

Accelerating the Multifrontal Method

Information Sciences Institute

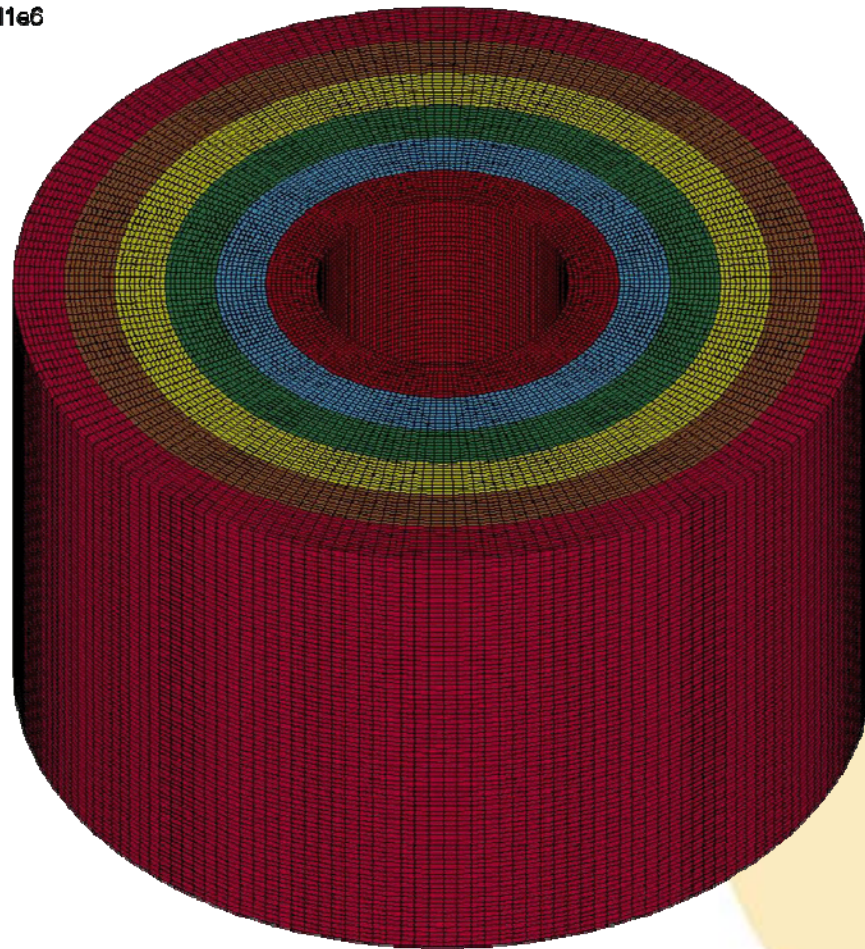


22 June 2012

Bob Lucas, Gene Wagenbreth, Dan Davis, Roger Grimes
{rflucas,genew,ddavis}@isi.edu and grimes@lstc.com

3D Finite Element Test Problem

Test Problem: cylinders cyl1e6



LS-DYNA Implicit Half Million Elements

Timing information
CPU(seconds) %CPU Clock(seconds) %Clock

Initialization	2.6513E+01	0.24	2.7598E+01	1.53
Element processing ...	8.1698E+01	0.73	7.6144E+01	4.22
Binary databases	2.0471E+00	0.02	4.8657E-01	0.03
ASCII database	1.3164E-04	0.00	1.5000E-05	0.00
Contact algorithm	2.9563E+01	0.26	5.7157E+00	0.32
Interf. ID 1	7.6828E+00	0.07	1.2124E+00	0.07
Interf. ID 2	4.8080E+00	0.04	9.7007E-01	0.05
Interf. ID 3	4.4934E+00	0.04	9.8310E-01	0.05
Interf. ID 4	5.5574E+00	0.05	1.2091E+00	0.07
Interf. ID 5	7.0113E+00	0.06	1.3380E+00	0.07
Contact entities	0.0000E+00	0.00	0.0000E+00	0.00
Rigid bodies	8.5901E-04	0.00	3.1600E-04	0.00
Implicit Nonlinear ...	8.7434E+00	0.08	6.5756E+00	0.36
Implicit Lin. Alg. ...	1.1063E+04	98.67	1.6898E+03	93.55
Totals	1.1212E+04	100.00	1.8063E+03	100.00

- **All of LS-DYNA**
 - **Linear Solver**
 - **Reordering (Metis) & symbolic factorization**
 - **Redistribute & permute sparse matrix**
 - **Factor, i.e., for all supernodes:**
 - **Assemble the frontal matrix**
 - **Factor the frontal matrix**
 - **Stack its Schur complement**
 - **Solve**



Time in solver (sec) Eight Nehalem Threads

grep WCT mes0000 =>

WCT: symbolic factorization = 18.952
WCT: matrix redistribution = 5.221
WCT: factorization = 795.354
WCT: numeric solve = 2.705
WCT: total imfslv_mf2 = 826.735
WCT: symbolic factorization = 16.560
WCT: matrix redistribution = 5.087
WCT: factorization = 804.471
WCT: numeric solve = 2.819
WCT: total imfslv_mf2 = 829.281



Start by accelerating factorization

- **All of LS-DYNA**
 - **Multifrontal Linear Solver**
 - **Reordering (Metis) & symbolic factorization**
 - **Redistribute & permute sparse matrix**
 - **Factor, i.e., for all supernodes**
 - **Assemble the frontal matrix**
 - **Factor the frontal matrix with accelerator**
 - **Stack its Schur complement**
 - **Solve**

Multifontal Method Toy Problem

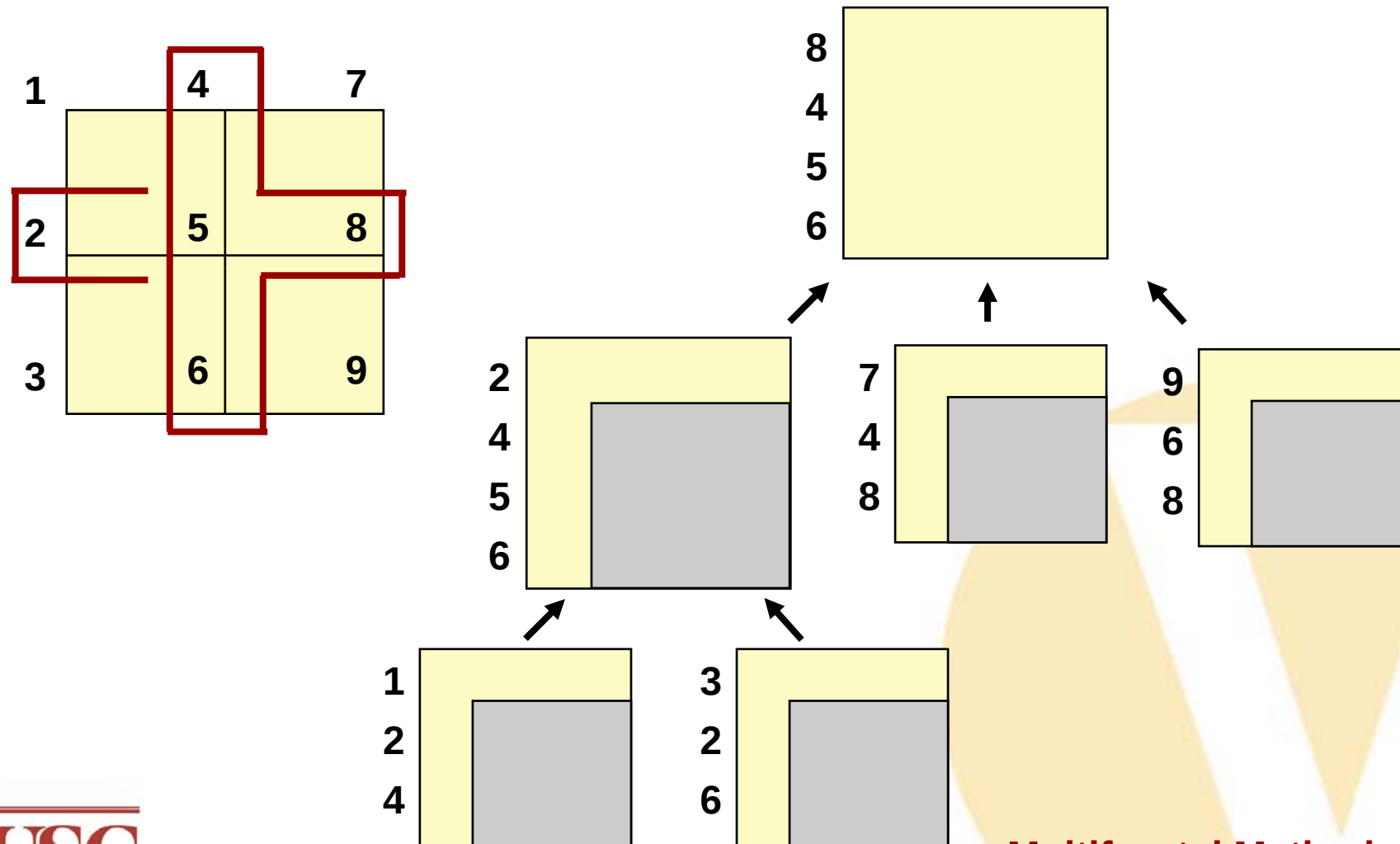
```

do 4 k = 1, 9
  do 1 i = k + 1, 9
    a(i, k) = a(i,k) / a(k,k)
1  continue
    do 3 j = k + 1, 9
      do 2 i = k + 1, 9
        a(i,j) = a(i,j) -
2          a(i,k) *
3          a(k,j)
4  continue
  continue
  continue
  continue
  
```

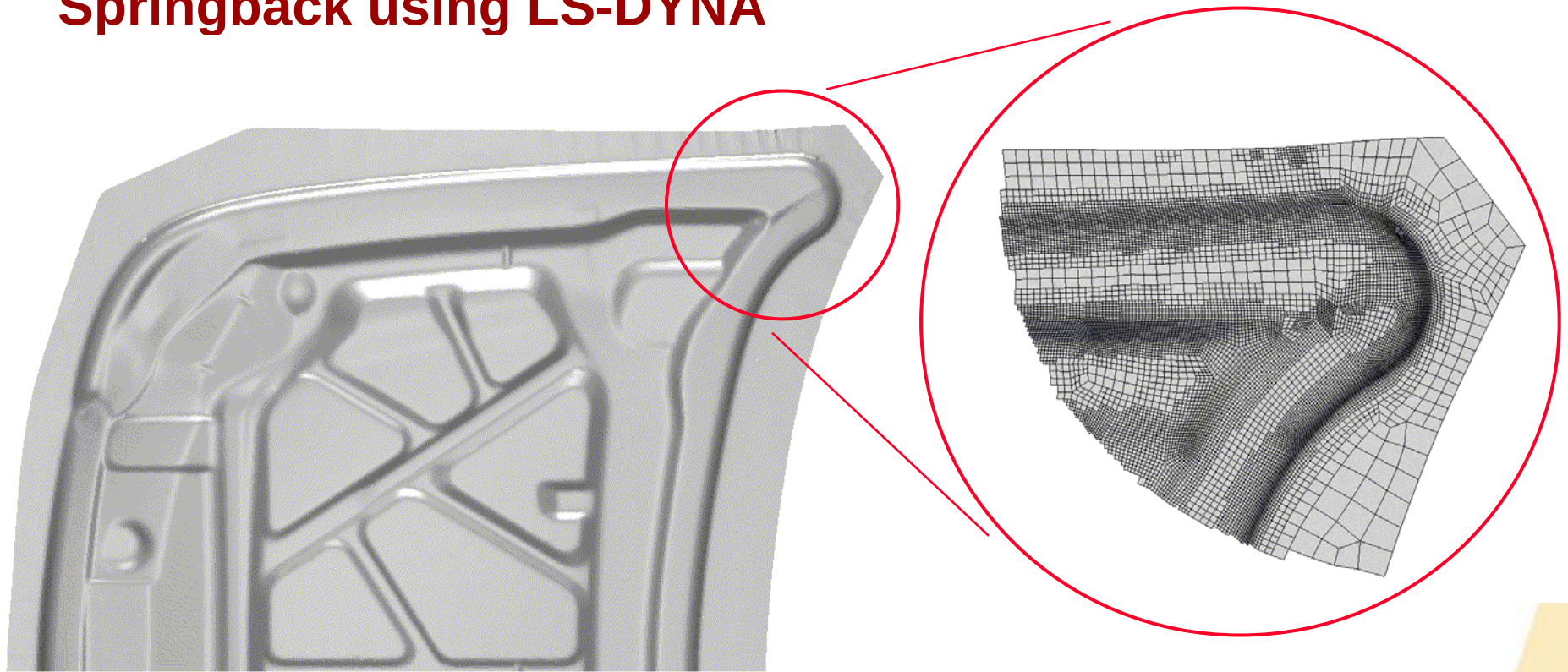
1	4	7
2	5	8
3	6	9

1	X	X		X		
3		XX				X
2	XXX			*	X	*
7		X	XX			
9			XX			X
8		XXX	*	X	*	
4	X	*	X	*	XX	*
5		X		XXXX		
6	X	*	X	*	*	XX

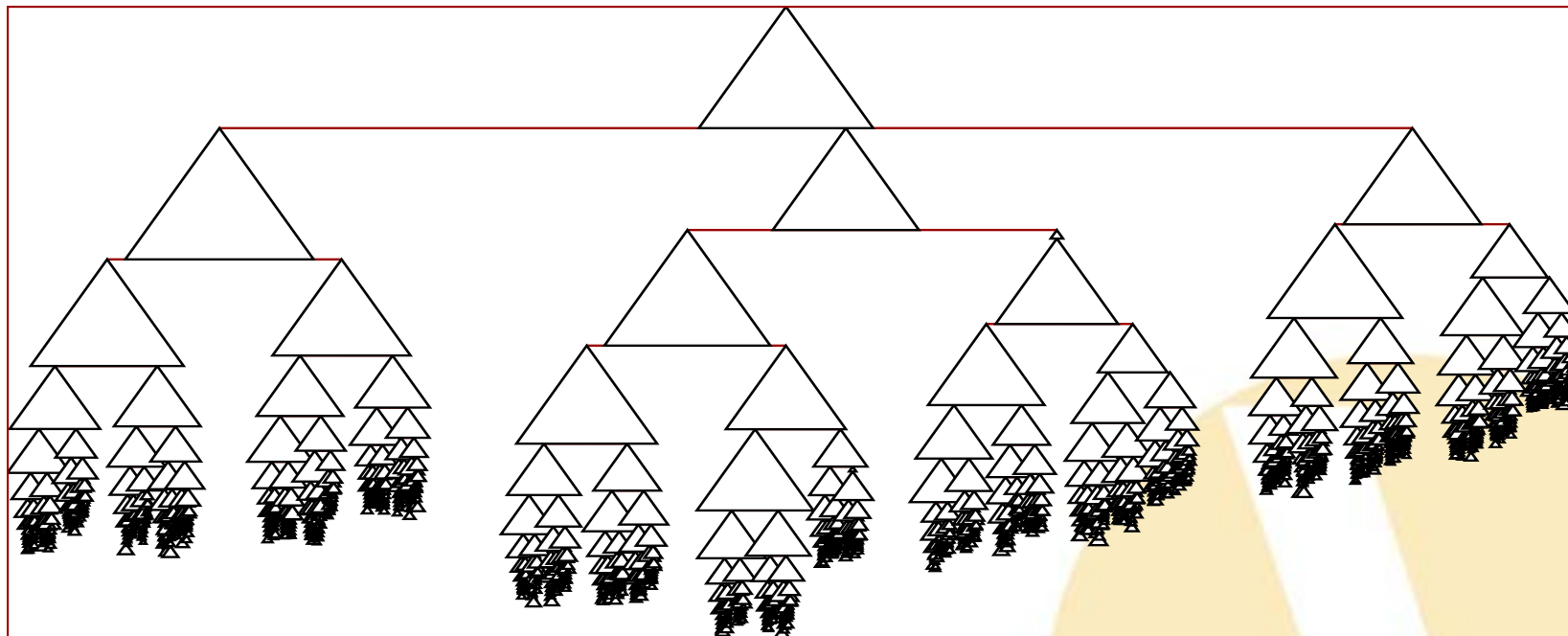
Multifrontal View of a Toy Matrix



Automotive Hood Inner Panel Springback using LS-DYNA

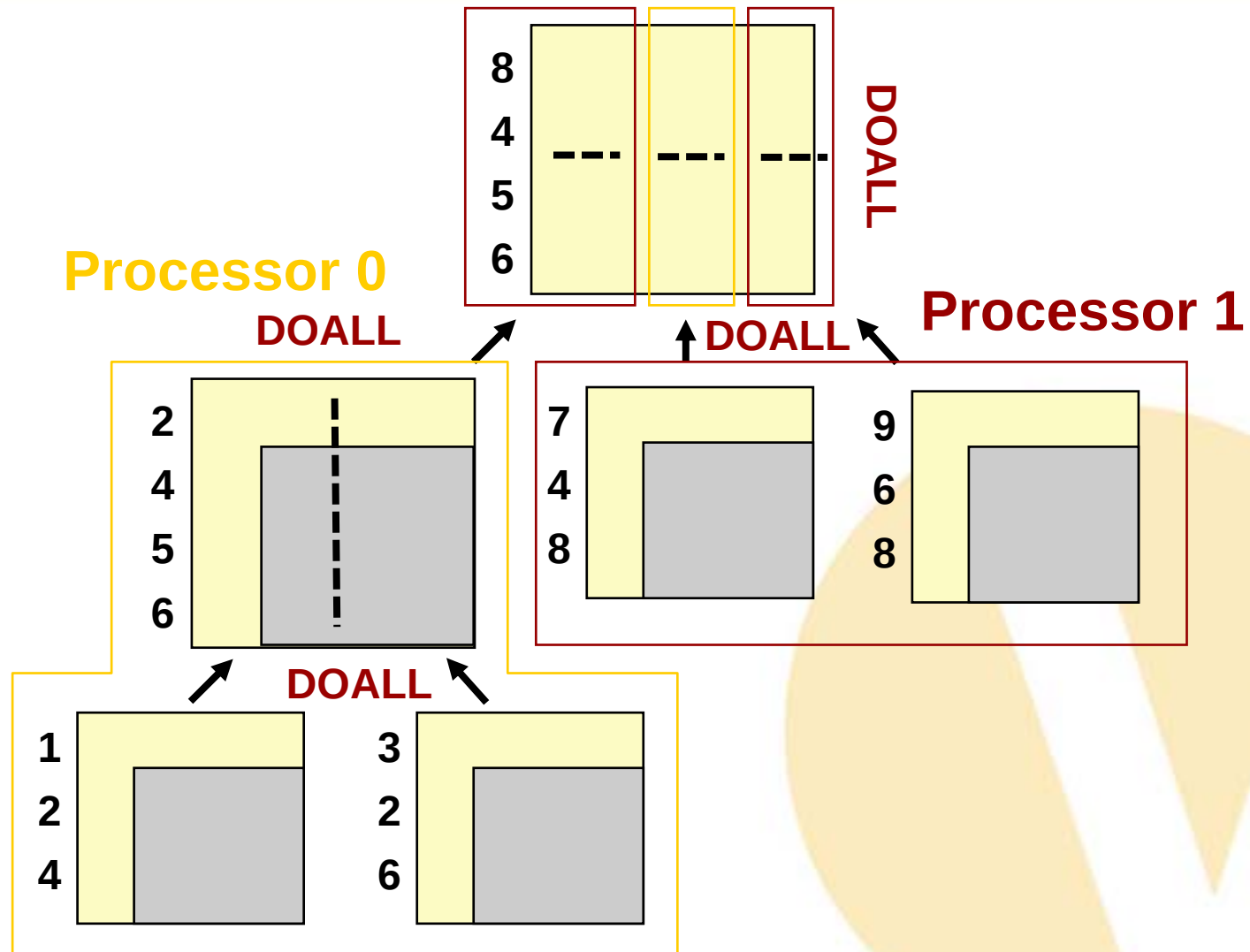


“Hood” Elimination Tree

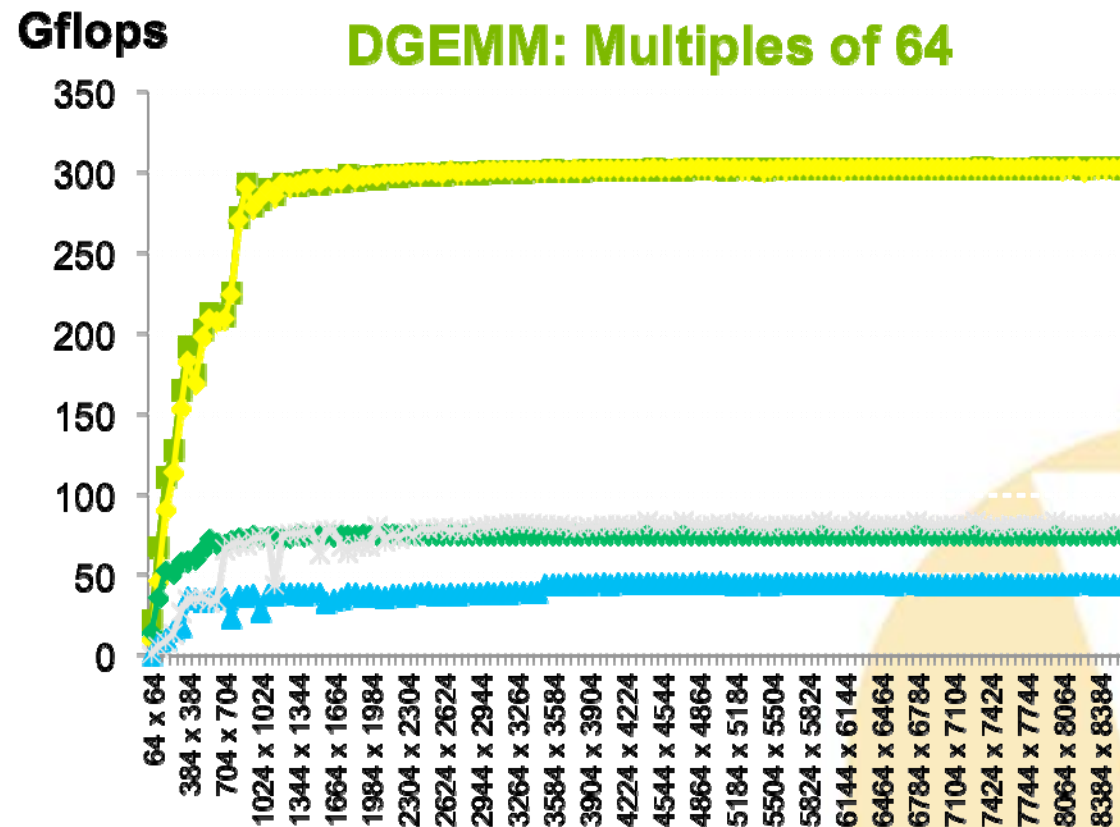


**Each frontal matrix's triangle scaled
by operations required to factor it.**

Exploiting Concurrency Toy with MPI & OpenMP



Why Accelerators?



Left-Looking Accelerator

Factorization of a Frontal Matrix

Assemble frontal matrix on host CPU

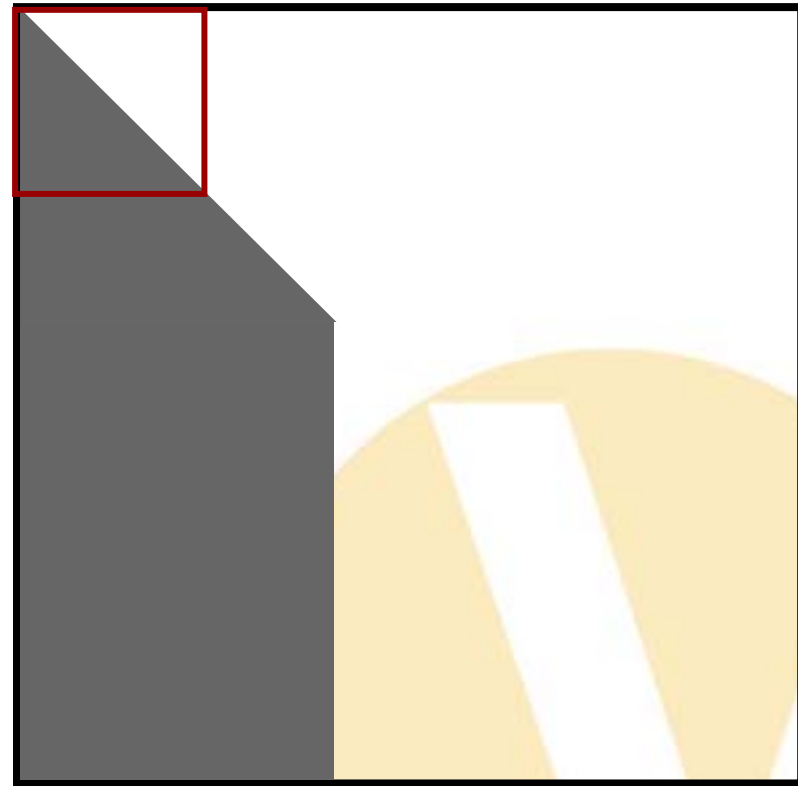
Initialize by sending panel of assembled frontal matrix

Only large frontal matrices due to high cost of sending data to and from accelerator



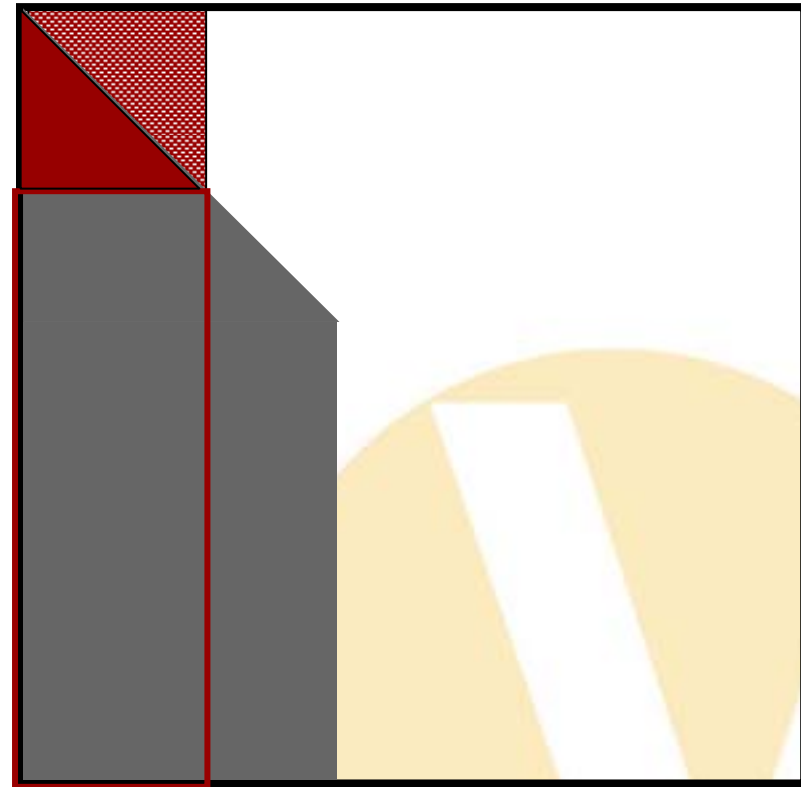
Eliminate panels

Factor diagonal block

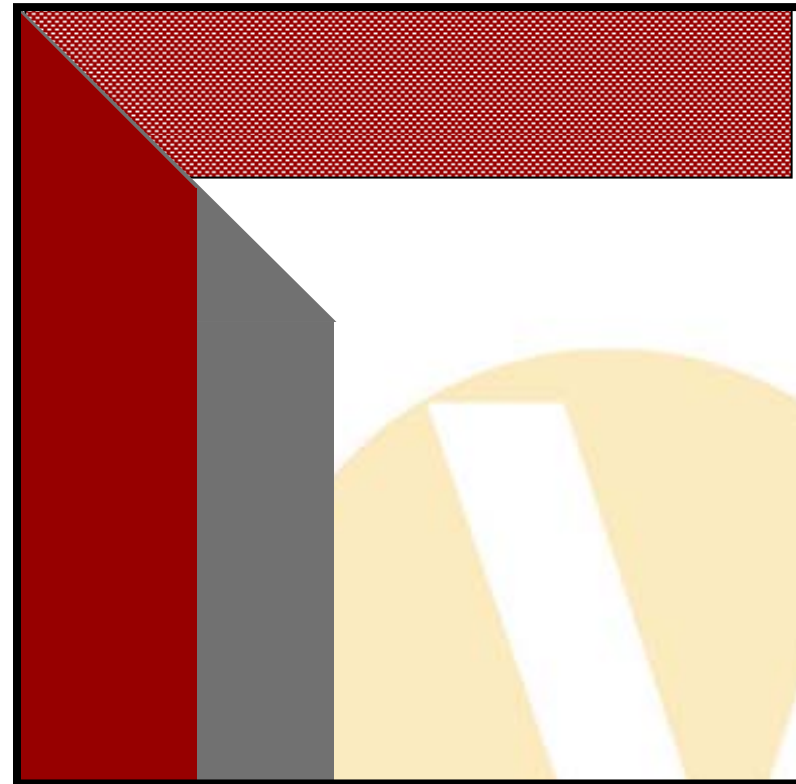


Eliminate panels

Eliminate off-diagonal panel



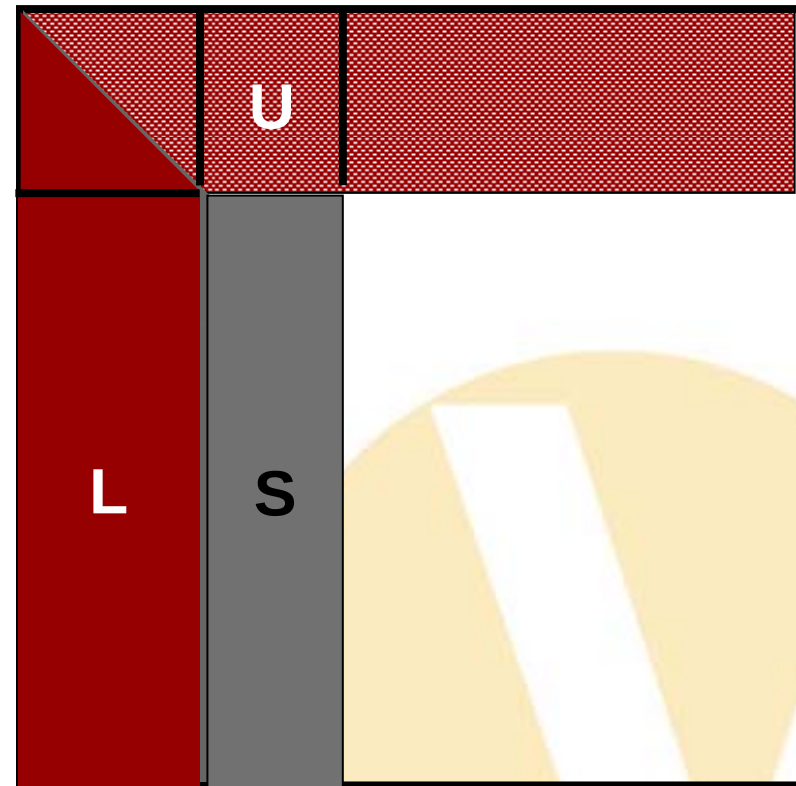
Fill Upper Triangle



Update panels with DGEMM

$$S = S - L * U$$

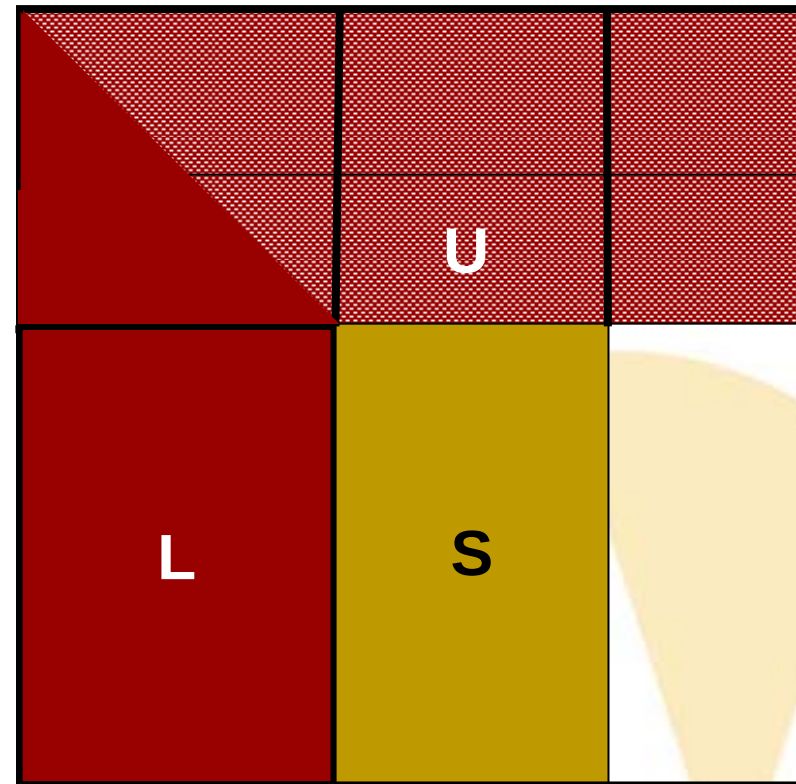
Assumes vendor math library
DGEMM is extremely fast



Wider panels in Schur complement

DGEMM is even faster

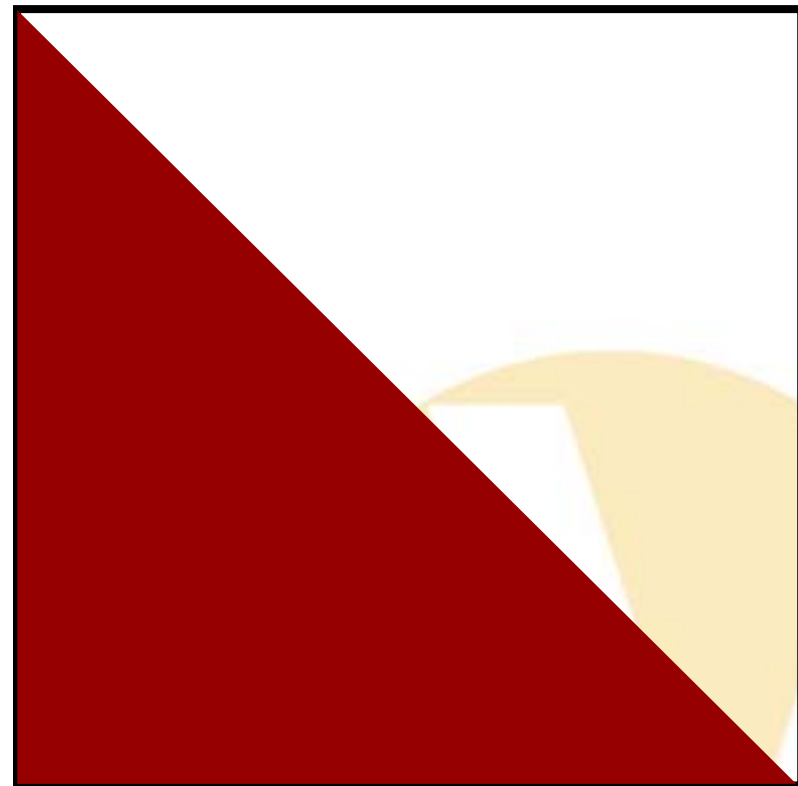
$$S = S - L * U$$



**Return error if diagonal of
0.0 encountered or pivot
threshold exceeded**

**Otherwise complete frontal
matrix is returned**

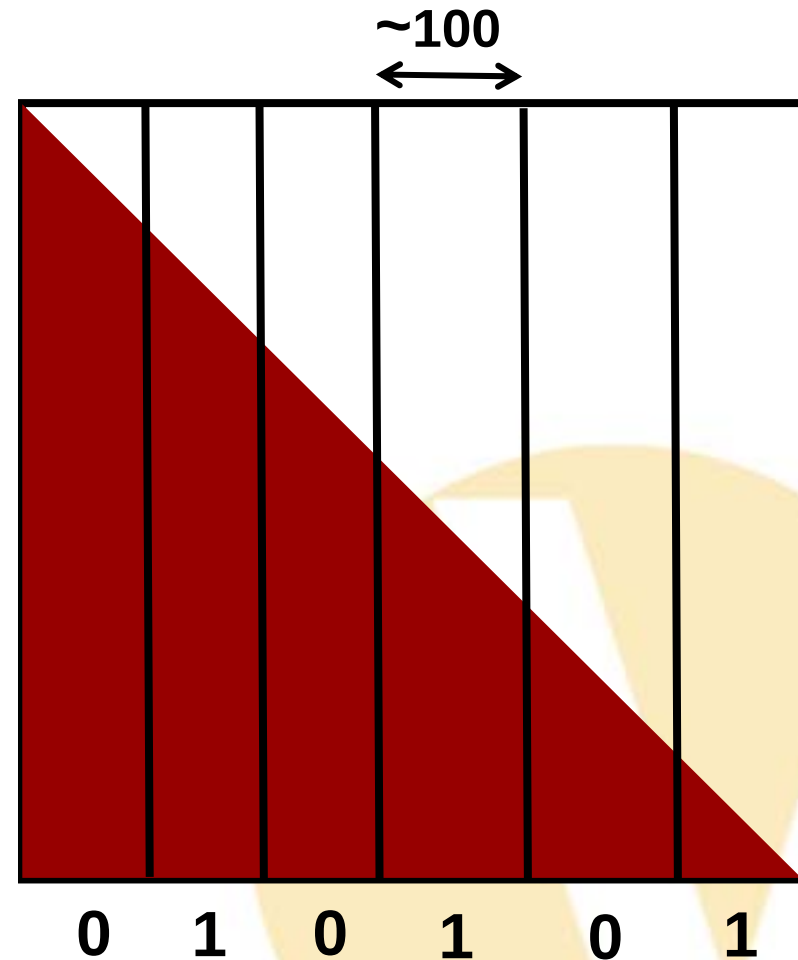
**Schur complement added to
initial values on host CPU**



**Factor panels broadcast
and sent to GPU.**

**GPU eliminates factor panel
and performs large-rank
updates to the Schur-
complement.**

**Load imbalance is major
bottleneck, so this needs
revisiting**



OpenMP on MIC (and on Xeon host)

```
do j = jl, jr
  do i = jl + 1, ld
    x = 0.0
    do k = jl, j - 1
      x = x + s(i, k) * s(k, j)
    end do
    s(i, j) = s(i, j) - x
  end do
end do
```

```
sq = jr - jl + 1
c$omp parallel do
c$omp& shared (jl, jr, sq, ld, s)
c$omp& private (ii, j, i, il, ir)
  do ii = jl + 1, ld, sq
    do j = jl, jr
      il = ii
      ir = min(ii + sq - 1, ld)
      do i = il, ir
        x = 0.0
        do k = jl, j - 1
          x = x + s(i, k) * s(k, j)
        end do
        s(i, j) = s(i, j) - x
      end do
    end do
  end do
end do
```

Fortran vs CUDA

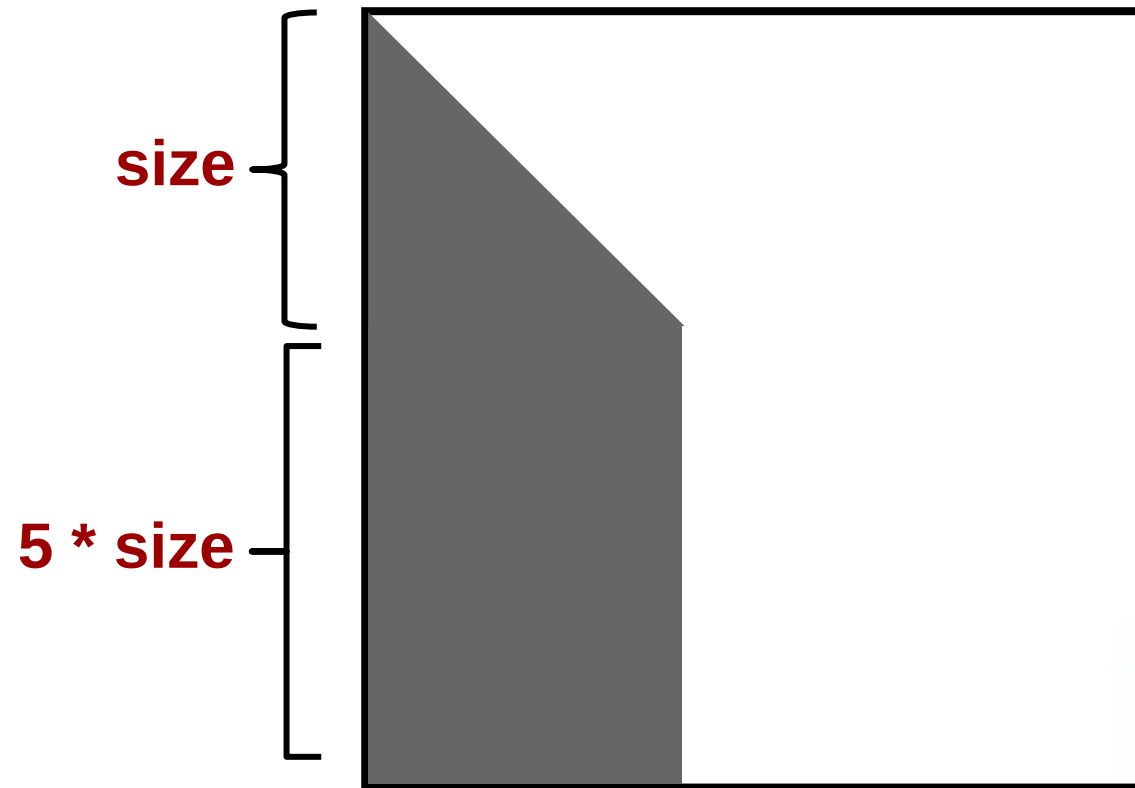
```
do j = jl, jr
  do i = jr + 1, ld
    x = 0.0
    do k = jl, j - 1
      x = x + s(i, k) * s(k, j)
    end do
    s(i, j) = s(i, j) - x
  end do
end do
```

```
ip=0;
for (j = jl; j <= jr; j++) {
  if(ltid <= (j-1)-jl){
    gpulskj(ip+ltid) = s[IDX(jl+ltid,j)];
  }
  ip = ip + (j - 1) - jl + 1;
}

__syncthreads();

for (i = jr + 1 + tid; i <= ld;
     i += GPUL_THREAD_COUNT) {
  for (j = jl; j <= jr; j++) {
    gpuls(j-jl,ltid) = s[IDX(i,j)];
  }
  ip=0;
  for (j = jl; j <= jr; j++) {
    x = 0.0f;
    for (k = jl; k <= (j-1); k++) {
      x = x + gpuls(k-jl,ltid) * gpulskj(ip);
      ip = ip + 1;
    }
    gpuls(j-jl,ltid) -= x;
  }
  for (j = jl; j <= jr; j++) {
    s[IDX(i,j)] = gpuls(j-jl,ltid);
  }
}
```

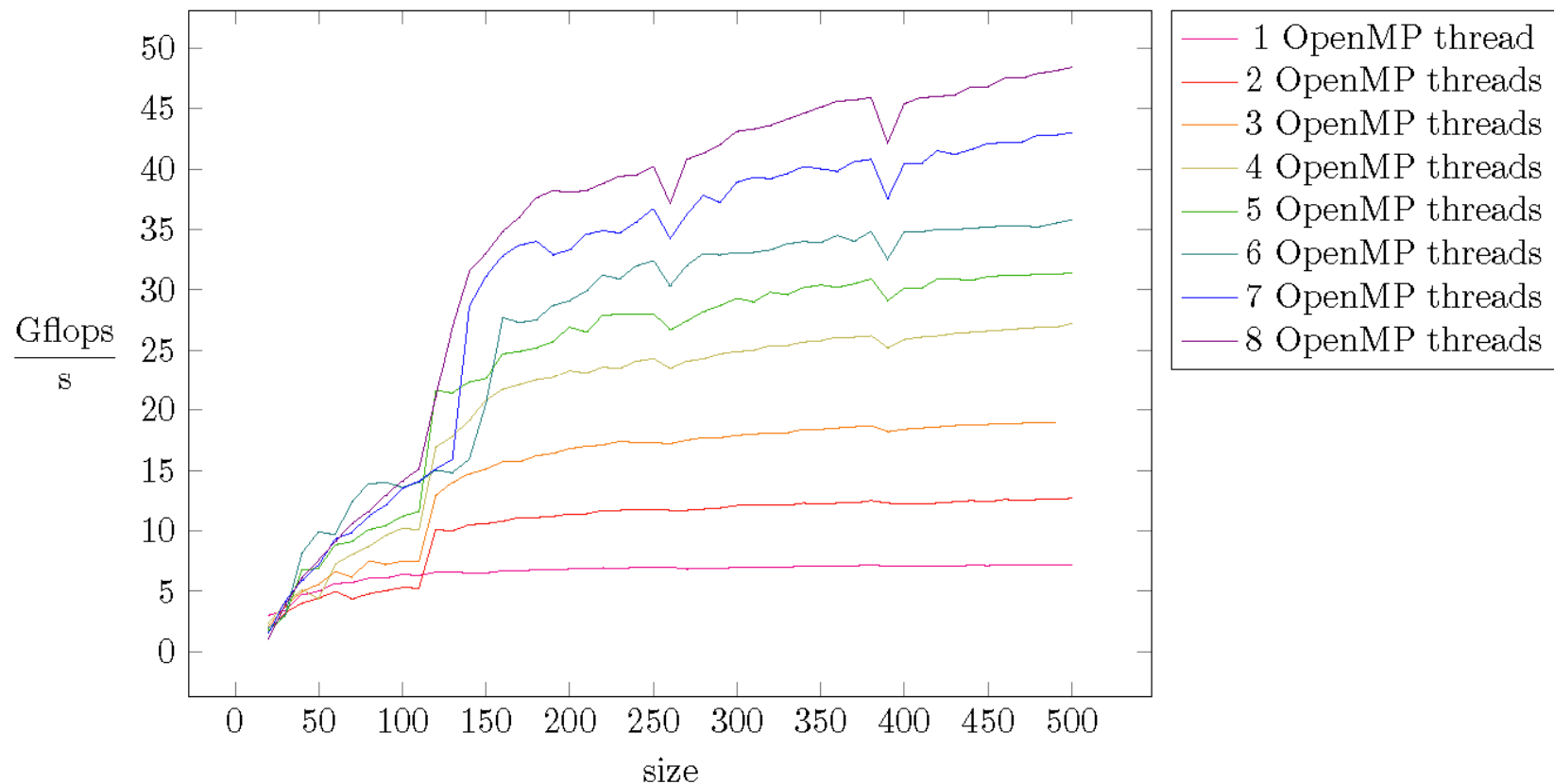
Factorization Performance on Individual Frontal Matrices



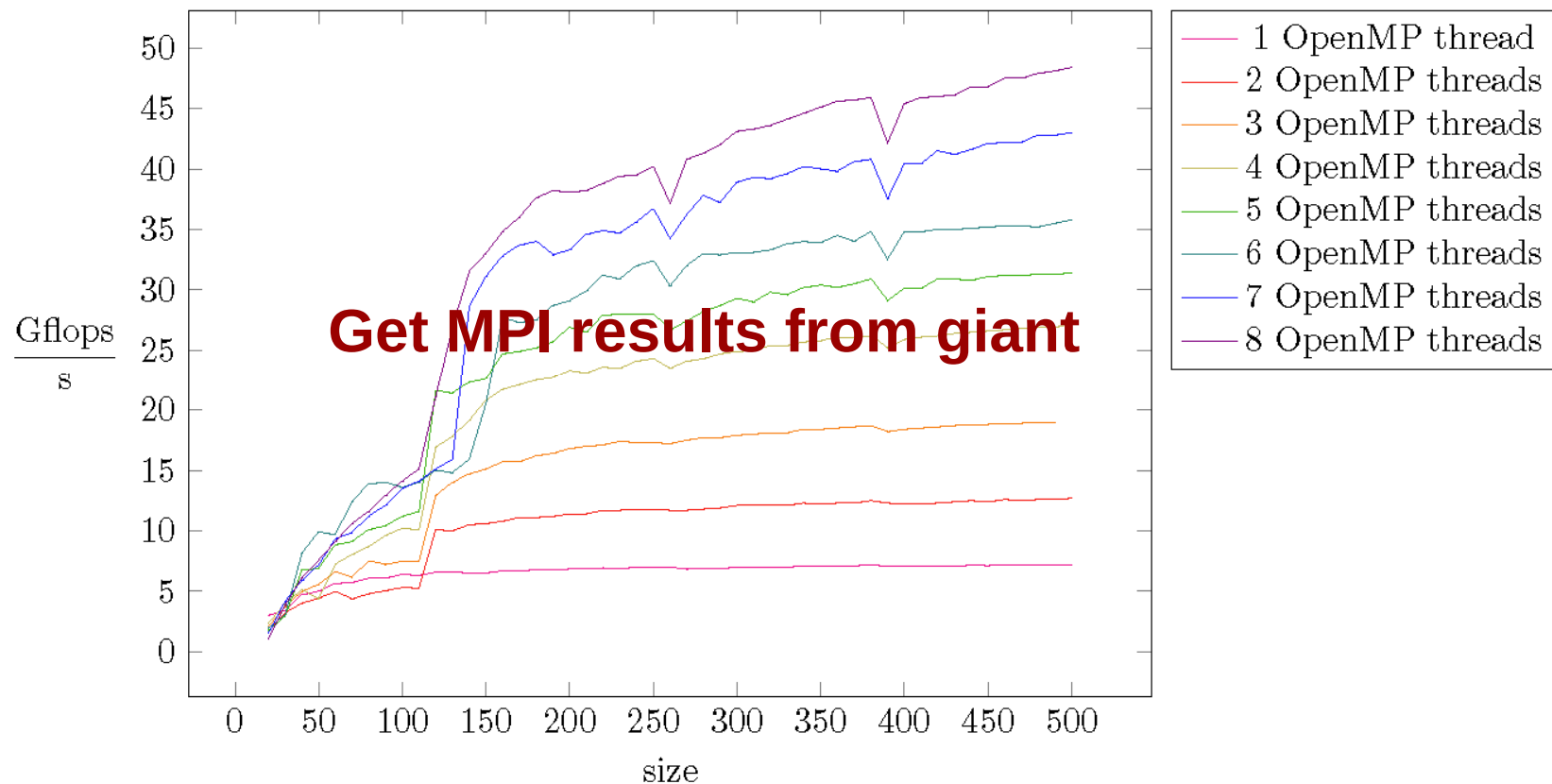
**Construct a family
of model frontal
matrices**

**Eliminate “size”
equations from 5X
their number**

Multithreaded Host Throughput Factoring a Model Frontal Matrix

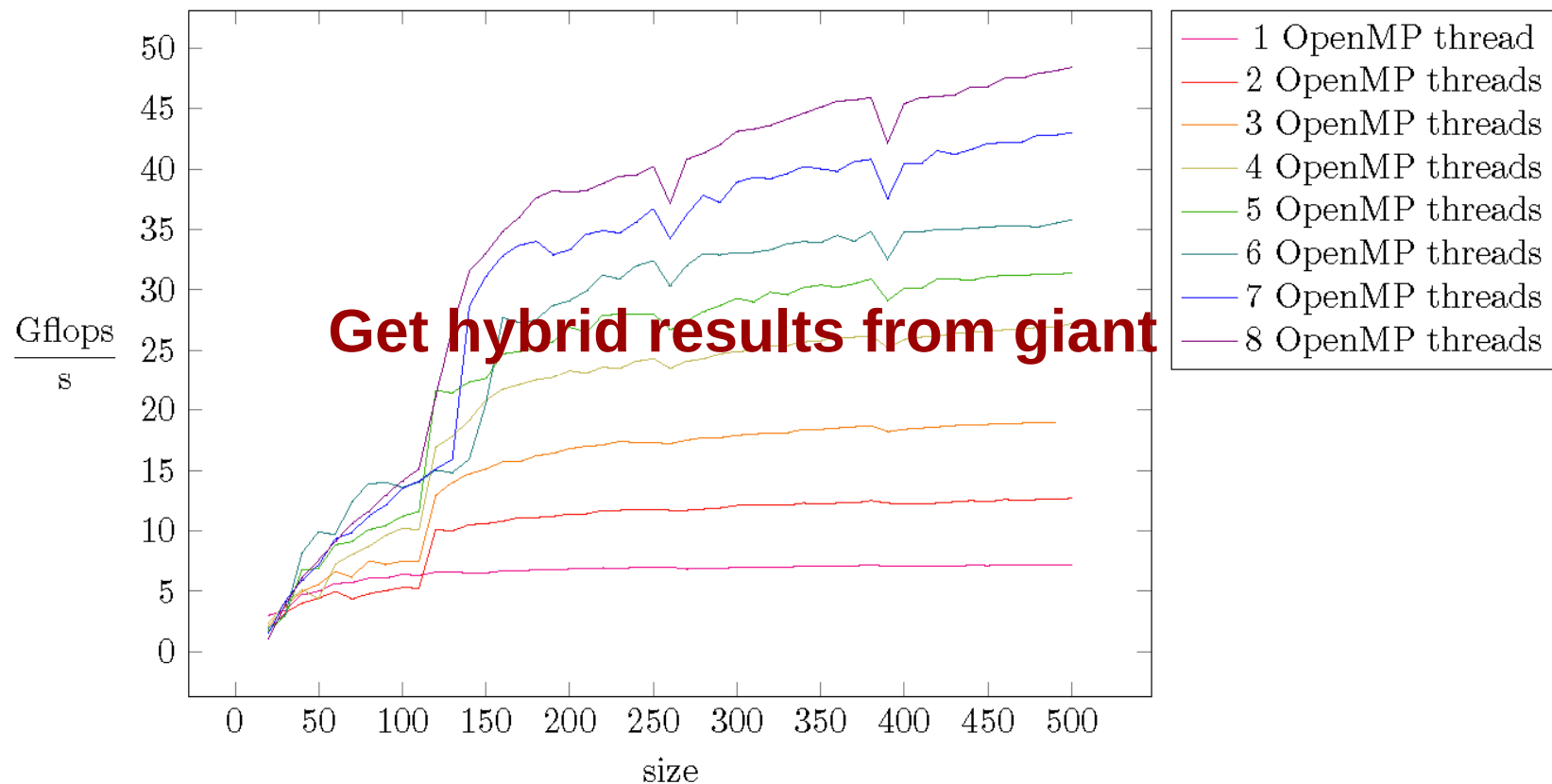


Host Throughput with MPI Factoring a Model Frontal Matrix



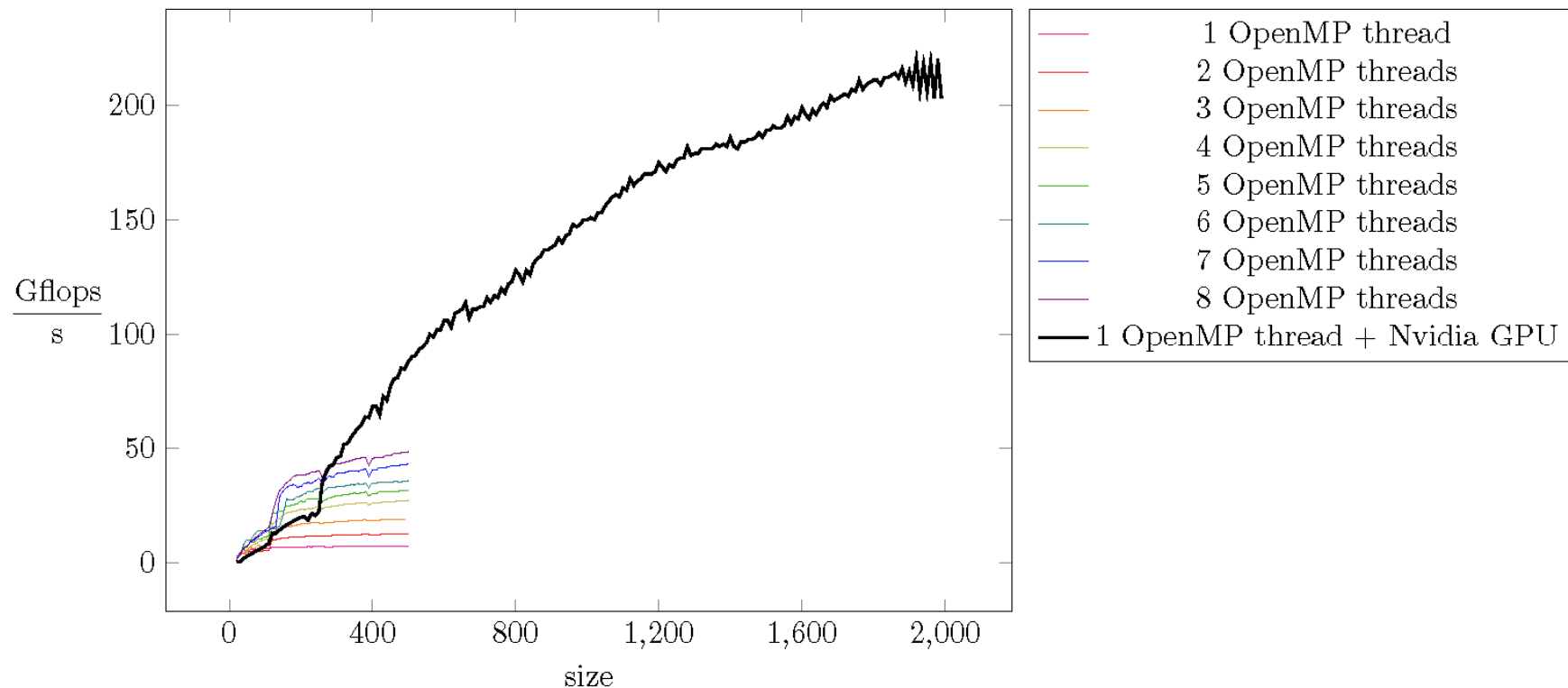
For larger frontal matrices,
XX Gflop/s have been observed

Hybrid (MPI + OpenMP) Throughput Factoring a Model Frontal Matrix



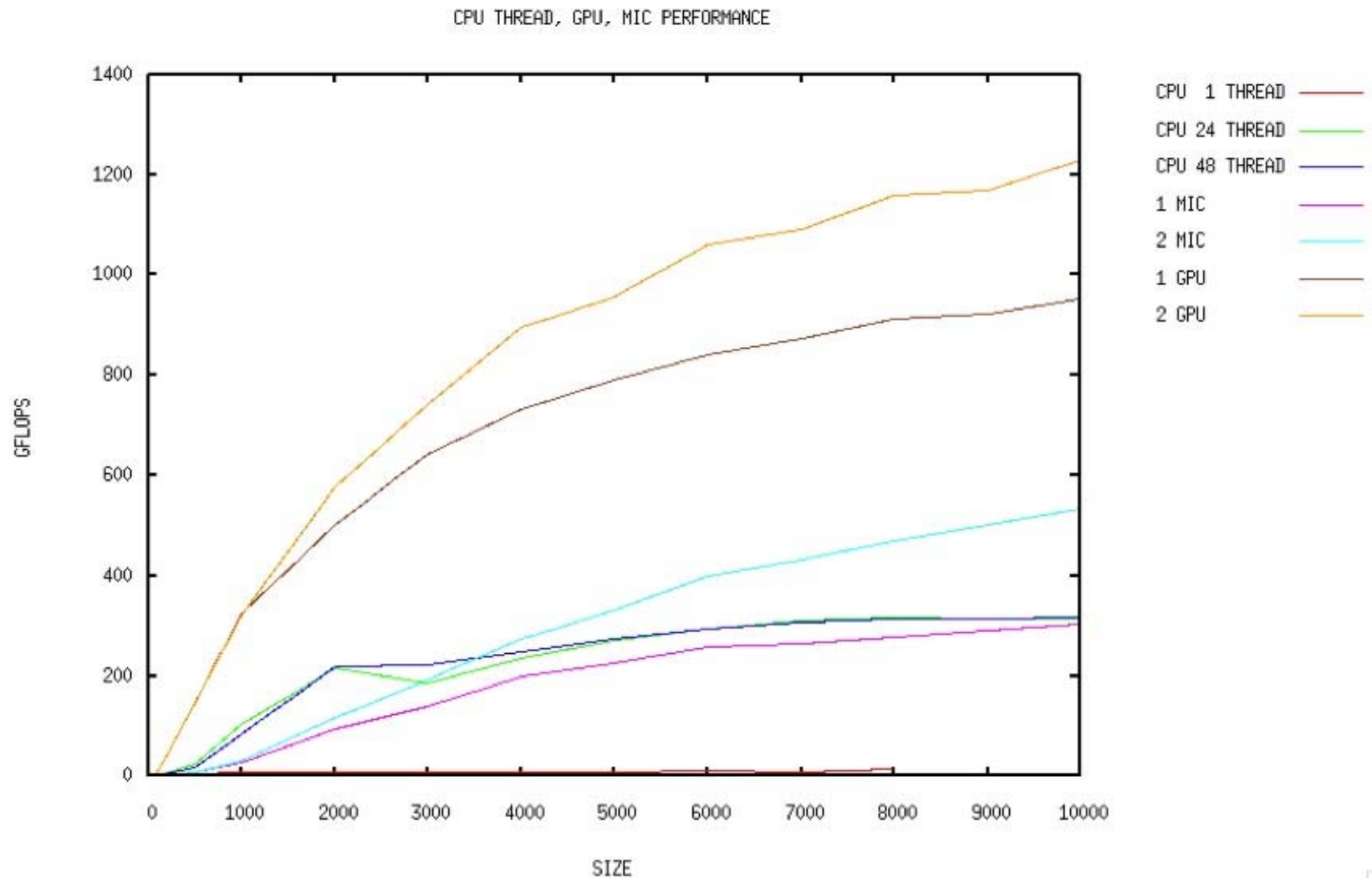
For larger frontal matrices,
XX Gflop/s have been observed

Nehalem & Tesla Throughput Factoring a Model Frontal Matrix

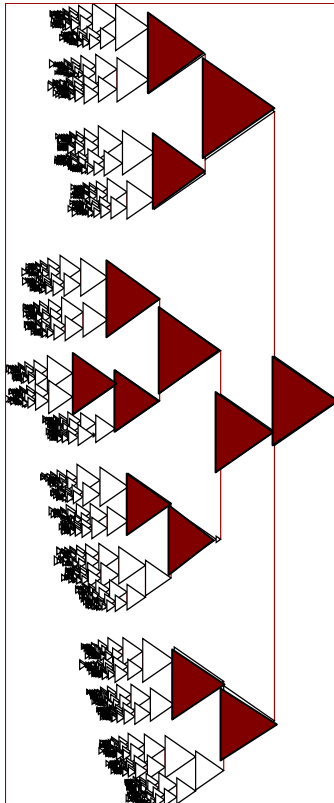


**For larger frontal matrices,
>250 Gflop/s have been observed**

Ivy Bridge & Accelerator Factoring a Model Frontal Matrix



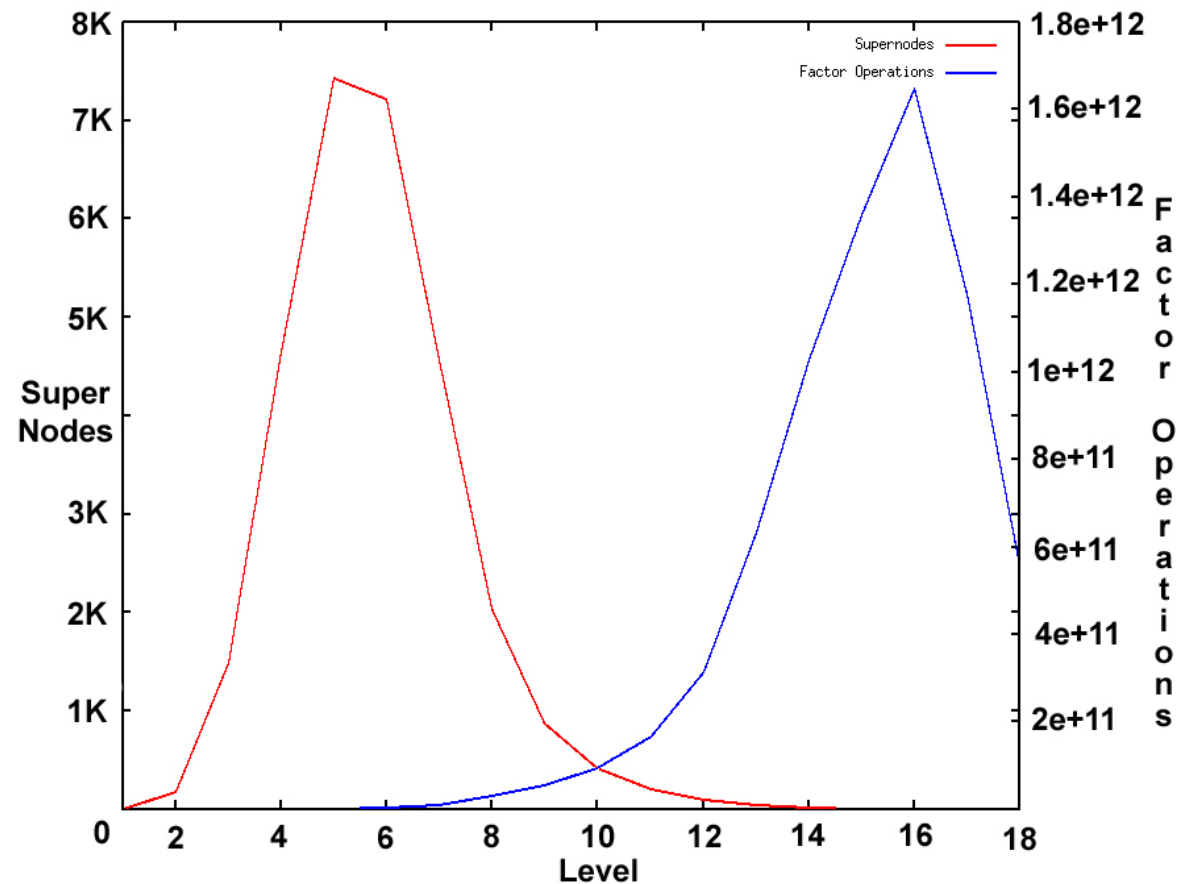
GPUs Get the Larger Supernodes



Hood's tree

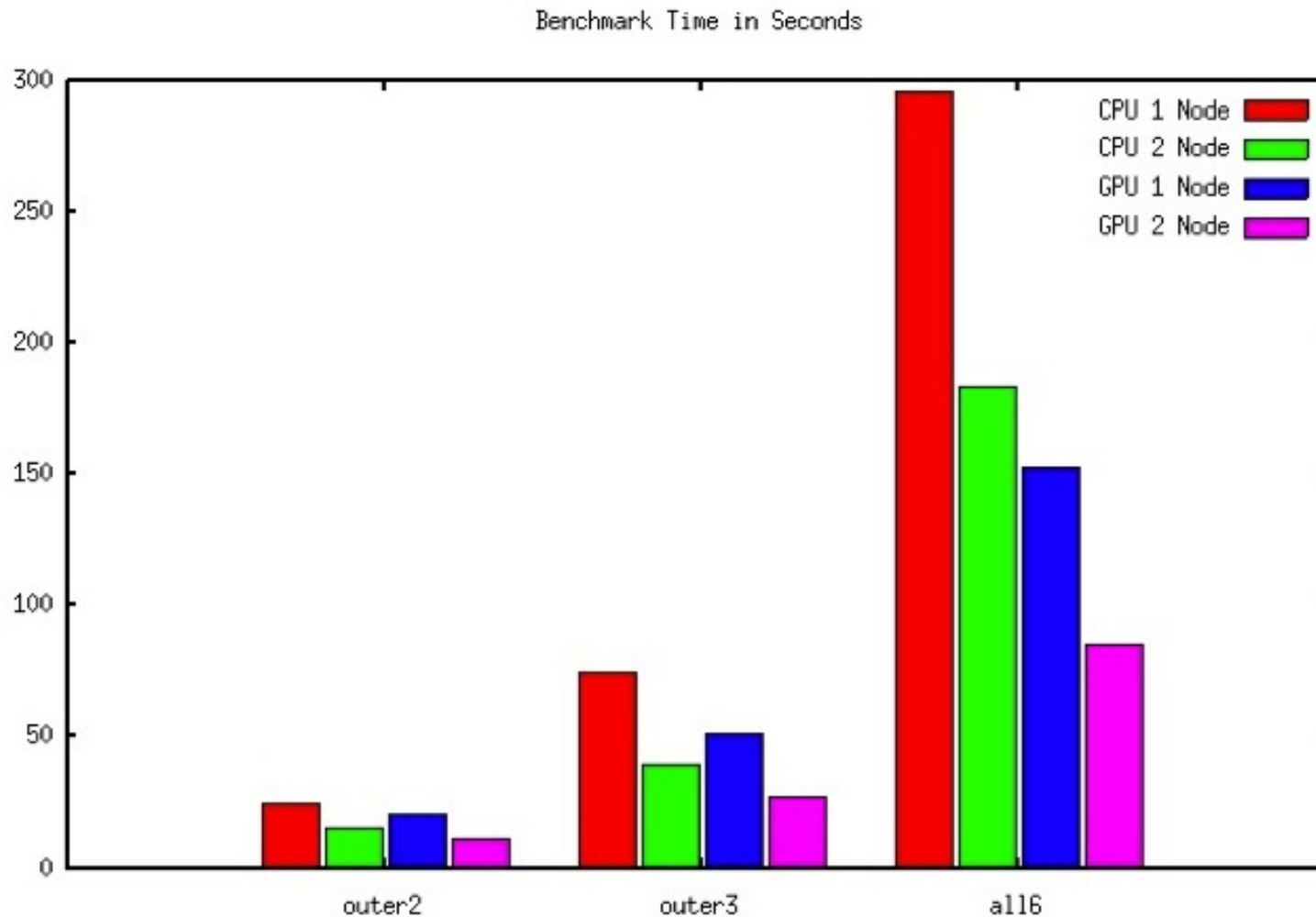


Number of SuperNodes and Factor Operations per Level



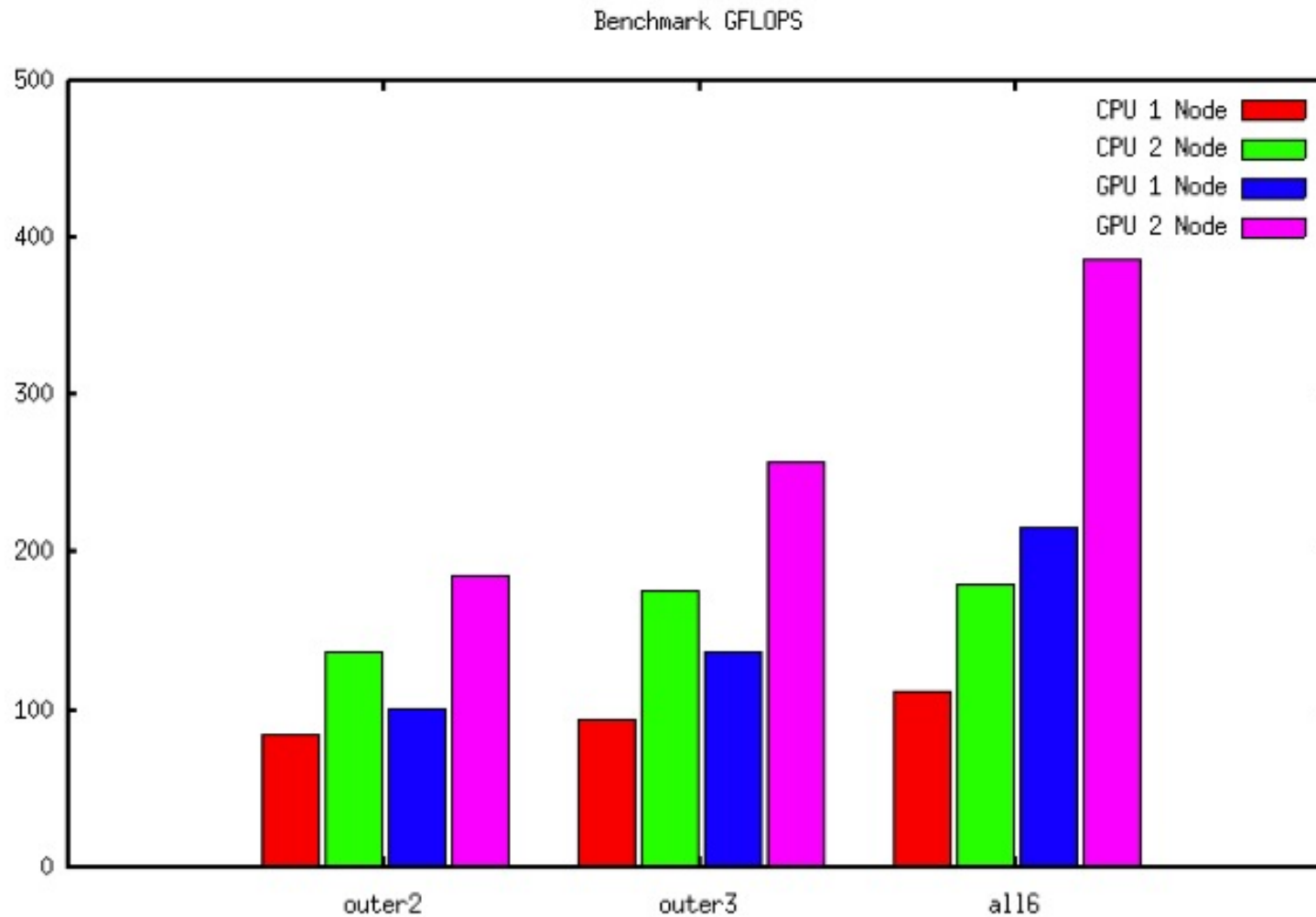
Three cylinders benchmark

Relative Time (sec)



USC Viterbi Relative Performance (GFlop/s)

School of Engineering



Nested cylinders benchmarks

- **Near Term**
 - Further accelerator performance optimization
 - Better MPI + accelerators
 - Pivoting on accelerators for numerical stability
- **Longer Term**
 - Factor many small frontal matrices on the accelerators
 - Assemble initial values
 - Maintain real stack
 - If the entire matrix fits on the accelerator
 - Forward and back solves
 - Exploit GDRAM memory bandwidth

Research Initially Funded by US JFCOM and AFRL

This material is based on research sponsored by the U.S. Joint Forces Command via a contract with the Lockheed Martin Corporation and SimIS, Inc., and on research sponsored by the Air Force Research Laboratory under agreement numbers F30602-02-C-0213 and FA8750-05-2-0204. The U.S. Government is authorized to reproduce and distribute reprints for Governmental purposes notwithstanding any copyright notation thereon. The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies or endorsements, either expressed or implied, of the U.S. Government. Approved for public release; distribution is unlimited.

Bonus Slides

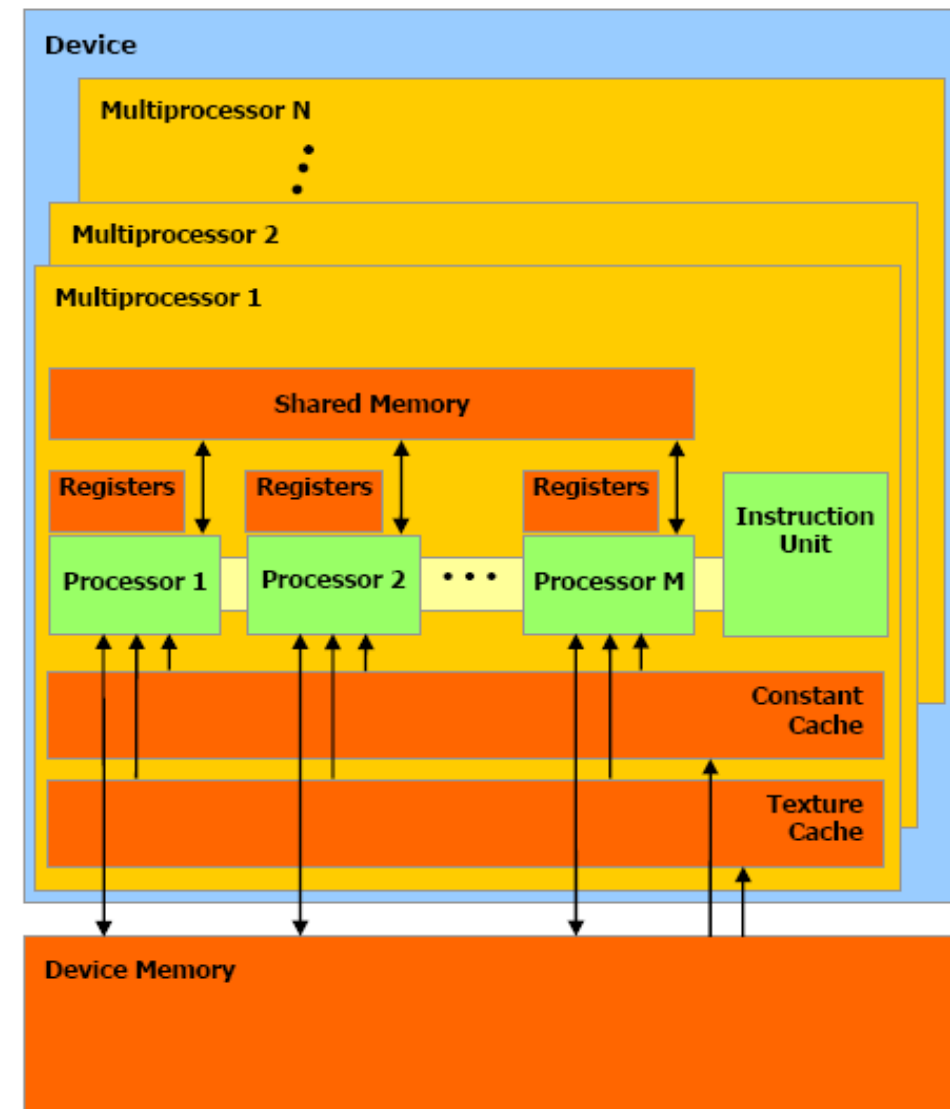


Multiple SIMD cores

Multithreaded
O(1000) per GPU

Banked shared memory
32 Kbytes Tesla
48 Kbytes Fermi

Simple thread model
Only sync at host



A set of SIMD multiprocessors with on-chip shared memory.

Figure 3-1. Hardware Model

Courtesy NVIDIA

Fortran vs CUDA

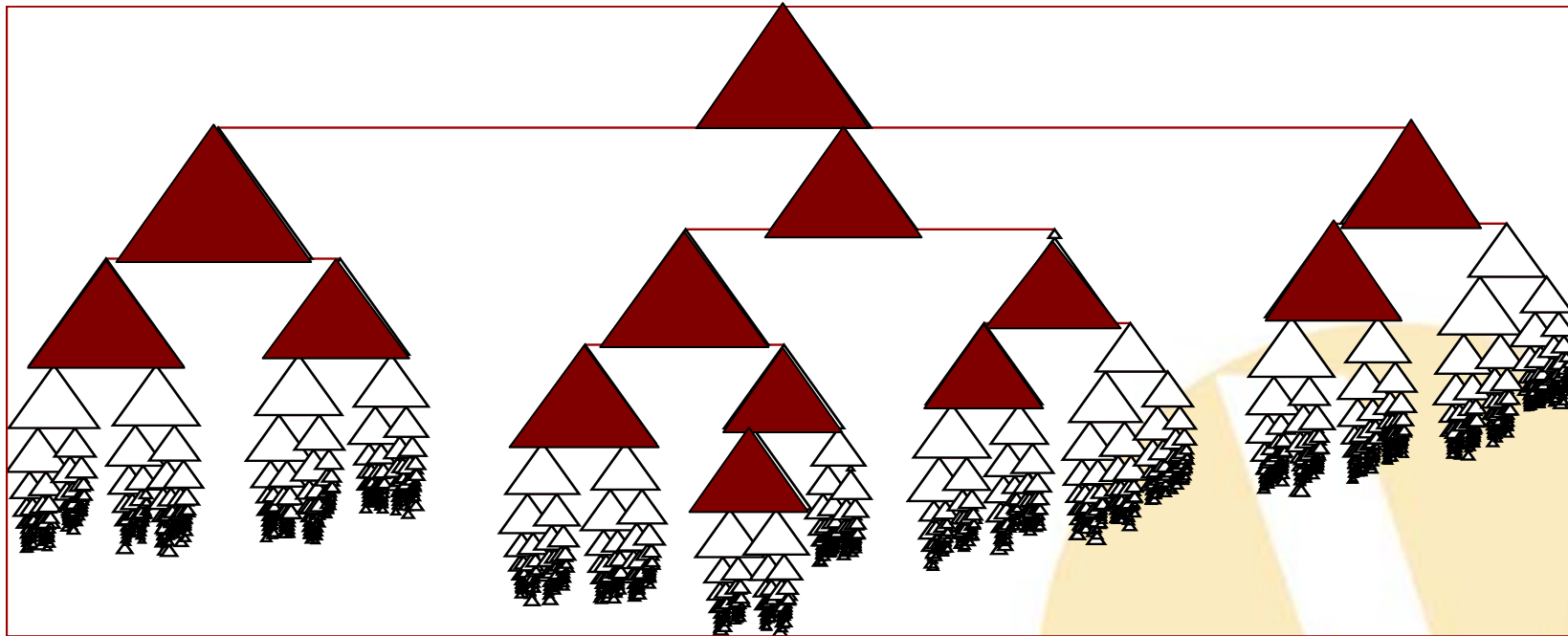
```
do j = jl, jr
  do i = jr + 1, ld
    x = 0.0
    do k = jl, j - 1
      x = x + s(i, k) * s(k, j)
    end do
    s(i, j) = s(i, j) - x
  end do
end do
```

```
ip=0;
for (j = jl; j <= jr; j++) {
  if(ltid <= (j-1)-jl){
    gpulskj(ip+ltid) = s[IDXs(jl+ltid,j)];
  }
  ip = ip + (j - 1) - jl + 1;
}

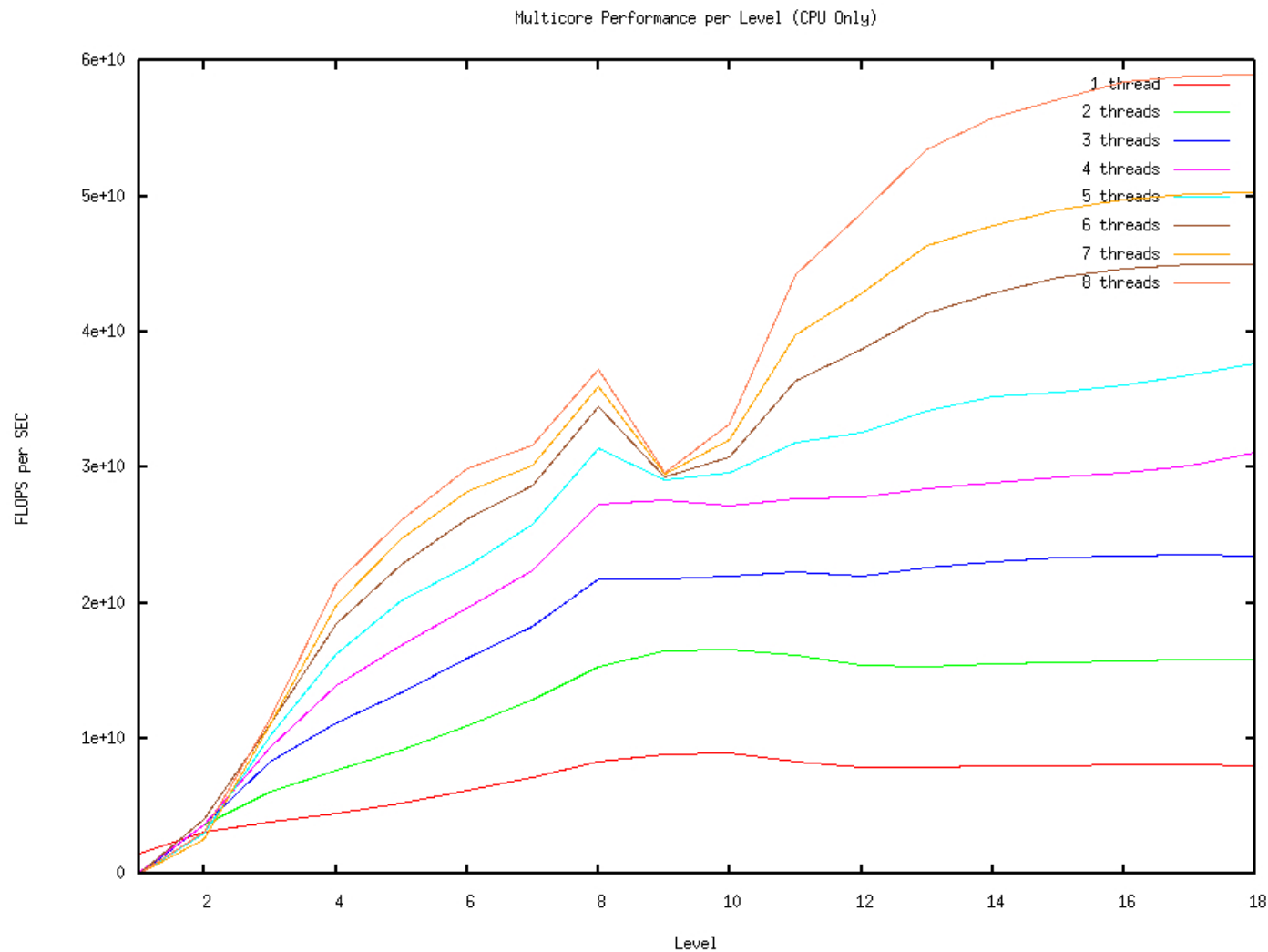
__syncthreads();

for (i = jr + 1 + tid; i <= ld;
     i += GPUL_THREAD_COUNT) {
  for (j = jl; j <= jr; j++) {
    gpuls(j-jl,ltid) = s[IDXs(i,j)];
  }
  ip=0;
  for (j = jl; j <= jr; j++) {
    x = 0.0f;
    for (k = jl; k <= (j-1); k++) {
      x = x + gpuls(k-jl,ltid) * gpulskj(ip);
      ip = ip + 1;
    }
    gpuls(j-jl,ltid) -= x;
  }
  for (j = jl; j <= jr; j++) {
    s[IDXs(i,j)] = gpuls(j-jl,ltid);
  }
}
```

GPU gets larger Supernodes



Multicore Performance (i4r4) vs. the Elimination Tree



Relative Performance (sec) half-million elements, in-core

MPI, OpenMP, GPU	1 st Factorization	LS-DYNA
8 threads	795	1806
8 threads + GPU	283	775
2 MPI x 4 threads	772	1659
2 MPI x 4 + 2 GPU	211	537

Relative Performance (sec) million elements, out-of-core

MPI, OpenMP, GPU	1 st Factorization	LS-DYNA
8 threads	6131	15562
8 threads + GPU	2828	9143
2 MPI x 4 threads	6248	15287
2 MPI x 4 + 2 GPU	2567	8214

Time in out-of-core solver (sec) Eight threads plus GPU

```
grep WCT mes0000 =>
```

```
WCT: symbolic factorization = 51.764
WCT: matrix redistribution = 10.875
WCT: factorization          = 2828.933
WCT: numeric solve         = 1409.523
WCT: total imfslv_mf2       = 4316.063
WCT: symbolic factorization = 40.244
WCT: matrix redistribution = 10.685
WCT: factorization          = 2679.422
WCT: numeric solve         = 1543.953
WCT: total imfslv_mf2       = 4346.260
```

**Note: almost half the elapsed time for LS-DYNA
(~4400 sec. out of 9143) is spent in disk I/O.**