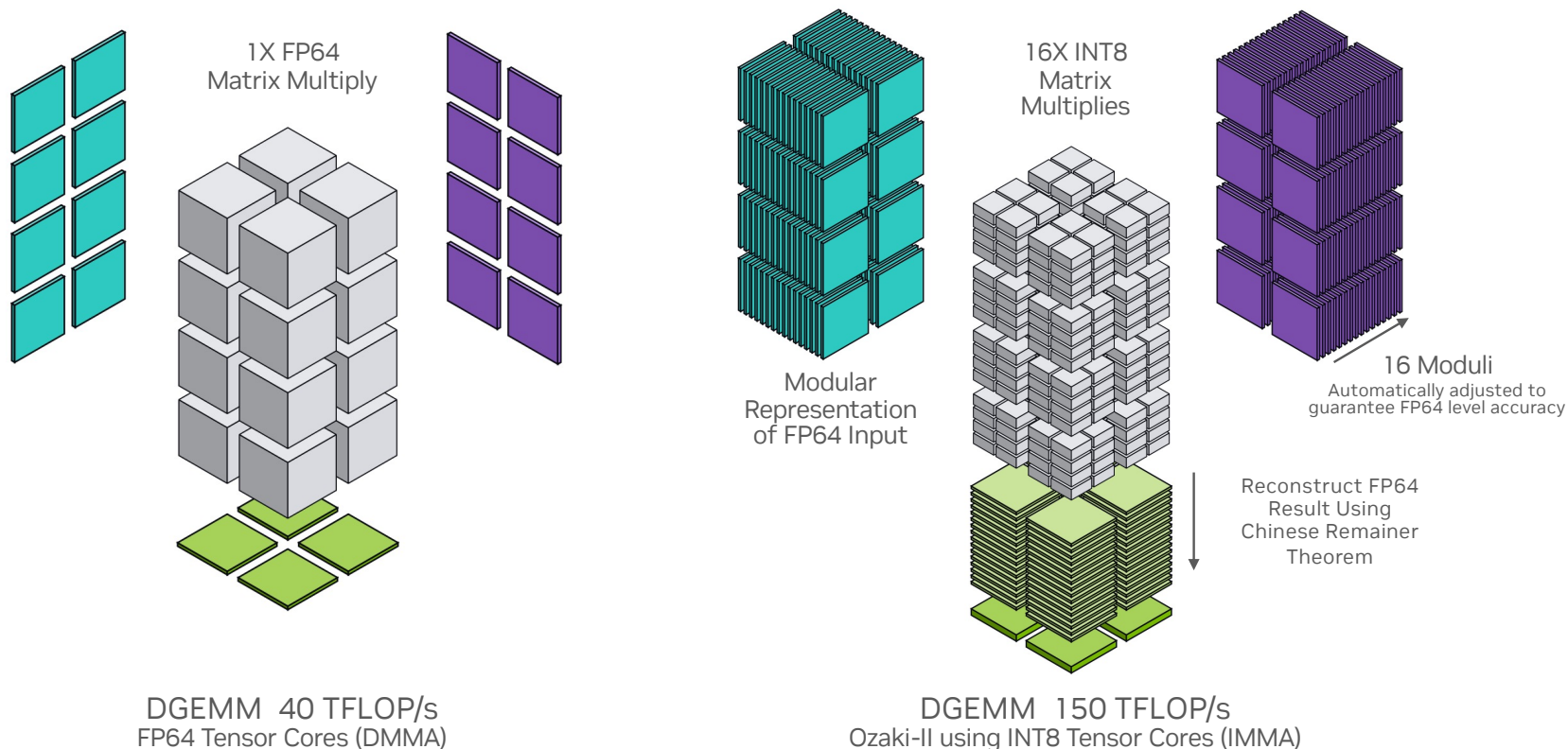
Sciences, Shibaura Institute of Technol-
ogy, Japan

u, Minuma-ku, Saitama-shi, Saitama 337-
c/o

Prepared using sagej.cls [Version: 2017/01/17 v]

Matrix Multiplication Emulation With High Throughput Tensor Cores

Ozaki-II method delivers 3.5X speed-up on GB200 for FP64 matrix multiplies



Complementing Existing Research & Innovations to Productize Great Technology

arXiv:2306.11975v4 [cs.DC] 30 Mar 2024



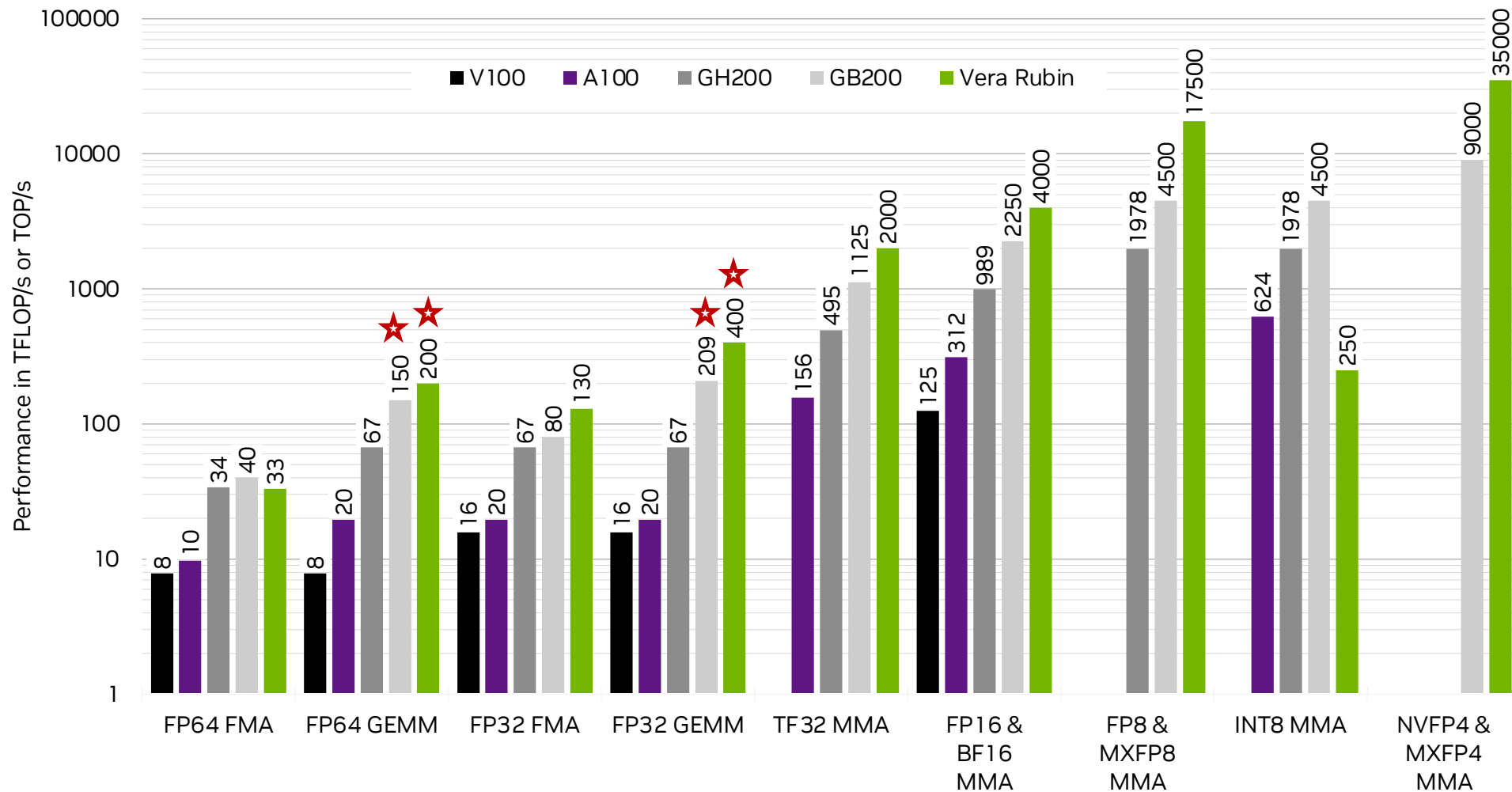
arXiv:2504.08009v3 [cs.LG] 27 Apr 2025



Productization in Libraries like cuBLAS, cuSOLVER,

Evolution of Peak Performance

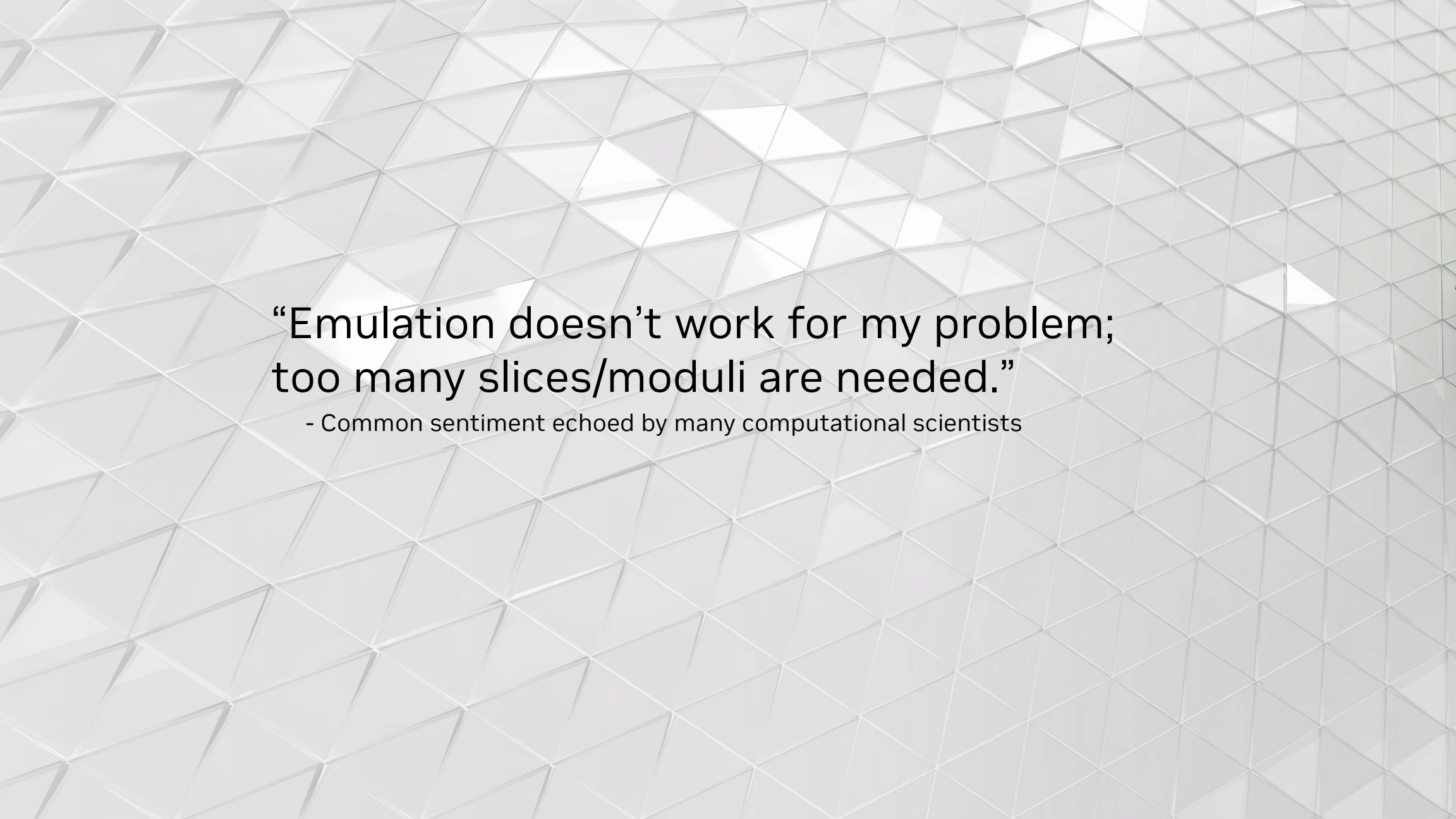
Hardware Specs & Emulation



Five Perceived Challenges With Ozaki Scheme

Hampering adoption by some in Scientific Computing

1. Emulation doesn't work for my problem; too many slices/moduli are needed
2. Results are not as accurate as IEEE FP64
3. Emulation uses too much memory
4. Emulation doesn't deliver performance for non-square matrices
5. Emulation doesn't deliver performance for small matrices

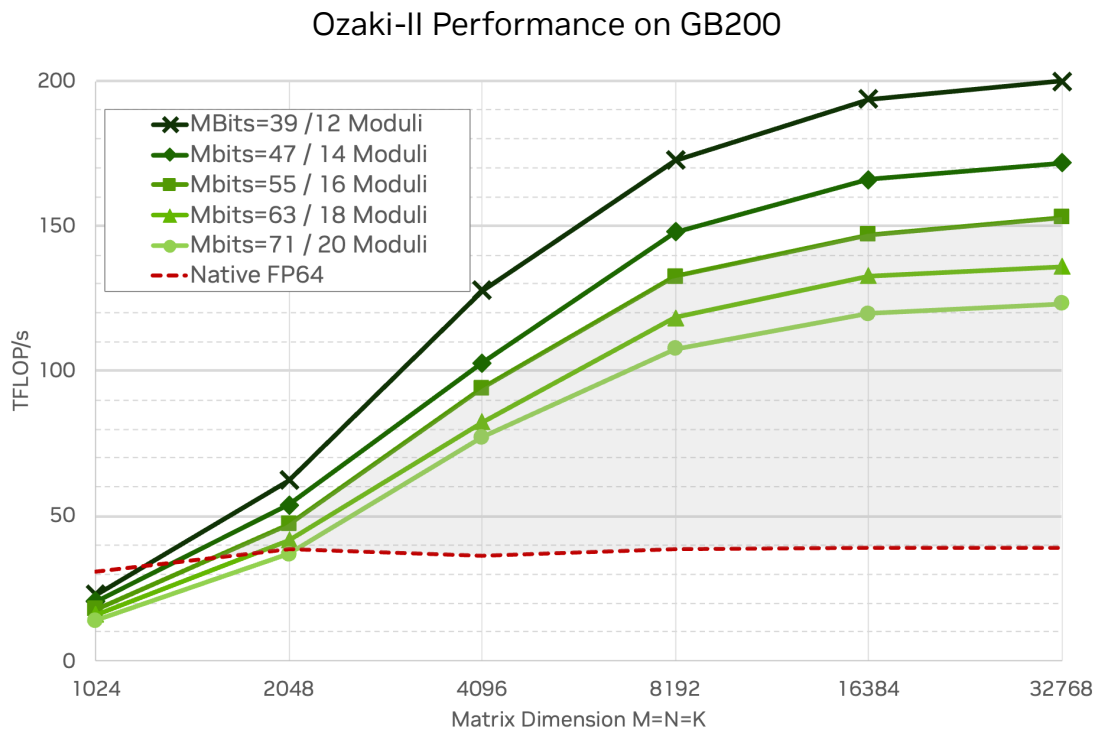


“Emulation doesn’t work for my problem;
too many slices/moduli are needed.”

- Common sentiment echoed by many computational scientists

Ozaki-II Performance

Does emulated DGEMM behave like FP64 as the problem gets harder? Yes, but at a price

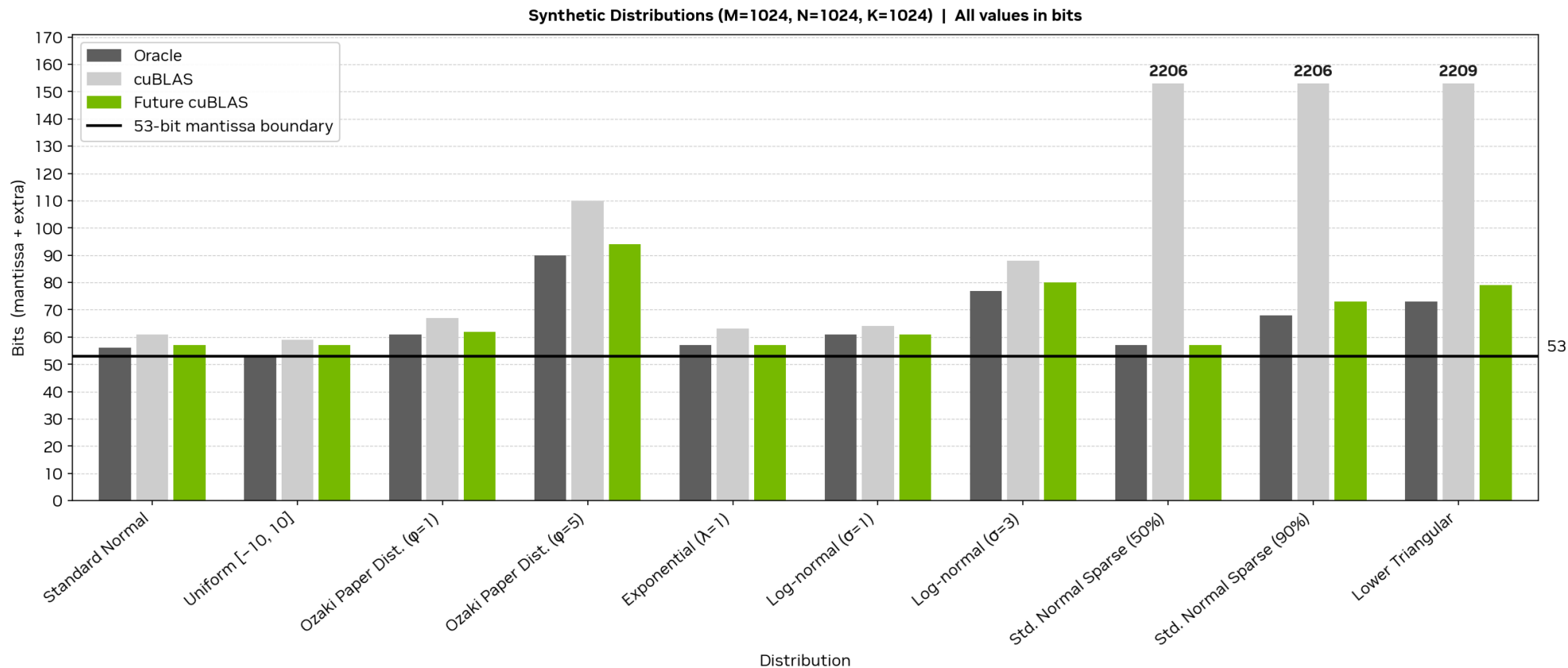


The current ESC algorithm is conservative with a loose bound and for some input data it is overly conservative, leading to poor speed-ups or fallback to native FP64

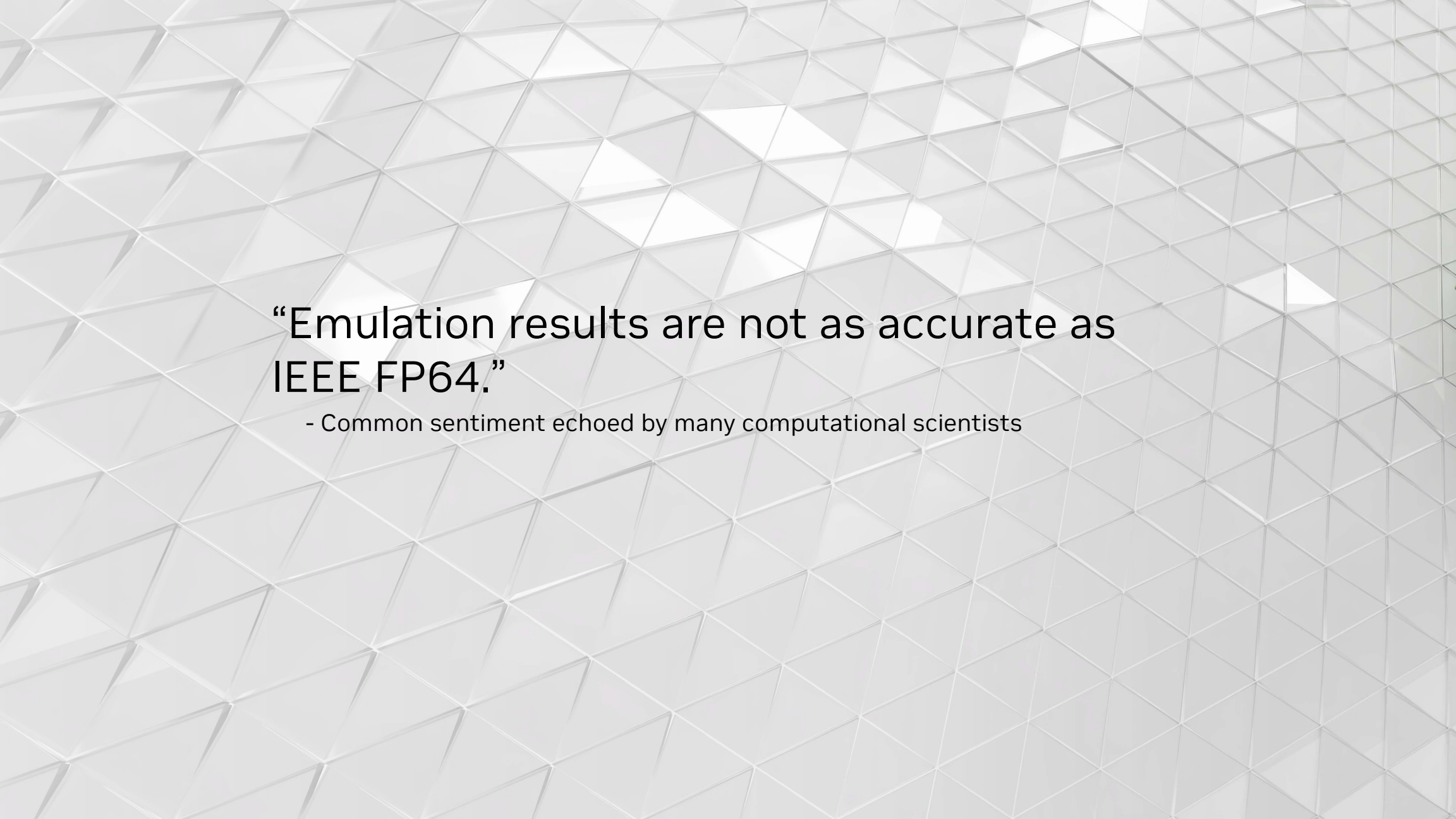
Improving the Exponent Span Capacity (ESC) Algorithm

The current ESC in cuBLAS has very loose bound (too conservative)

We have a new ESC algorithm we will productize later in 2026 that delivers significant improvements (details of the new algorithm will be published when ready as was done with the current version)



Preliminary data | Data subject to change



“Emulation results are not as accurate as
IEEE FP64.”

- Common sentiment echoed by many computational scientists

IEEE 754 Compliance of Matmul Emulation

Current Limitations

NaN/Inf/+0,-0

- We can detect if these are present on the input matrices and fall back to FP64 native
- We can NOT detect nor propagate these when they happen during the dot product (not necessarily a bad thing)
- We treat +0 and -0 as the same value

We do not observe strict IEEE compliance

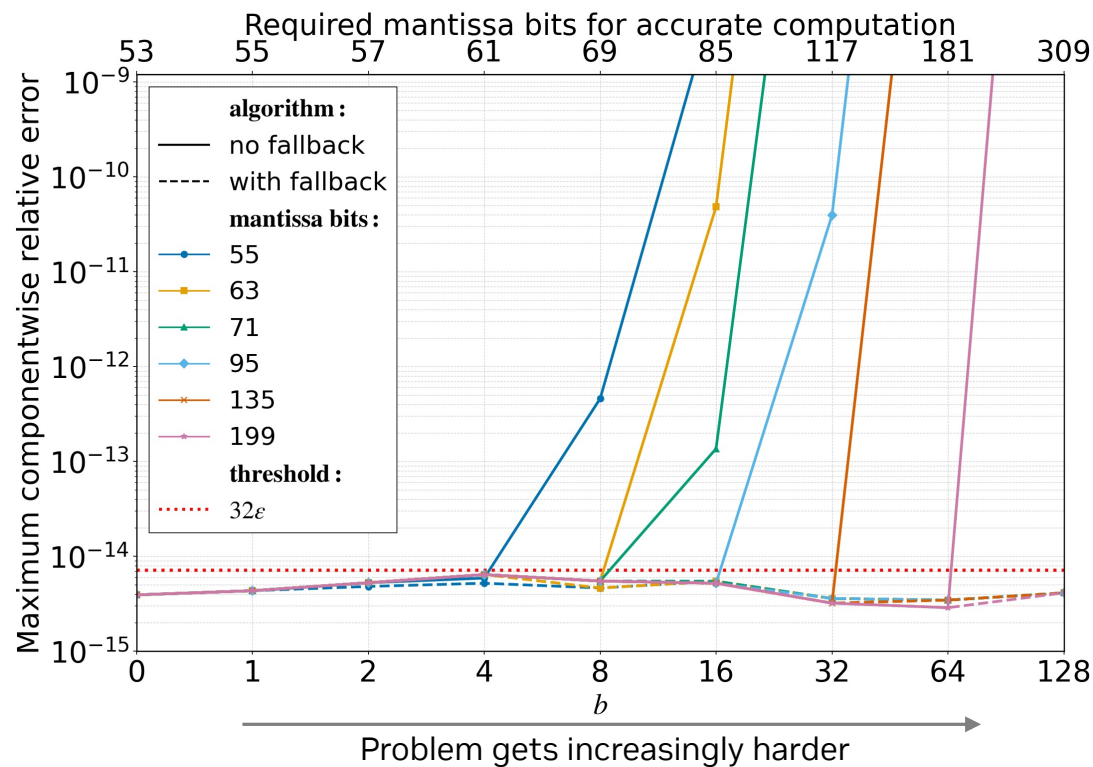
- Since exponents are coupled with results only at certain points (number of times depends on k dimension of matrix)
- Because we store only one exponent per row/column of the matrix and do not check for overflow at a scalar product level, we can generate NaNs and "sticky" +/-Infs, but not at the IEEE-prescribed granularity

These limitations don't speak to the accuracy of the results which are FP64 level accurate



Putting ESC to the Test

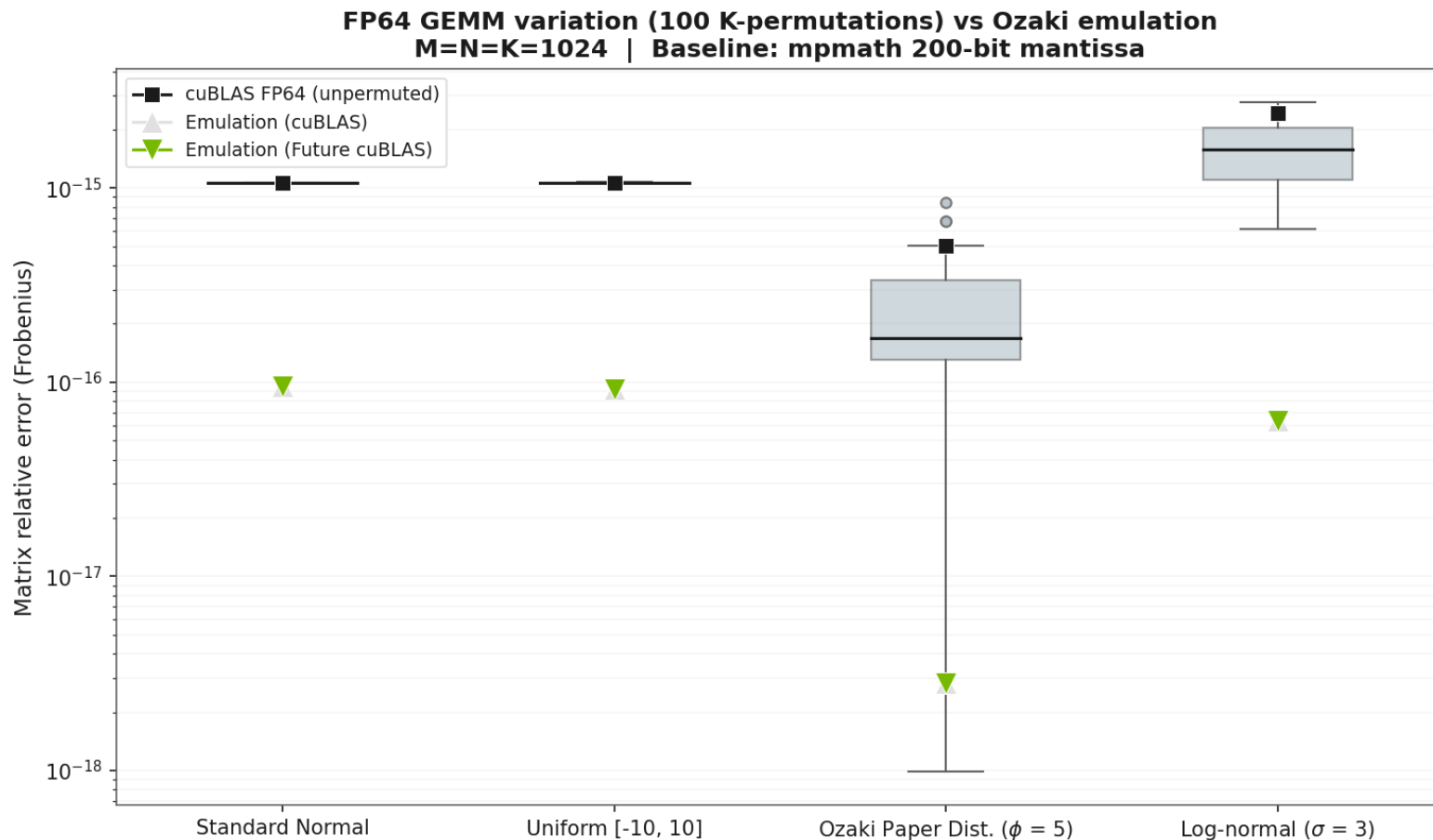
Does emulated DGEMM behave like FP64 as the problem gets harder? Yes, but at a price



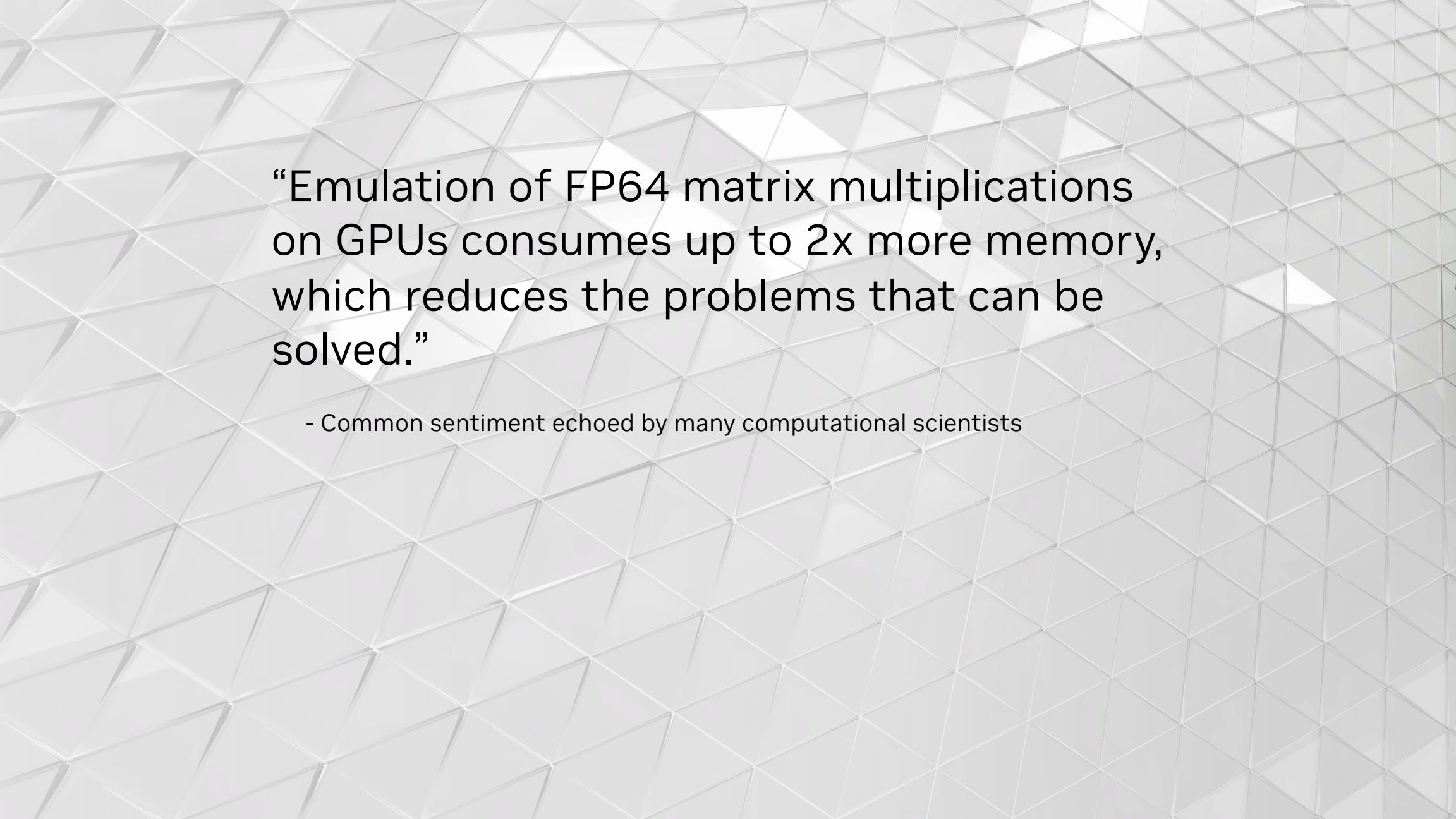
The current ESC algorithm is conservative with a loose bound and for some input data it is overly conservative, leading to poor speed-ups or fallback to native FP64

Putting the new ESC to the Test

Does emulated DGEMM behave like FP64 as the problem gets harder? Yes, and will deliver much better performance



Preliminary data | Data subject to change

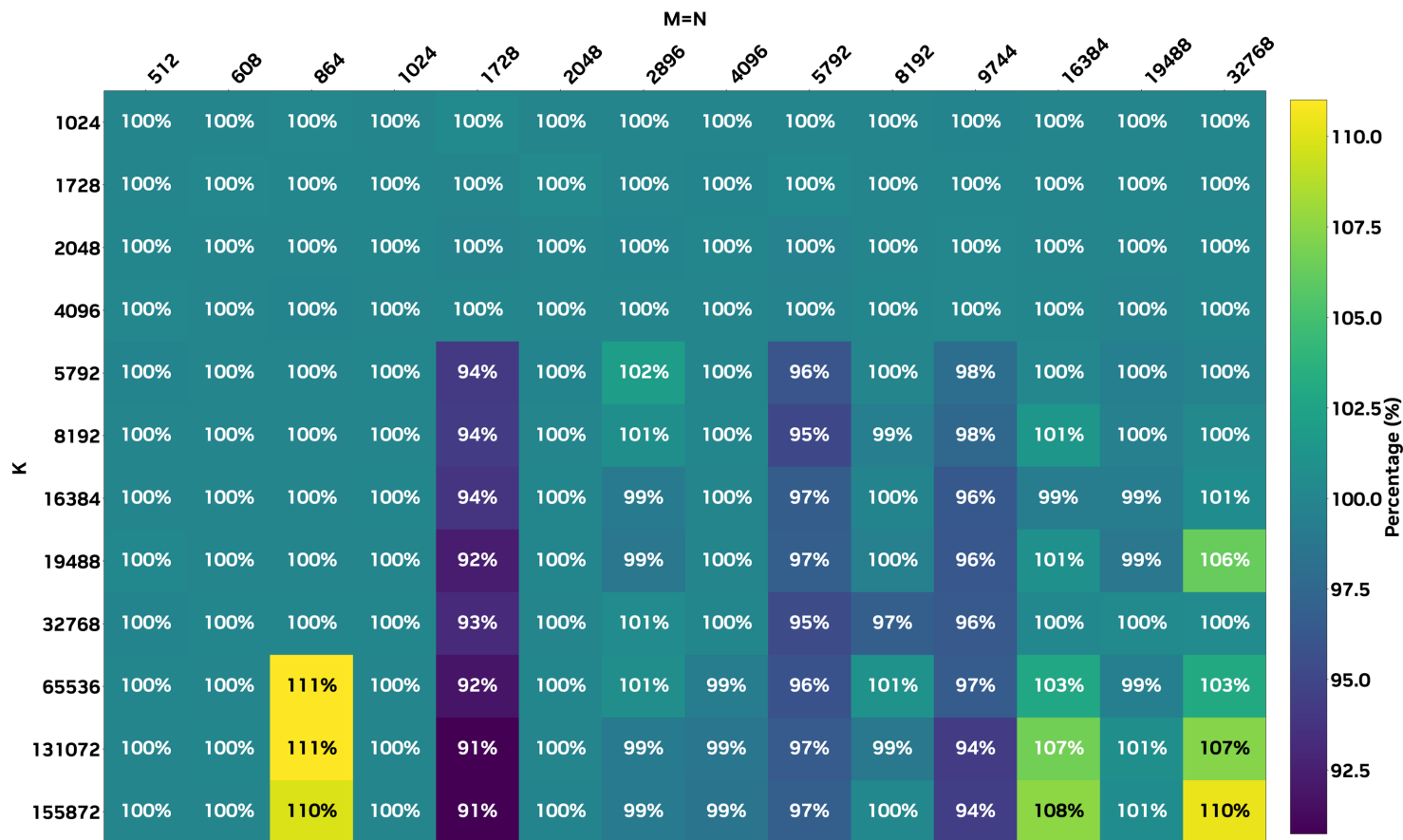


“Emulation of FP64 matrix multiplications on GPUs consumes up to 2x more memory, which reduces the problems that can be solved.”

- Common sentiment echoed by many computational scientists

Impact of Available Memory on DGEMM Performance

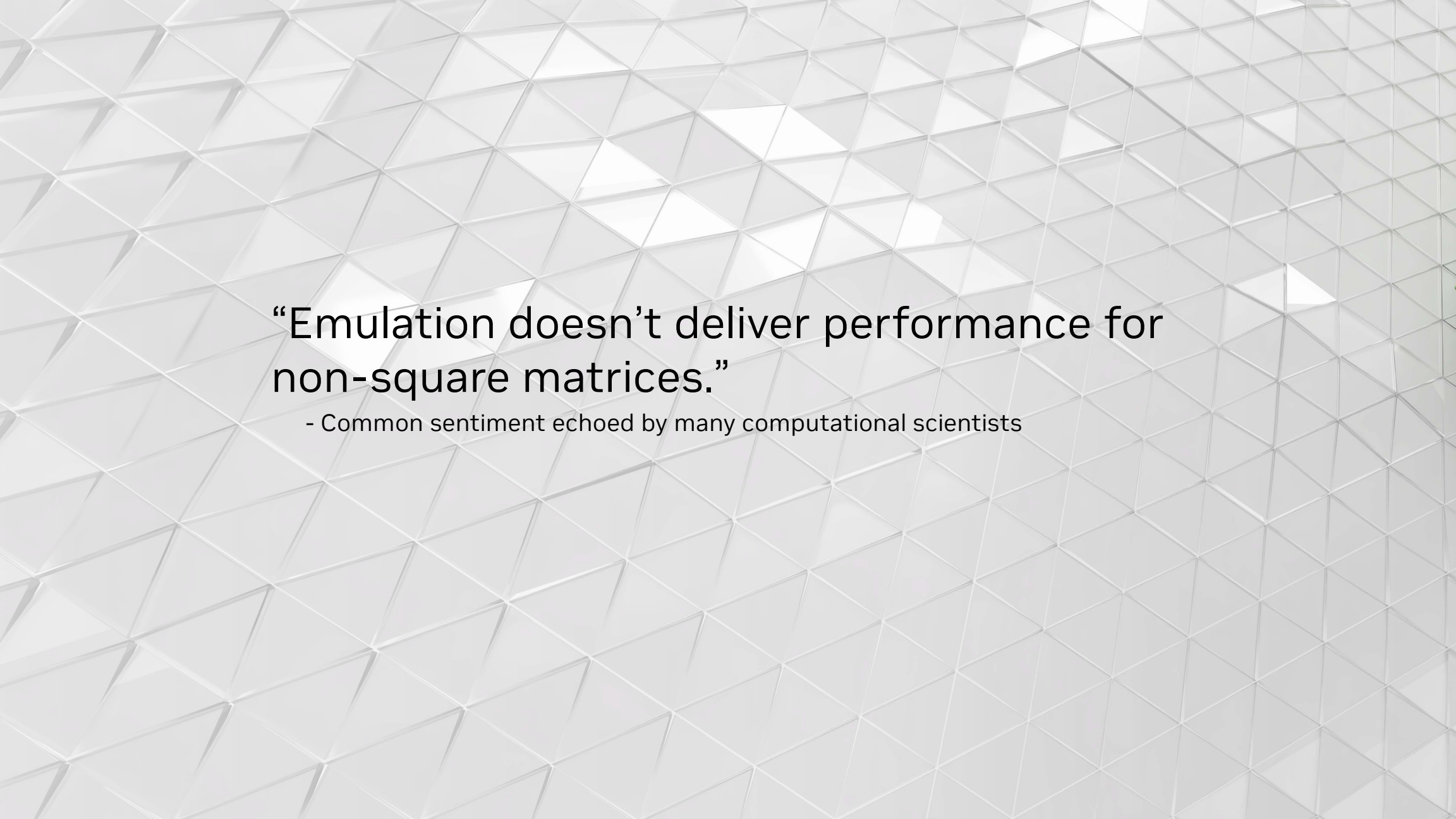
Performance Ratio Using cuBLAS (8GB : Unlimited Memory Tiling) on the B200 GPU



The new blocking strategy using the limited memory policy leads to noticeable performance uplift in some cases

Coming soon in a 2026 cuBLAS release

Preliminary data | Data subject to change

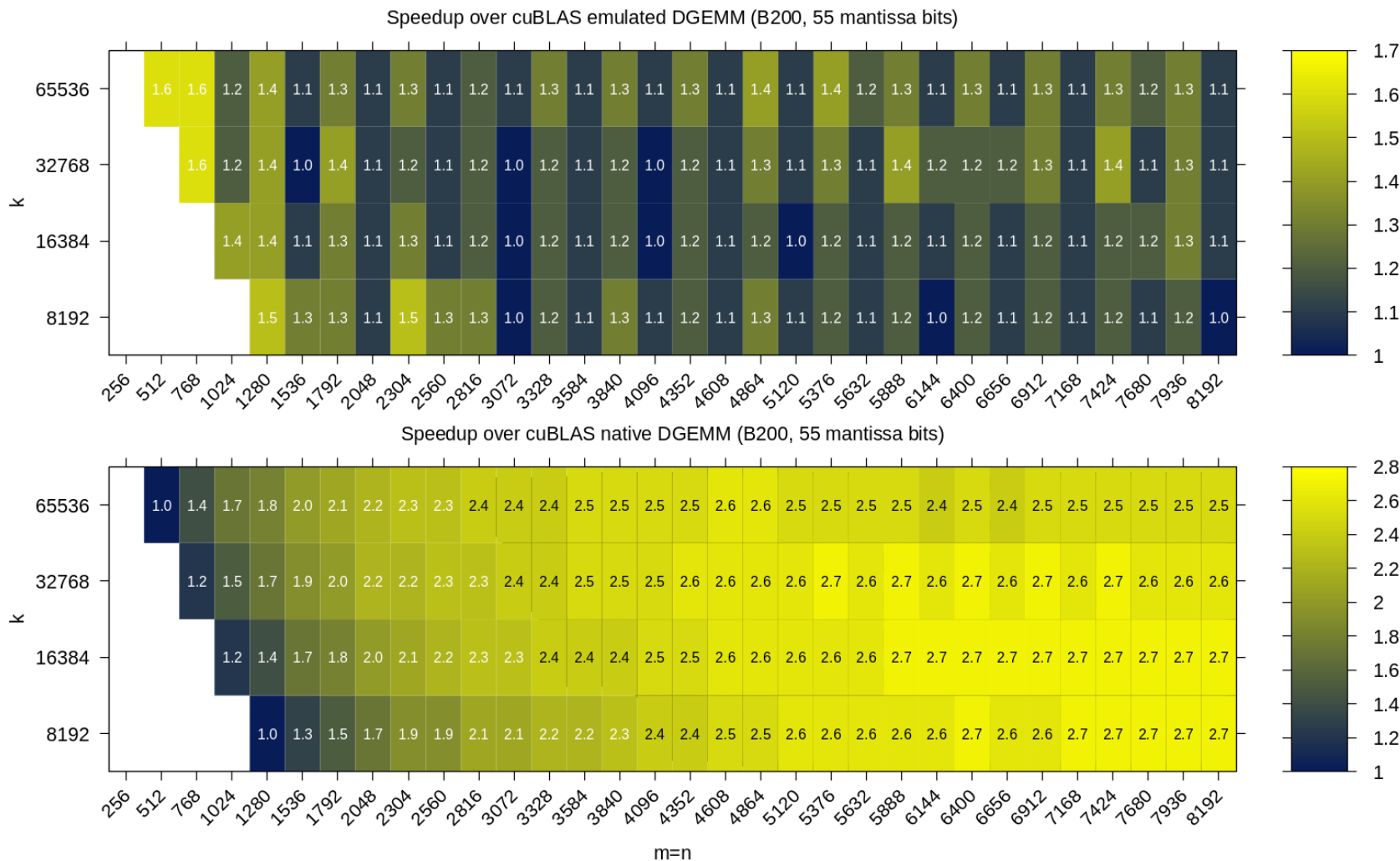


“Emulation doesn’t deliver performance for non-square matrices.”

- Common sentiment echoed by many computational scientists

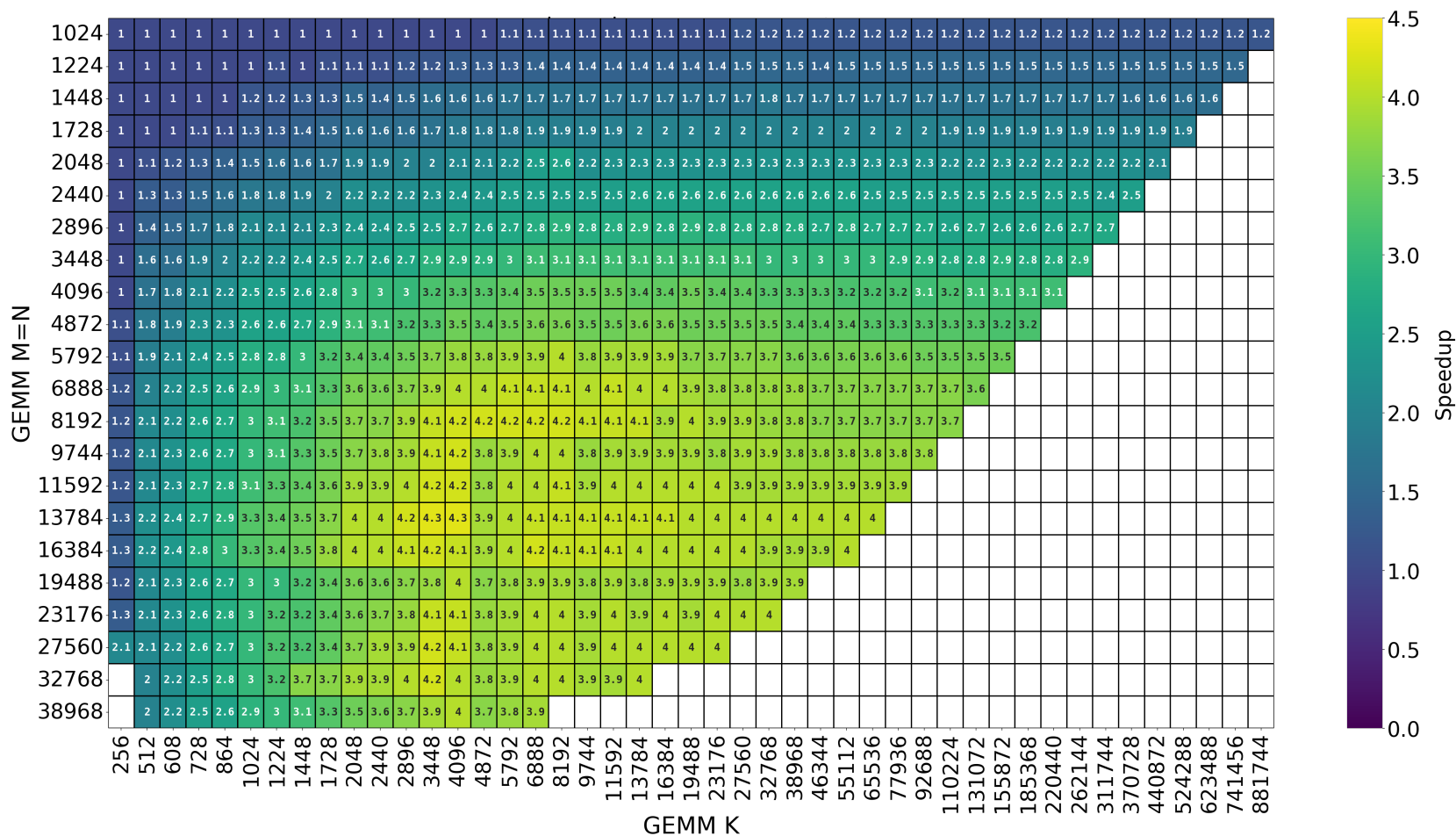
Emulation Performance for Non-Square Matrices (Ozaki-I)

Overheads plus memory vs compute bound regime pose challenges



Performance of DGEMM on GB200 GPUs

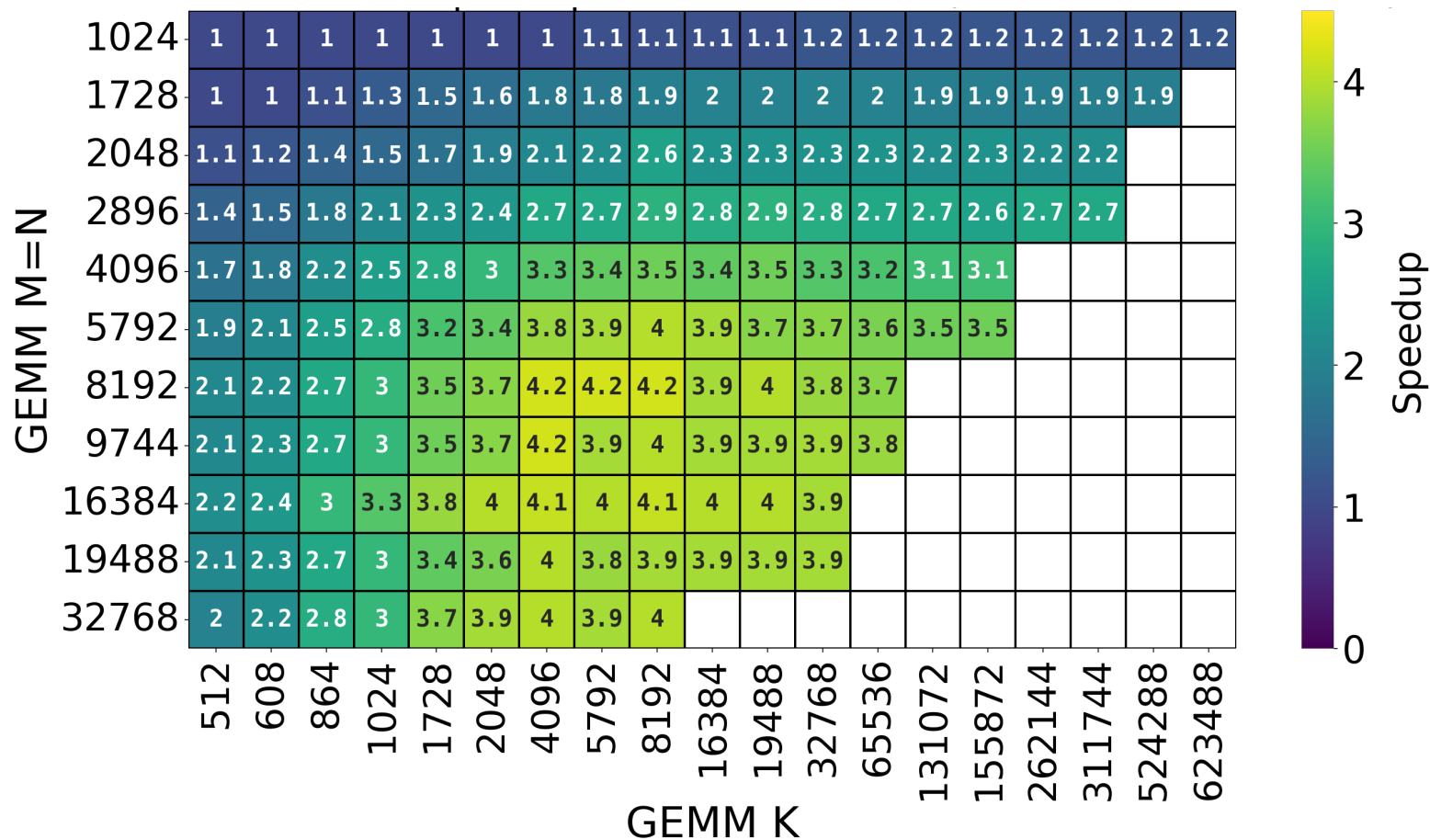
Speedup of Ozaki-II Over Native FP64 (Coming to cuBLAS Later in 2026)



Preliminary data | Data subject to change

Performance of DGEMM on GB200 GPUs

Speedup of Ozaki-II Over Native FP64 (Coming to cuBLAS Later in 2026)



Preliminary data | Data subject to change

Five Perceived Challenges With Ozaki Schemes

Hampering adoption by some in Scientific Computing

SOLVED

Emulation doesn't work for my problem, too many slices/moduli are needed **(coming soon)**

SOLVED

Results are not as accurate as IEEE FP64

SOLVED

Emulation uses too much memory **(coming very soon)**

SOLVED

Emulation doesn't deliver performance for non-square matrices

5. Emulation doesn't deliver performance for small matrices ← **Likely requires new APIs**

Impact of Improvements on Quantum Chemistry

cuEST library using the Ozaki-II New ESC Algorithm in cuBLAS

