

Automatic Generation of Numerical Codes for GPUs Using LLMs¹⁾

Daichi Mukunoki¹ Shunichiro Hayashi² Tetsuya Hoshino¹ Ryo Mikasa³
Koki Morita² Takahiro Katagiri¹

¹Information Technology Center, Nagoya University, ²Graduate School of Informatics, Nagoya University,
³School of Informatics, Nagoya University

JHPCN Field Workshop
State-of-the-Art in Code Generative AI for High-Performance Computing
Dec. 5, 2025.

¹⁾This research was supported by the Joint Usage/Research Center for Interdisciplinary Large-scale Information Infrastructures (JHPCN) and the High Performance Computing Infrastructure (HPCI) (Project ID: jh250015). It was also supported by JSPS KAKENHI Grant Numbers JP23K11126 and JP24K02945.

HPC-GENIE^a – High-Performance Computing with Generative Neural Intelligence for Execution

- A project for HPC code generation using LLMs at the Information Technology Center, Nagoya University
 - ▶ Project leader: Takahiro Katagiri
 - ▶ Sub-leader: Daichi Mukunoki

Key Missions

- Automatic HPC code optimization
- Fortran to GPU **for Fugaku NEXT**
- Technologies for local LLMs
- And more ...

^ahttps://www.hpc.itc.nagoya-u.ac.jp/menu/hpc_genie.html



BLAS Code Generation Using o4-mini & GPT-4.1²⁾

Number of Correct Codes (in 10 Generations)

Level	Routine	Non-Optimized Code		Optimized Code	
		GPT-4.1	o4-mini	GPT-4.1	o4-mini
1	dasum	10	10	9	7
	daxpy	10	10	10	10
	ddot	10	10	8	9
	idamax	3	10	2	9
	dnrm2	10	10	7	5
	drot	10	10	8	9
2	dgemv	8	9	3	6
	dger	10	10	7	7
	dsymv	8	8	0	2
	dsyr2	8	9	2	7
	dtrmv	3	4	0	5
	dtrsv	8	9	0	4
3	dgemm	10	10	3	0
	dsymm	4	8	3	3
	dsyrk	3	10	0	1
	dsyr2k	3	10	0	0
	dtrmm	0	1	0	0
	dtrsm	0	0	0	0

- OpenAI GPT-4.1 & o4-mini – entry-level general-purpose LLMs
- **Generation with one simple prompt** (for optimized code): *“Implement #ROUTINE# routine in BLAS in C language. Thread parallelization, SIMD vectorization, and cache blocking should be considered for speedup.”*
- **Given just a routine name, working code can often be generated.**

²⁾D. Mukunoki, S. Hayashi, T. Hoshino, T. Katagiri, “Performance Evaluation of General Purpose Large Language Models for Basic Linear Algebra Subprograms Code Generation”, arXiv preprint arXiv:2507.04697, 2025.

Reference Fortran code (part)

```

1      IF (NOTB) THEN
2          IF (NOTA) THEN
3              *      Form C := alpha*A*B + beta*C.
4                  DO 90 J = 1,N (略)
5
6
7                  DO 80 L = 1,K
8                      TEMP = ALPHA*B(L,J)
9                      DO 70 I = 1,M
10                         C(I,J) = C(I,J) + TEMP*A(I,L)
11
12                         CONTINUE
13
14                         CONTINUE
15
16                         CONTINUE
17
18                         ELSE
19
20                         *      Form C := alpha*A**T*B + beta*C
21                             DO 120 J = 1,N
22                                 DO 110 I = 1,M
23                                     TEMP = ZERO
24                                     DO 100 L = 1,K
25                                         TEMP = TEMP + A(L,I)*B(L,J)
26
27                                         CONTINUE
28                                         IF (BETA.EQ.ZERO) THEN
29                                             C(I,J) = ALPHA*TEMP
30                                         ELSE
31                                             C(I,J) = ALPHA*TEMP + BETA*C
32                                                 (I,J)
33
34                                         END IF
35
36                                     CONTINUE
37
38                                     CONTINUE
39
40                                 END IF
41
42                                 CONTINUE
43
44                                 ELSE
45
46                                 IF (NOTA) THEN

```

Generated C code (an example)

```

1      #include <stdio.h>
2
3      void GPTBLAS_dgemm(const char *transa, const char *
4          transb, const int *m, const int *n, const int *k,
5          const double *alpha, const double *a, const int *
6          lda, const double *b, const int *ldb, const double
7          *beta, double *c, const int *ldc)
8
9      {
10         printf("[gptblas]");
11         int i, j, l;
12         int nota = (*transa == 'N' || *transa == 'n');
13         int notb = (*transb == 'N' || *transb == 'n');
14         for (j = 0; j < *n; j++) {
15             for (i = 0; i < *m; i++) {
16                 double cij = (*beta) * c[i + j * (*ldc)];
17                 double sum = 0.0;
18                 for (l = 0; l < *k; l++) {
19                     double a_val = nota ? a[i + l * (*lda)]
20                         : a[l + i * (*lda)];
21                     double b_val = notb ? b[l + j * (*ldb)]
22                         : b[j + l * (*ldb)];
23                     sum += a_val * b_val;
24                 }
25                 c[i + j * (*ldc)] = cij + (*alpha) * sum;
26             }
27         }
28     }

```

Code structure is not necessarily the same as the reference Fortran code

Level-1 BLAS Performance (GB/s, n=16777216)

Routine	Ref	Optimized C code from routine name only				Optimized C code from reference Fortran code			
		GPT-4.1		o4-mini		GPT-4.1		o4-mini	
dasum	6.0	68.0	(11.4x)	63.5	(10.7x)	61.5	(10.3x)	63.6	(10.7x)
daxpy	17.2	77.6	(4.5x)	68.8	(4.0x)	71.9	(4.2x)	67.8	(3.9x)
ddot	11.7	65.8	(5.6x)	63.1	(5.4x)	61.0	(5.2x)	60.8	(5.2x)
idamax	7.1	63.8	(9.0x)	65.1	(9.2x)	60.3	(8.6x)	62.8	(8.9x)
dnrm2	5.1	66.5	(12.9x)	64.0	(12.5x)	60.6	(11.8x)	63.5	(12.4x)
drot	10.6	40.6	(3.8x)	34.1	(3.2x)	34.9	(3.3x)	34.4	(3.3x)
drotm	11.1	39.6	(3.6x)	36.6	(3.3x)	34.7	(3.1x)	34.1	(3.1x)

- “Ref”: reference Fortran code (non-parallelized, non-optimized)
- Xeon Gold 6230 (20 cores, Cascade Lake) $\times$ 2, 40 threads, gcc/gfortran 11.3.0 -march=native
- Best result among the 10 generated codes
- Blank entries indicate that no working code was generated

Level-2 BLAS Performance (GB/s, m=n=8129)

Routine	Parameters	Ref	Optimized C code from routine name only				Optimized C code from reference Fortran code			
			GPT-4.1		o4-mini		GPT-4.1		o4-mini	
dgemv	trans=N	9.1	5.4	(0.6x)	30.0	(3.3x)	32.3	(3.5x)	6.6	(0.7x)
	trans=T	6.7	68.9	(10.4x)	66.3	(10.0x)	67.6	(10.2x)	65.9	(9.9x)
dger		18.5	46.3	(2.5x)	64.9	(3.5x)	66.4	(3.6x)	64.2	(3.5x)
dsymv	uplo=L	6.1			4.6	(0.8x)			4.3	(0.7x)
	uplo=U	6.1			5.3	(0.9x)			4.7	(0.8x)
dsyr	uplo=L	17.4	32.4	(1.9x)	64.0	(3.7x)	64.5	(3.7x)	60.9	(3.5x)
	uplo=U	17.5			61.5	(3.5x)	64.4	(3.7x)	58.9	(3.4x)
dsyr2	uplo=L	7.6	26.4	(3.5x)	63.7	(8.3x)	61.0	(8.0x)	60.2	(7.9x)
	uplo=U	15.5	29.0	(1.9x)	59.6	(3.9x)	62.1	(4.0x)	58.6	(3.8x)
dtrmv	uplo=L, trans=N, diag=N	7.6			7.3	(1.0x)	3.0	(0.4x)	3.0	(0.4x)
	uplo=L, trans=N, diag=U	7.5			7.7	(1.0x)	2.9	(0.4x)	2.9	(0.4x)
	uplo=L, trans=T, diag=N	6.5			56.9	(8.8x)	3.3	(0.5x)	3.2	(0.5x)
	uplo=L, trans=T, diag=U	6.5			56.3	(8.7x)	3.3	(0.5x)	3.2	(0.5x)
	uplo=U, trans=N, diag=N	9.0			7.6	(0.8x)	3.1	(0.3x)	3.1	(0.3x)
	uplo=U, trans=N, diag=U	8.8			7.6	(0.9x)	3.1	(0.3x)	3.1	(0.4x)
	uplo=U, trans=T, diag=N	6.1			56.7	(9.3x)	3.2	(0.5x)	3.2	(0.5x)
	uplo=U, trans=T, diag=U	6.1			56.4	(9.3x)	3.2	(0.5x)	3.2	(0.5x)

- "Ref": reference Fortran code (non-parallelized, non-optimized)
- Xeon Gold 6230 (20 cores, Cascade Lake) × 2, 40 threads, gcc/gfortran 11.3.0 -march=native
- Best result among the 10 generated codes
- Blank entries indicate that no working code was generated

Level-3 BLAS Performance (GFlops/s, m=n=k=2048)

Routine	Parameters	Ref	Optimized C code from routine name only				Optimized C code from reference Fortran code			
			GPT-4.1		o4-mini		GPT-4.1		o4-mini	
dgemm	transa=N, transb=N	2.7	17.4	(6.5x)	20.5	(7.7x)	18.0	(6.8x)	20.8	(7.8x)
	transa=N, transb=T	2.5	16.5	(6.5x)			16.1	(6.3x)	20.7	(8.1x)
	transa=T, transb=N	1.8	16.8	(9.3x)	0.8	(0.5x)	29.2	(16.1x)	23.8	(13.1x)
	transa=T, transb=T	0.3	15.7	(46.3x)			3.5	(10.4x)	20.3	(59.8x)
dsymm	side=L, uplo=L	3.6	13.6	(3.7x)	17.2	(4.7x)	19.9	(5.5x)	22.3	(6.2x)
	side=L, uplo=U	3.5	13.9	(4.0x)	17.3	(4.9x)	19.9	(5.7x)	22.4	(6.4x)
	side=R, uplo=L	2.7	13.4	(4.9x)	16.4	(6.0x)	19.0	(7.0x)	21.2	(7.8x)
	side=R, uplo=U	2.7	13.4	(4.9x)	21.3	(7.8x)	18.7	(6.8x)	21.1	(7.7x)
dsyrk	uplo=L, trans=N	1.0			2.9	(2.9x)	18.2	(18.2x)	16.6	(16.6x)
	uplo=L, trans=T	1.9	20.1	(10.6x)	21.8	(11.5x)	27.7	(14.7x)	23.0	(12.2x)
	uplo=U, trans=N	2.3			3.8	(1.6x)	16.5	(7.1x)	16.5	(7.1x)
	uplo=U, trans=T	1.9	36.5	(19.2x)	21.0	(11.0x)	34.7	(18.2x)	22.5	(11.9x)
dsyr2k	uplo=L, trans=N	1.8					31.0	(16.8x)	18.4	(10.0x)
	uplo=L, trans=T	3.1	42.9	(13.8x)			29.6	(9.5x)	23.1	(7.4x)
	uplo=U, trans=N	3.1					30.2	(9.9x)	18.1	(5.9x)
	uplo=U, trans=T	3.1	43.2	(13.7x)			27.4	(8.7x)	23.0	(7.3x)

- “Ref”: reference Fortran code (non-parallelized, non-optimized)
- Xeon Gold 6230 (20 cores, Cascade Lake) × 2, 40 threads, gcc/gfortran 11.3.0 -march=native
- Best result among the 10 generated codes
- Blank entries indicate that no working code was generated

Optimized Code / Summary

- Performance optimization is very poor...
- OpenMP is applied, but not always with the optimal strategy
- SIMD is often applied with `#pragma omp simd` or AVX2/AVX-512 intrinsics
 - Code branching is implemented with macros (`__AVX__`, `__AVX512F__`)
- Block size is often arbitrarily set to 64 (some cases use 128 or 256)
- Difficult to generate sufficiently optimized code in one shot from a single simple prompt
- GEMM is not necessarily easy

Auto-Tuning System Using Local LLMs on Consumer-Grade PCs

Why Local LLMs? – Issues with Commercial Services

- **High costs:** Cloud-based API usage fees or subscription charges
- **Black box:** Closed source, making custom improvements and verification difficult for R&D
- **Security concerns:** Code and research data processed on external servers

Objective 1: Building a Lightweight System **Runnable on Consumer-Grade Local PCs**

- High-end open-source models are available but require large-scale computing resources
- For research purpose: to examine techniques to compensate for limited model performance

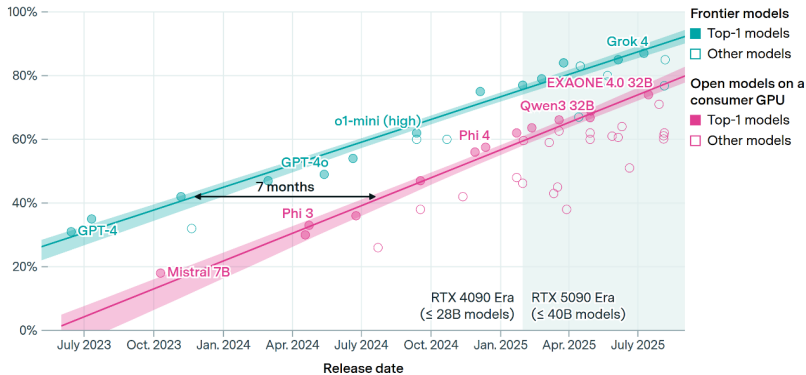
Objective 2: Specializing in Performance Optimization of HPC Code

- HPC code development is more challenging than general code development
- Performance optimization is not just about “making it work”
- An iterative code improvement process (implement → evaluate → modify) is required

Expectations for Local LLMs Runnable on Consumer-Grade PCs

Models that fit on a single consumer GPU trail the absolute frontier by less  EPOCH AI than a year.

GPQA-Diamond accuracy



CC-BY

epoch.ai

3)

³⁾Source: <https://epoch.ai/data-insights/consumer-gpu-model-gap>

Our Prototype System

Simple Automatic Code Optimization System for Local PCs

- This is more like a toy, still just practice ...
- **Multi-agent system** to compensate for limited model performance
- **Automatic iterative prompt generation** for autonomous optimization (non-stop without human intervention)
- **Using gpt-oss-120b** (for now)

OpenAI gpt-oss-120b⁴⁾ (August 2025)

- Open model with approximately 120 billion parameters (about 80 GB needed)
- **Comparable performance to o4-mini** (recall the results of BLAS generation)
- Can run on AMD Ryzen AI Max+ 395 with 128GB unified memory⁵⁾
- PC price: approximately \$2,000 → Can be considered **“runnable on consumer-grade PCs”**

⁴⁾<https://github.com/openai/gpt-oss>

⁵⁾<https://www.amd.com/en/products/processors/laptop/ryzen/ai-300-series/amd-ryzen-ai-max-plus-395.html>

System Workflow: Iterative Optimization Phase

- Initial prompt: “Optimize this code for the target processor” (actually more complicated)
- 1. **Programmer (PG) 1 – PG 5: Code generation**
 - ▶ Each PG independently creates code
- 2. **Testing and benchmarking each code**
 - ▶ Generated code is compiled and verified/evaluated using test programs
- 3. **Debugger (DG): Debugging**
 - ▶ DG analyzes the cause when errors occur and PG modifies the code (up to 2 retries)
- 4. **PG 1 – PG 5: Code analysis**
 - ▶ Analyzes which optimizations contributed to performance improvement or degradation
- 5. **Project Manager (PM): Optimization strategy planning for next iteration**
 - ▶ Collects analysis results from PGs and **generates different prompts for each PG**
 - ▶ Adopts the fastest code as the base for the next generation
- Repeat from step 1 until maximum iterations or time limit is reached

System Specification

- The user does not give any target-code-specific instructions to the system – any prompts are given; the system automatically generates iterative prompts from the initial prompt to proceed with optimization
- Optimization target is one file (computation kernel in our evaluation)
- Makefile, compiler options, execution way cannot be changed
- Benchmark and verification program is given by user (system does not edit this)
- Using gpt-oss-120b (temperature = 1.5 for high randomness) via Cerebras inference API

Performance Evaluation Environment

Evaluation Environment

- Xeon Gold 6230 (20 cores) $\times$ 2 (2688 GFlops/s in FP64, 281.5 GB/s)
- gcc 11.3.0, -O3 -march=native -fopenmp -lm -lpthread (system cannot change)
- MKL links and paths configured for verification and performance comparison (practically MKL can be used in the optimization code)

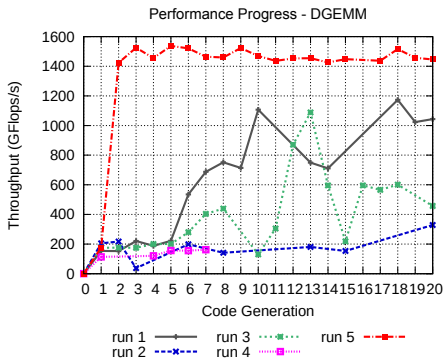
Conditions

- Strings that could reveal code content are not passed to the LLM
- Use of MKL is prohibited in optimization (but violated sometimes...)

Test 1: DGEMM

Input Code

```
1 void exp1_opt(int N,  
2             double *A,  
3             double *B,  
4             double *C) {  
5     for (int i = 0; i < N; i++) {  
6         for (int j = 0; j < N; j++) {  
7             double sum = 0.0;  
8             for (int k = 0; k < N; k++) {  
9                 sum += A[i*N+k] * B[k*N+j];  
10            }  
11            C[i*N+j] = sum;  
12        }  
13    }  
14 }
```



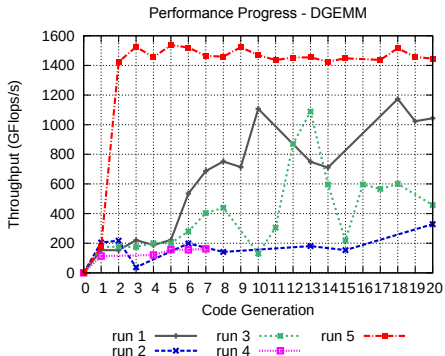
Attention!: how to read the figure

- 5 experiments (run 1 – 5) were performed with the same input and under the same conditions, due to random nature of LLM outputs
- The horizontal axis represents the number of iterations (code generations)
- Each iteration generates 5 codes, and the performance of the fastest code is plotted

Test 1: DGEMM

Input Code

```
1 void exp1_opt(int N,  
2             double *A,  
3             double *B,  
4             double *C) {  
5     for (int i = 0; i < N; i++) {  
6         for (int j = 0; j < N; j++) {  
7             double sum = 0.0;  
8             for (int k = 0; k < N; k++) {  
9                 sum += A[i*N+k] * B[k*N+j];  
10            }  
11            C[i*N+j] = sum;  
12        }  
13    }  
14 }
```



- Performance evaluated with $n = 1024$ square matrices (small for target)
- Fastest: 1537.06 GFlops/s in run 5, but this **called MKL's dgemm**
→ LLM recognized this code as a mat-mul (but use of BLAS was prohibited in prompt)
- Full-scratch fastest: **1173.42 GFlops/s** at Gen. 18 of run 1
- Large variation in results across runs

Fastest Generated DGEMM Code (1/3)

```
1 static inline void pack_B_panel(const double *B_src,
2 double *B_pack, int N, int kp, int jp, int k_max,
3 int blk_width) {
4     for (int k = kp; k < k_max; ++k) {
5         const double *src = &B_src[k * N + jp];
6         double *dst = &B_pack[(k - kp) * blk_width];
7         int j = 0;
8         _mm_prefetch((const char *)src, _MM_HINT_T0);
9         for (; j + 8 <= blk_width; j += 8) {
10             __m512d v = __mm512_loadu_pd(&src[j]);
11             __mm512_storeu_pd(&dst[j], v);
12         }
13         for (; j < blk_width; ++j) { dst[j] = src[j]; }
14     }
15     _mm_sfence();
16 }
17 void expl_opt(int N, double *restrict A, double *
18 restrict B, double *restrict C) {
19     const int BLOCK = 64;
20     omp_set_num_threads(omp_get_max_threads());
21     #pragma omp parallel for collapse(2) schedule(static)
22     for (int ii = 0; ii < N; ii += BLOCK) {
23         for (int jj = 0; jj < N; jj += BLOCK) {
24             int i_max = (ii + BLOCK > N) ? N : ii + BLOCK;
25             int j_max = (jj + BLOCK > N) ? N : jj + BLOCK;
26             int blk_width = j_max - jj;
27             int vec_width = blk_width & ~31;
28             double C_local[BLOCK * BLOCK] __attribute__((
29                 aligned(64)));
30             double *C_local_aligned = (double *)
31                 __builtin_assume_aligned(C_local, 64);
32             for (int i = 0; i < BLOCK * BLOCK; ++i)
33                 C_local_aligned[i] = 0.0;
34             for (int kk = 0; kk < N; kk += BLOCK) {
35                 int k_max = (kk + BLOCK > N) ? N : kk +
36                     BLOCK;
37                 double B_pack_local[BLOCK * BLOCK]
38                     __attribute__((aligned(64)));
39                 pack_B_panel(B, B_pack_local, N, kk, jj,
40                             k_max, blk_width);
41                 for (int i = ii; i + 3 < i_max; i += 4) {
42                     const double *a_row0 = &A[(i + 0) * N +
43                         kk];
```

```
44 const double *a_row1 = &A[(i + 1) * N +
45     kk];
46 const double *a_row2 = &A[(i + 2) * N +
47     kk];
48 const double *a_row3 = &A[(i + 3) * N +
49     kk];
50 double *c_loc0 = &C_local_aligned[(i - ii
51     ) * BLOCK];
52 double *c_loc1 = &C_local_aligned[(i - ii
53     + 1) * BLOCK];
54 double *c_loc2 = &C_local_aligned[(i - ii
55     + 2) * BLOCK];
56 double *c_loc3 = &C_local_aligned[(i - ii
57     + 3) * BLOCK];
58 for (int j = 0; j < vec_width; j += 16) {
59     __m512d c00 = __mm512_loadu_pd(&c_loc0[
60         j + 0]);
61     __m512d c01 = __mm512_loadu_pd(&c_loc0[
62         j + 8]);
63     __m512d c10 = __mm512_loadu_pd(&c_loc1[
64         j + 0]);
65     __m512d c11 = __mm512_loadu_pd(&c_loc1[
66         j + 8]);
67     __m512d c20 = __mm512_loadu_pd(&c_loc2[
68         j + 0]);
69     __m512d c21 = __mm512_loadu_pd(&c_loc2[
70         j + 8]);
71     __m512d c30 = __mm512_loadu_pd(&c_loc3[
72         j + 0]);
73     __m512d c31 = __mm512_loadu_pd(&c_loc3[
74         j + 8]);
75     for (int k = kk; k < k_max; ++k) {
76         const int k_off = k - kk;
77         __m512d b0 = __mm512_load_pd(&
78             B_pack_local[k_off * blk_width
79                 + j + 0]);
80         __m512d b1 = __mm512_load_pd(&
81             B_pack_local[k_off * blk_width
82                 + j + 8]);
83         __m512d a0 = __mm512_set1_pd(a_row0[
84             k_off]);
85         __m512d a1 = __mm512_set1_pd(a_row1[
86             k_off]);
```

Fastest Generated DGEMM Code (2/3)

```
56     __m512d a2 = _mm512_set1_pd(a_row2[
57         k_off]);
58     __m512d a3 = _mm512_set1_pd(a_row3[
59         k_off]);
60     c00 = _mm512_fmadd_pd(a0, b0, c00);
61     c01 = _mm512_fmadd_pd(a0, b1,
62         c01);
63     c10 = _mm512_fmadd_pd(a1, b0, c10);
64     c11 = _mm512_fmadd_pd(a1, b1,
65         c11);
66     c20 = _mm512_fmadd_pd(a2, b0, c20);
67     c21 = _mm512_fmadd_pd(a2, b1,
68         c21);
69     c30 = _mm512_fmadd_pd(a3, b0, c30);
70     c31 = _mm512_fmadd_pd(a3, b1,
71         c31);
72 }
73 _mm512_storeu_pd(&c_loc0[j + 0], c00);
74 _mm512_storeu_pd(&c_loc0[j + 8],
75     c01);
76 _mm512_storeu_pd(&c_loc1[j + 0], c10);
77 _mm512_storeu_pd(&c_loc1[j + 8],
78     c11);
79 _mm512_storeu_pd(&c_loc2[j + 0], c20);
80 _mm512_storeu_pd(&c_loc2[j + 8],
81     c21);
82 _mm512_storeu_pd(&c_loc3[j + 0], c30);
83 _mm512_storeu_pd(&c_loc3[j + 8],
84     c31);
85 }
86 for (int j = vec_width; j < blk_width; ++
87     j) {
88     double s0 = c_loc0[j]; double s1 =
89     c_loc1[j];
90     double s2 = c_loc2[j]; double s3 =
91     c_loc3[j];
92     for (int k = kk; k < k_max; ++k) {
93         const int k_off = k - kk;
94         double b = B_pack_local[k_off *
95             blk_width + j];
96         s0 += a_row0[k_off] * b; s1 +=
97         a_row1[k_off] * b;
```

```
75         s2 += a_row2[k_off] * b; s3 +=
76         a_row3[k_off] * b;
77     }
78     c_loc0[j] = s0; c_loc1[j] = s1; c_loc2
79     [j] = s2; c_loc3[j] = s3;
80 }
81 for (int i = i_max - ((i_max - ii) % 4); i <
82     i_max; ++i) {
83     const double *a_ptr = &A[i * N + kk];
84     double *c_loc = &C_local_aligned[(i - ii)
85         * BLOCK];
86     for (int j = 0; j < vec_width; j += 8) {
87         __m512d c_vec = _mm512_loadu_pd(&c_loc
88             [j]);
89         for (int k = kk; k < k_max; ++k) {
90             const int k_off = k - kk;
91             __m512d b_vec = _mm512_load_pd(&
92                 B_pack_local[k_off * blk_width
93                     + j]);
94             __m512d a_vec = _mm512_set1_pd(
95                 a_ptr[k_off]);
96             c_vec = _mm512_fmadd_pd(a_vec,
97                 b_vec, c_vec);
98         }
99         _mm512_storeu_pd(&c_loc[j], c_vec);
100     }
101     for (int j = vec_width; j < blk_width; ++
102         j) {
103         double sum = c_loc[j];
104         for (int k = kk; k < k_max; ++k) {
105             const int k_off = k - kk;
106             sum += a_ptr[k_off] * B_pack_local[
107                 k_off * blk_width + j];
108         }
109         c_loc[j] = sum;
110     }
111 }
112 for (int i = ii; i < i_max; ++i) {
113     double *c_ptr = &C[i * N + jj];
114     double *c_loc = &C_local_aligned[(i - ii) *
115         BLOCK];
```

Fastest Generated DGEMM Code (3/3)

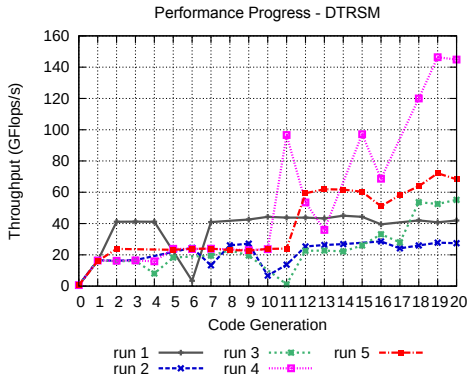
```
106         c_ptr = (double *)__builtin_assume_aligned(  
107             c_ptr, 64);  
108         c_loc = (double *)__builtin_assume_aligned(  
109             c_loc, 64);  
110         int j = 0;  
111         for (; j < vec_width; j += 8) {  
112             _mm512d tmp = _mm512_loadu_pd(&c_loc[j]);  
113             _mm512_storeu_pd(&c_ptr[j], tmp);  
114         }  
115         for (; j < blk_width; ++j) {  
116             c_ptr[j] = c_loc[j];  
117         }  
118     }  
119 }
```

*Due to space limitations, auto-inserted comments, blank lines, and header includes were removed from the actual generated code, and indentation/line breaks were partially modified.

Test 2: DTRSM

Input Code

```
1 void exp2_opt(int N,  
2     double *A, double *B,  
3     double *X, double alpha) {  
4     for (int i = 0; i < N; i++) {  
5         for (int j = 0; j < N; j++) {  
6             X[i*N+j] = alpha * B[i*N+j];  
7         }  
8     }  
9     for (int j = 0; j < N; j++) {  
10        for (int i = 0; i < N; i++) {  
11            for (int k = 0; k < i; k++) {  
12                X[i*N+j] -= A[i*N+k]*X[k*N+j];  
13            }  
14            X[i*N+j] /= A[i*N+i];  
15        }  
16    }  
17 }
```



- Solving equations with triangular matrix coefficients (lower triangular, $n = 1024$); more complex data dependencies than DGEMM; high optimization difficulty
- Fastest: **146.33 GFlops/s** at Gen.19 of run 4 (MKL: 305.74 GFlops/s, **48% of MKL**)
- Not observed in this experiment, but cases using MKL were seen during test runs

Fastest Generated DTRSM Code (part)

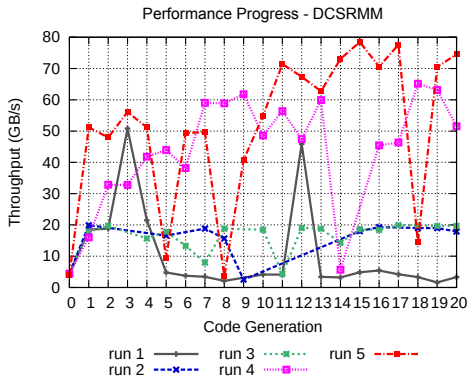
```
1  #define COL_BLOCK 8
2  void exp2_opt(int N, double * restrict A, double *
   restrict B,
3      double * restrict X, double alpha) {
4      double inv_diag[N] __attribute__((aligned(64)));
5      double *trans_X = (double*)aligned_alloc(64, (size_t)
   N * N * sizeof(double));
6      #pragma omp parallel for schedule(static)
7      for (int i = 0; i < N; ++i) {
8          double *xrow = X + i * N;
9          double *brow = B + i * N;
10         int j = 0;
11         __m512d valpha = __m512d_set1_pd(alpha);
12         for (; j <= N-8; j += 8) {
13             __m512d v = __m512d_loadu_pd(brow + j);
14             v = __m512d_mul_pd(v, valpha);
15             __m512d_storeu_pd(xrow + j, v);
16         }
17         for (; j < N; ++j) xrow[j] = alpha * brow[j];
18         inv_diag[i] = 1.0 / A[i * N + i];
19     }
20     const int BS = 64;
21     #pragma omp parallel for collapse(2) schedule(static)
22     for (int i0 = 0; i0 < N; i0 += BS) {
23         for (int j0 = 0; j0 < N; j0 += BS) {
24             int i_max = i0 + BS < N ? i0 + BS : N;
25             int j_max = j0 + BS < N ? j0 + BS : N;
26             for (int i = i0; i < i_max; ++i) {
27                 double *xrow = X + i * N;
28                 for (int j = j0; j < j_max; ++j) {
29                     trans_X[j * N + i] = xrow[j];
30                 }
31             }
32         }
33     }
34     #pragma omp parallel
35     {
36         #pragma omp for schedule(runtime) nowait
37         for (int jb = 0; jb < N; jb += COL_BLOCK) {
38             int j_end = (jb + COL_BLOCK < N) ? jb +
   COL_BLOCK : N;
39             #pragma omp task depend(inout: trans_X[jb * N])
40             {
41                 for (int j = jb; j < j_end; ++j) {
```

```
42         double *xcol = trans_X + j * N;
43         #pragma omp simd
44         for (int i = 0; i < N; ++i) {
45             const double *arow = A + i * N;
46             __m512d vsum = __m512d_setzero_pd();
47             int limit = i & ~7;
48             if (limit >= 8) {
49                 __mm_prefetch(arow + 8, _MM_HINT_T0);
50                 __mm_prefetch(xcol + 8, _MM_HINT_T0);
51                 __m512d a_prev = __m512d_loadu_pd(
   arow);
52                 __m512d x_prev = __m512d_loadu_pd(
   xcol);
53                 int k = 8;
54                 for (; k < limit; k += 8) {
55                     __mm_prefetch(arow + k + 8,
   _MM_HINT_T0);
56                     __mm_prefetch(xcol + k + 8,
   _MM_HINT_T0);
57                     __m512d a_cur = __m512d_loadu_pd(
   arow + k);
58                     __m512d x_cur = __m512d_loadu_pd(
   xcol + k);
59                     vsum = __m512d_fmadd_pd(a_prev,
   x_prev, vsum);
60                     a_prev = a_cur; x_prev = x_cur;
61                 }
62                 vsum = __m512d_fmadd_pd(a_prev,
   x_prev, vsum);
63             }
64             int rem = i - limit;
65             if (rem) {
66                 __mm_prefetch(arow + limit + 8,
   _MM_HINT_T0);
67                 __mm_prefetch(xcol + limit + 8,
   _MM_HINT_T0);
68                 __mmask8 m = (1U << rem) - 1U;
69                 __m512d a = __m512d_maskz_loadu_pd(m,
   arow + limit);
70                 __m512d x = __m512d_maskz_loadu_pd(m,
   xcol + limit);
71                 vsum = __m512d_fmadd_pd(a, x, vsum);
```

Test 3: DCSRMM

Input Code

```
1 void exp3_opt(int m, int k, int n,  
2     const int *row_ptr,  
3     const int *col_indices,  
4     const double *values,  
5     const double *B, double *C) {  
6     memset(C, 0, m*n*sizeof(double));  
7     for (int i = 0; i < m; i++) {  
8         int row_start = row_ptr[i];  
9         int row_end = row_ptr[i + 1];  
10        for (int idx=row_start;  
11            idx<row_end; idx++) {  
12            int col = col_indices[idx];  
13            double val = values[idx];  
14            for (int j = 0; j < n; j++) {  
15                C[i*n+j] += val*B[col*n+j];  
16            }  
17        }  
18    }  
19 }
```



- Sparse-dense mat-mul in CSR format (gridgena.mtx: $n = 48962$, $nnz = 512084$, RHS width $m = 128$); irregular memory access; depends on sparse structure
- Fastest: **78.45 GB/s** at Gen. 15 of run 5 (MKL: 86.88 GB/s, **90% of MKL**)
- Attempted FP16 (failed), CSC and ELLPACK format changes (failed) → sparse matrix structure was recognized

Fastest Generated DCSRMM Code (part)

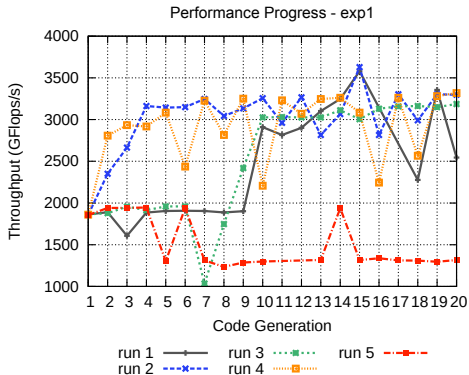
```
1  #define UNROLL_FACTOR 4
2  #define VEC_WIDTH 8
3  #define UNROLL_STEP (UNROLL_FACTOR * VEC_WIDTH)
4  void exp3_opt(int m, int k, int n,
5               const int *__restrict row_ptr,
6               const int *__restrict col_indices,
7               const double *__restrict values,
8               const double *__restrict B,
9               double *__restrict C) {
10     int block_n = 64;
11     const char *env_n = getenv("BLOCK_N");
12     if (env_n) { int v = atoi(env_n); if (v > 0) block_n
13                 = v; }
14     int block_m = 128;
15     const char *env_m = getenv("BLOCK_M");
16     if (env_m) { int v = atoi(env_m); if (v > 0) block_m
17                 = v; }
18     int block_shift = 0;
19     while ((1 << block_shift) < block_n) ++block_shift;
20     #pragma omp parallel for schedule(static)
21     for (int i0 = 0; i0 < m; i0 += block_m) {
22         int i1 = i0 + block_m;
23         if (i1 > m) i1 = m;
24         for (int i = i0; i < i1; ++i) {
25             double *c_row = C + (size_t)i * n;
26             int row_start = row_ptr[i];
27             int row_end = row_ptr[i + 1];
28             for (int idx = row_start; idx < row_end; ++idx) {
29                 int col = col_indices[idx];
30                 double val = values[idx];
31                 __m512d val_vec = __m512d_set1_pd(val);
32                 int block_id = col >> block_shift;
33                 int col_off = col & (block_n - 1);
34                 const double *b_ptr = B +
35                     ((size_t)block_id * block_n + (size_t)
36                     col_off) * n;
37                 int j = 0;
38                 int unroll_limit = n - (n % UNROLL_STEP);
39                 for (; j < unroll_limit; j += UNROLL_STEP) {
40                     __m512d c0 = __m512d_loadu_pd(c_row + j);
41                     __m512d b0 = __m512d_loadu_pd(b_ptr + j);
42                     __m512d c1 = __m512d_loadu_pd(c_row + j +
43                     VEC_WIDTH);
```

```
44                     __m512d b1 = __m512d_loadu_pd(b_ptr + j +
45                     VEC_WIDTH);
46                     __m512d c2 = __m512d_loadu_pd(c_row + j +
47                     2 * VEC_WIDTH);
48                     __m512d b2 = __m512d_loadu_pd(b_ptr + j +
49                     2 * VEC_WIDTH);
50                     __m512d c3 = __m512d_loadu_pd(c_row + j +
51                     3 * VEC_WIDTH);
52                     __m512d b3 = __m512d_loadu_pd(b_ptr + j +
53                     3 * VEC_WIDTH);
54                     c0 = __m512d_fmadd_pd(val_vec, b0, c0);
55                     c1 = __m512d_fmadd_pd(val_vec, b1, c1);
56                     c2 = __m512d_fmadd_pd(val_vec, b2, c2);
57                     c3 = __m512d_fmadd_pd(val_vec, b3, c3);
58                     __m512d_storeu_pd(c_row + j, c0);
59                     __m512d_storeu_pd(c_row + j + VEC_WIDTH,
60                     c1);
61                     __m512d_storeu_pd(c_row + j + 2 *
62                     VEC_WIDTH, c2);
63                     __m512d_storeu_pd(c_row + j + 3 *
64                     VEC_WIDTH, c3);
65                 }
66                 int tail_rem = n - j;
67                 int tail_vec = tail_rem & ~7;
68                 for (int t = 0; t < tail_vec; t += VEC_WIDTH) {
69                     __m512d c_vec = __m512d_loadu_pd(c_row + j
70                     + t);
71                     __m512d b_vec = __m512d_loadu_pd(b_ptr + j
72                     + t);
73                     c_vec = __m512d_fmadd_pd(val_vec, b_vec,
74                     c_vec);
75                     __m512d_storeu_pd(c_row + j + t, c_vec);
76                 }
77                 int tail_elem = tail_rem - tail_vec;
78                 if (tail_elem) {
79                     __mmask8 mask = (1U << tail_elem) - 1;
80                     __m512d c_mask = __m512d_maskz_loadu_pd(
81                     mask, c_row + j + tail_vec);
82                     __m512d b_mask = __m512d_maskz_loadu_pd(
83                     mask, b_ptr + j + tail_vec);
84                     c_mask = __m512d_fmadd_pd(val_vec, b_mask,
85                     c_mask);
```

Test 1: DGEMM (CUDA)

Input Code

```
1  __global__ void exp1_kernel(int N, const
2      double *A, const double *B, double *C) {
3      int row = blockIdx.y * blockDim.y +
4          threadIdx.y;
5      int col = blockIdx.x * blockDim.x +
6          threadIdx.x;
7      if (row < N && col < N) {
8          double sum = 0.0;
9          for (int k = 0; k < N; k++) {
10             sum += A[row * N + k] * B[k * N +
11                 col];
12         }
13         C[row * N + col] = sum;
14     }
15 }
16
17 extern "C" void exp1_opt_kernel(int N, const
18     double *d_A, const double *d_B, double *
19     d_C) {
20     dim3 blockDim(16, 16);
21     dim3 gridDim((N + blockDim.x - 1) /
22         blockDim.x, (N + blockDim.y - 1) /
23         blockDim.y);
24     exp1_kernel<<<gridDim, blockDim>>>(N, d_A
25         , d_B, d_C);
26     cudaDeviceSynchronize();
27 }
```



Target: Tesla V100 (one GPU, 7000 GFlops/s in FP64)

Best code: 3628.29 GFlops/s

cuBLAS: 6168.09 GFlops/s

Fastest Generated DGEMM Code (CUDA) (1/3)

```
1  #include <cuda_runtime.h>
2  #include <stdlib.h>
3
4  extern "C" cudaError_t cudaGetDeviceProperties_v2(
5      cudaDeviceProp *prop, int device) {
6      return cudaGetDeviceProperties(prop, device);
7  }
8
9  extern "C" long __isoc23_strtol(const char *nptr, char
10     **endptr, int base) {
11     return strtol(nptr, endptr, base);
12 }
13
14 #define TILE 32 // ×3232 タイル
15 #define K_UNROLL 32
16 // 手動アンロールマクロ
17 #define UNROLL_K(iter)
18
19     do {
20         \
21         double a0 = As[calc][threadIdx.y*2 + 0][iter];
22         double a1 = As[calc][threadIdx.y*2 + 1][iter];
23         double b0 = Bs[calc][iter][threadIdx.x*2 + 0];
24         double b1 = Bs[calc][iter][threadIdx.x*2 + 1];
25
26         regC[0][0] += a0 * b0;
27         regC[0][1] += a0 * b1;
28         regC[1][0] += a1 * b0;
29         regC[1][1] += a1 * b1;
30     } while (0)
31
32 __global__ void exp1_kernel(int N, const double *
33     __restrict__ A,
34     const double *__restrict__ B
35     ,
36     double *__restrict__ C) {
```

```
31     double regC[2][2] = {{0.0, 0.0},
32                           {0.0, 0.0}}; //
33                                     // ×22 出力
34
35     __shared__ double As[2][TILE][TILE];
36     __shared__ double Bs[2][TILE][TILE];
37
38     int numTiles = (N + TILE - 1) / TILE;
39
40     int baseRow = blockIdx.y * TILE + threadIdx.y * 2;
41     // 2 行相当
42     int baseCol = blockIdx.x * TILE + threadIdx.x * 2;
43     // 2 列相当
44
45     // ---- 1st タイルを load (ping = 0) ----
46     int ping = 0;
47     for (int dy = 0; dy < 2; ++dy) {
48         int aRow = baseRow + dy;
49         int aCol = 0 * TILE + threadIdx.x * 2;
50         double4 a_pair = ((double4*)(A + aRow * N +
51                                     aCol));
52         double a0 = a_pair.x;
53         double a1 = a_pair.y;
54         As[ping][threadIdx.y * 2 + dy][threadIdx.x * 2]
55             = a0;
56         As[ping][threadIdx.y * 2 + dy][threadIdx.x * 2 +
57             1] = a1;
58     }
59     for (int dy = 0; dy < 2; ++dy) {
60         int bRow = 0 * TILE + threadIdx.y * 2 + dy;
61         int bCol = baseCol;
62         double4 b_pair = ((double4*)(B + bRow * N +
63                                     bCol));
64         double b0 = b_pair.x;
65         double b1 = b_pair.y;
66         Bs[ping][threadIdx.y * 2 + dy][threadIdx.x * 2]
67             = b0;
68         Bs[ping][threadIdx.y * 2 + dy][threadIdx.x * 2 +
69             1] = b1;
70     }
71 }
72
73 __syncthreads();
74
75     // 初期タイルロード
76     待ち
77     for (int t = 0; t < numTiles; ++t) {
```

Fastest Generated DGEMM Code (CUDA) (2/3)

```
65     int calc = ping;                                     // 現在計算に使用する
66     バッファ
67     // ---- 計算 (calc バッファ) ----
68     UNROLL_K(0); UNROLL_K(1); UNROLL_K(2);
69     UNROLL_K(3); UNROLL_K(4); UNROLL_K(5); UNROLL_K(6);
70     UNROLL_K(7); UNROLL_K(8); UNROLL_K(9); UNROLL_K(10);
71     UNROLL_K(11); UNROLL_K(12); UNROLL_K(13); UNROLL_K(14);
72     UNROLL_K(15); UNROLL_K(16); UNROLL_K(17); UNROLL_K(18);
73     UNROLL_K(19); UNROLL_K(20); UNROLL_K(21); UNROLL_K(22);
74     UNROLL_K(23); UNROLL_K(24); UNROLL_K(25); UNROLL_K(26);
75     UNROLL_K(27); UNROLL_K(28); UNROLL_K(29); UNROLL_K(30);
76     UNROLL_K(31);
77     // 次のタイルを prefetch (ping バッファ) すでに最後の場合はスキップ
78     ping = 1 - ping;                                     // バッファを反転
79
80     if (t + 1 < numTiles) {
81         for (int dy = 0; dy < 2; ++dy) {
82             int aRow = baseRow + dy;
83             int aCol = (t + 1) * TILE + threadIdx.x
84                 * 2;
85             double4 a_pair = *((double4*)(A + aRow *
86                 N + aCol));
87             double a0 = a_pair.x;
88             double a1 = a_pair.y;
89             As[ping][threadIdx.y * 2 + dy][threadIdx
90                 .x * 2] = a0;
91             As[ping][threadIdx.y * 2 + dy][threadIdx
92                 .x * 2 + 1] = a1;
93         }
94         for (int dy = 0; dy < 2; ++dy) {
95             int bRow = (t + 1) * TILE + threadIdx.y
96                 * 2 + dy;
97             int bCol = baseCol;
```

```
92     double4 b_pair = *((double4*)(B + bRow *
93         N + bCol));
94     double b0 = b_pair.x;
95     double b1 = b_pair.y;
96     Bs[ping][threadIdx.y * 2 + dy][threadIdx
97         .x * 2] = b0;
98     Bs[ping][threadIdx.y * 2 + dy][threadIdx
99         .x * 2 + 1] = b1;
100     }
101     __syncthreads();                                     // 次計算に備える
102
103     }
104     // ---- 書き戻し ----
105     C[baseRow * N + baseCol] = regC
106     [0][0];
107     C[baseRow * N + baseCol + 1] = regC
108     [0][1];
109     C[(baseRow + 1) * N + baseCol] = regC
110     [1][0];
111     C[(baseRow + 1) * N + baseCol + 1] = regC
112     [1][1];
113 }
114 extern "C" void exp1_opt_kernel(int N, const double *d_A
115     ,
116     const double *d_B,
117     double *d_C) {
118     if (N % 8 != 0) {
119         int Npad = ((N + 7) / 8) * 8;
120         size_t pitch_pad = (size_t)Npad * sizeof(double)
121             ;
122         size_t pitch_orig = (size_t)N * sizeof(double);
123         double *d_Apad = nullptr, *d_Bpad = nullptr, *
124             d_Cpad = nullptr;
125         cudaMalloc(&d_Apad, pitch_pad * Npad);
126         cudaMalloc(&d_Bpad, pitch_pad * Npad);
127         cudaMalloc(&d_Cpad, pitch_pad * Npad);
128         cudaMemcpy(d_Apad, 0, pitch_pad * Npad);
129         cudaMemcpy(d_Bpad, 0, pitch_pad * Npad);
130         cudaMemcpy(d_Cpad, 0, pitch_pad * Npad);
```

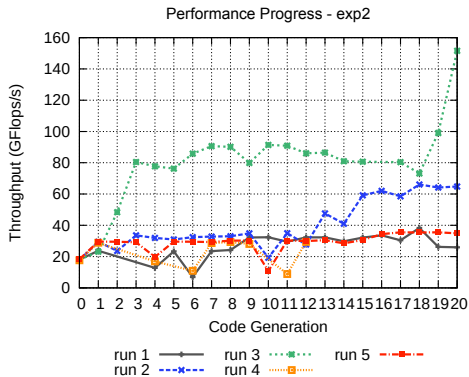
Fastest Generated DGEMM Code (CUDA) (2/3)

```
122     cudaMemcpy2D(d_Apad, pitch_pad, d_A, pitch_orig,
123                 pitch_orig, N,
124                 cudaMemcpyDeviceToDevice);
125     cudaMemcpy2D(d_Bpad, pitch_pad, d_B, pitch_orig,
126                 pitch_orig, N,
127                 cudaMemcpyDeviceToDevice);
128
129     dim3 blockDim(TILE / 2, TILE / 2);
130     dim3 gridDim((Npad + TILE - 1) / TILE, (Npad +
131                 TILE - 1) / TILE);
132     exp1_kernel<<<gridDim, blockDim>>>(Npad, d_Apad,
133                 d_Bpad, d_Cpad);
134     cudaDeviceSynchronize();
135
136     cudaMemcpy2D(d_C, pitch_orig, d_Cpad, pitch_pad,
137                 pitch_orig, N,
138                 cudaMemcpyDeviceToDevice);
139
140     cudaFree(d_Apad);
141     cudaFree(d_Bpad);
142     cudaFree(d_Cpad);
143     return;
144 }
145
146 dim3 blockDim(TILE / 2, TILE / 2);
147 dim3 gridDim((N + TILE - 1) / TILE, (N + TILE - 1) /
148                 TILE);
149 exp1_kernel<<<gridDim, blockDim>>>(N, d_A, d_B, d_C);
150 cudaDeviceSynchronize();
151 }
```

```
#undef K_UNROLL
#undef UNROLL_K
```

Test 2: DTRSM (CUDA)

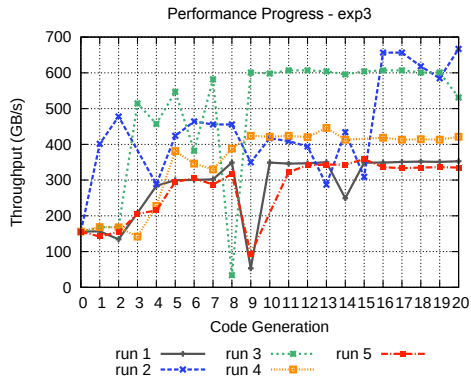
```
1  __global__ void kernel1(int N, const double *B,
2      double *X, double alpha) {
3      int idx = blockIdx.x * blockDim.x + threadIdx.x;
4      if (idx < N * N)
5          X[idx] = alpha * B[idx];
6  }
7  __global__ void kernel2(int N, int row, const double
8      *A, double *X) {
9      int col = blockIdx.x * blockDim.x + threadIdx.x;
10     if (col < N) {
11         double sum = 0.0;
12         for (int k = 0; k < row; k++) {
13             sum += A[row * N + k] * X[k * N + col];
14         }
15         X[row * N + col] = (X[row * N + col] - sum)
16             / A[row * N + row];
17     }
18 }
19 extern "C" void exp2_opt_kernel(int N, const double
20     *d_A, const double *d_B, double *d_X, double
21     alpha) {
22     int total = N * N;
23     int threads = 256;
24     int blocks = (total + threads - 1) / threads;
25     kernel1<<<blocks, threads>>>(N, d_B, d_X, alpha)
26     ;
27     cudaDeviceSynchronize();
28     threads = 256;
29     blocks = (N + threads - 1) / threads;
30     for (int i = 0; i < N; i++) {
31         kernel2<<<blocks, threads>>>(N, i, d_A, d_X)
32         ;
33         cudaDeviceSynchronize();
34     }
35 }
```



Best code: 151.59 GFlops/s
cuBLAS: 1278.75 GFlops/s

Test 3: DCSRMM (CUDA)

```
1  __global__ void exp3_kernel(int m, int n, const
    int *row_ptr, const int *col_indices, const
    double *values, const double *B, double *C) {
2      int row = blockIdx.x * blockDim.x + threadIdx
        .x;
3      int col = blockIdx.y * blockDim.y + threadIdx
        .y;
4      if (row < m && col < n) {
5          double sum = 0.0;
6          int row_start = row_ptr[row];
7          int row_end = row_ptr[row + 1];
8          for (int idx = row_start; idx < row_end;
                idx++) {
9              int sparse_col = col_indices[idx];
10             double val = values[idx];
11             sum += val * B[sparse_col * n + col];
12         }
13         C[row * n + col] = sum;
14     }
15 }
16 extern "C" void exp3_opt_kernel(int m, int k, int
    n, const int *d_row_ptr, const int *
    d_col_indices, const double *d_values, const
    double *d_B, double *d_C) {
17     cudaMemset(d_C, 0, m * n * sizeof(double));
18     dim3 blockDim(16, 16);
19     dim3 gridDim((m + blockDim.x - 1) / blockDim.
        x, (n + blockDim.y - 1) / blockDim.y);
20     exp3_kernel<<<gridDim, blockDim>>>(m, n,
        d_row_ptr, d_col_indices, d_values, d_B,
        d_C);
21     cudaDeviceSynchronize();
22 }
```



Best code: 666.78 GFlops/s

cuSparse: 397.58 GFlops/s

→ This is true (but possible as it is specialized for the given problem)

Summary (our sysytem)

- Our toy system with gpt-oss-120b is pretty impressive
- Matrix multiplication is easy, right? Yes, but still difficult.
- Actual more complicated code? – will be more difficult
- There are tons of ideas for improvement

Generatie AI

- A Must-Have Technology for Fugaku NEXT
- From an academic research perspective, there is concern about whether we can catch up with the rapid pace of technological evolution (and \$ cost for GPUs)
- Model performance and commercial services are evolving so rapidly – our modest efforts will soon be wasted ??? (though still meaningful as research)