

L13: Application Case Studies

CS6235



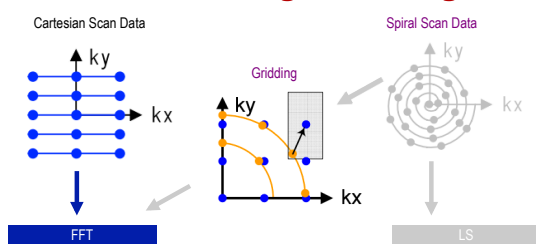
Outline

- Application Case Studies
- Advanced MRI Reconstruction
Reading: Kirk and Hwu, Chapter 8 or
From the sample book
<http://courses.ece.illinois.edu/ece498/al/textbook/Chapter7-MRI-Case-Study.pdf>

L12: Application Case Studies



Reconstructing MR Images



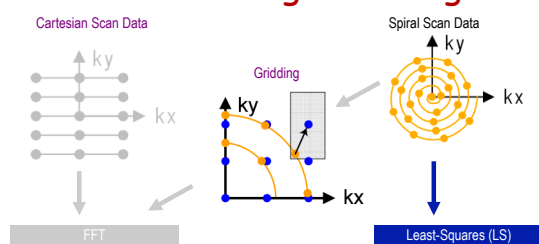
Cartesian scan data + FFT:
Slow scan, fast reconstruction, images may be poor

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007–2009
University of Illinois, Urbana-Champaign

3



Reconstructing MR Images



Spiral scan data + LS
Superior images at expense of significantly more computation

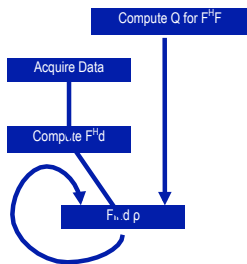
© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007–2009
University of Illinois, Urbana-Champaign

4



Least-Squares Reconstruction

$$F^H F \rho = F^H d$$



- Q : a data structure that allows us to compute $F^H F$
- Q depends only on scanner configuration
- $F^H d$ depends on scan data
- Aim: find ρ ; using linear solver
- Accelerate Q and $F^H d$ on GPU
 - Q: 1-2 days on CPU
 - $F^H d$: 6-7 hours on CPU
 - ρ : 1.5 minutes on CPU

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

5



```

for (m = 0; m < M; m++) {
    phiMag[m] = rPhi[m]*rPhi[m] +
               iPhi[m]*iPhi[m];

    for (n = 0; n < N; n++) {
        expQ = 2*PI*(kx[m]*x[n] +
                    ky[m]*y[n] +
                    kz[m]*z[n]);

        rQ[n] += phiMag[m]*cos(expQ);
        iQ[n] += phiMag[m]*sin(expQ);
    }
}
  
```

(a) Q computation

Q v.s. $F^H d$

```

for (m = 0; m < M; m++) {
    rMu[m] = rPhi[m]*rD[m] +
             iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] -
             iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] +
                      ky[m]*y[n] +
                      kz[m]*z[n]);

        cArg = cos(expFhD);
        sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg -
                  iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg +
                  rMu[m]*sArg;
    }
}
  
```

(b) $F^H d$ computation

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

6



Algorithms to Accelerate

```

for (m = 0; m < M; m++) {
    rMu[m] = rPhi[m]*rD[m] +
             iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] -
             iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] +
                      ky[m]*y[n] +
                      kz[m]*z[n]);

        cArg = cos(expFhD);
        sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg -
                  iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg +
                  rMu[m]*sArg;
    }
}
  
```

- Scan data
 - M = # scan points
 - kx, ky, kz = 3D scan data
- Pixel data
 - N = # pixels
 - x, y, z = input 3D pixel data
 - rFhD, iFhD = output pixel data
- Complexity is $O(MN)$
- Inner loop
 - 13 FP MUL or ADD ops
 - 2 FP trig ops
 - 12 loads, 2 stores

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

7



Step 1. Consider Parallelism to Evaluate Partitioning Options

```

for (m = 0; m < M; m++) {
    rMu[m] = rPhi[m]*rD[m] +
             iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] -
             iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] +
                      ky[m]*y[n] +
                      kz[m]*z[n]);

        cArg = cos(expFhD);
        sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg -
                  iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg +
                  rMu[m]*sArg;
    }
}
  
```

What about M total threads?

Note: M is O(millions)

(Step 2) What happens to data accesses with this strategy?



One Possibility

```
__global__ void cmpFhD(float* rPhi, iPhi, phiMag,
    kx, ky, kz, x, y, z, rMu, iMu, int N) {

    int m = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]);

        cArg = cos(expFhD);  sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg - iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg + rMu[m]*sArg;
    }
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

8



One Possibility

```
__global__ void cmpFhD(float* rPhi, iPhi, phiMag,
    kx, ky, kz, x, y, z, rMu, iMu, int N) {

    int m = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]);

        cArg = cos(expFhD);  sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg - iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg + rMu[m]*sArg;
    }
}
```

Major flaw: This code does not
work correctly! The accumulation
needs to use atomic operation.

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

8



Back to the Drawing Board - Maybe map the n loop to threads?

```
for (m = 0; m < M; m++) {

    rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m];

    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]);
        cArg = cos(expFhD);
        sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg - iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg + rMu[m]*sArg;
    }
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

9



Loop Fission

<pre>for (m = 0; m < M; m++) { rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m]; iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m]; for (n = 0; n < N; n++) { expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]); cArg = cos(expFhD); sArg = sin(expFhD); rFhD[n] += rMu[m]*cArg - iMu[m]*sArg; iFhD[n] += iMu[m]*cArg + rMu[m]*sArg; } }</pre>		<pre>for (m = 0; m < M; m++) { rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m]; iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m]; for (n = 0; n < N; n++) { expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]); cArg = cos(expFhD); sArg = sin(expFhD); rFhD[n] += rMu[m]*cArg - iMu[m]*sArg; iFhD[n] += iMu[m]*cArg + rMu[m]*sArg; } }</pre>
--	--	--

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

10



A Separate cmpMu Kernel

- After loop fission
 - computation is in 2 steps
 - 2 parallel kernels executed one after the other
 - first step implemented below:

```
__global__ void cmpMu(float* rPhi, iPhi, rD, iD, rMu, iMu)
{
    int m = blockIdx.x * MU_THREADS_PER_BLOCK + threadIdx.x;

    rMu[m] = rPhi[m]*rD[m] + iPhi[m]*iD[m];
    iMu[m] = rPhi[m]*iD[m] - iPhi[m]*rD[m];
}
```

© David Kirk/NVIDIA and
Wen-mei W. Hwu,
2007-2009
University of Illinois,
Urbana-Champaign

11



Loop Interchange

```
for (m = 0; m < M; m++) {
    for (n = 0; n < N; n++) {
        expFhD = 2*PI*(kx[m]*x[n] +
                      ky[m]*y[n] +
                      kz[m]*z[n]);

        cArg = cos(expFhD);
        sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg -
                  iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg +
                  rMu[m]*sArg;
    }
}
```

)} (a) before loop interchange

- Parallelizing second loop
- Options:
 - 1. one thread for each inner loop
 - $N * M$ threads
 - 2. loop interchange
 - each outer loop processes
an $rFhD[n]$ and $iFhD[n]$

Figure 7.9 Loop interchange of the F^hD computation

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

12



Loop Interchange

- Possible since all iterations of both level are independent

<pre>for (m = 0; m < M; m++) { for (n = 0; n < N; n++) { expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]); cArg = cos(expFhD); sArg = sin(expFhD); rFhD[n] += rMu[m]*cArg - iMu[m]*sArg; iFhD[n] += iMu[m]*cArg + rMu[m]*sArg; } }</pre>		<pre>for (n = 0; n < N; n++) { for (m = 0; m < M; m++) { expFhD = 2*PI*(kx[m]*x[n] + ky[m]*y[n] + kz[m]*z[n]); cArg = cos(expFhD); sArg = sin(expFhD); rFhD[n] += rMu[m]*cArg - iMu[m]*sArg; iFhD[n] += iMu[m]*cArg + rMu[m]*sArg; } }</pre>
--	--	--

)} (a) before loop interchange } (b) after loop interchange

Figure 7.9 Loop interchange of the F^hD computation

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

12



Step 2. New FhD kernel

```
__global__ void cmpFhD(float* kx,ky,kz, x,y,z, rMu,iMu, int
M){
    int n = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    for (m = 0; m < M; m++) {
        float expFhD = 2*PI*(kx[m]*x[n]+ky[m]*y[n]+kz[m]*z[n]);

        float cArg = cos(expFhD);
        float sArg = sin(expFhD);

        rFhD[n] += rMu[m]*cArg - iMu[m]*sArg;
        iFhD[n] += iMu[m]*cArg + rMu[m]*sArg;
    }
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

13



Step 3. Using Registers to Reduce Global Memory Traffic

```
__global__ void cmpFHD(float* rPhi, iPhi, phiMag,
    kx, ky, kz, x, y, z, rMu, iMu, int M) {

    int n = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    float xn_r = x[n]; float yn_r = y[n]; float zn_r = z[n];
    float rFhDn_r = rFhD[n]; float iFhDn_r = iFhD[n];

    for (m = 0; m < M; m++) {
        float expFhD = 2*PI*(kx[m]*xn_r+ky[m]*yn_r+kz[m]*zn_r);

        float cArg = cos(expFhD);
        float sArg = sin(expFhD);

        rFhDn_r += rMu[m]*cArg - iMu[m]*sArg;
        iFhDn_r += iMu[m]*cArg + rMu[m]*sArg;
    }
    rFhD[n] = rFhD_r; iFhD[n] = iFhD_r;
}
```

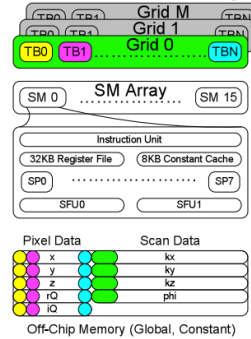
Still too much stress on memory! Note that kx, ky and kz are read-only and based on m

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

14



Tiling of Scan Data



LS recon uses multiple grids

- Each grid operates on all pixels
- Each grid operates on a distinct subset of scan data
- Each thread in the same grid operates on a distinct pixel

Thread n operates on pixel n:

```
for (m = 0; m < M/32; m++) {
    exQ = 2*PI*(kx[m]*x[n] +
                ky[m]*y[n] +
                kz[m]*z[n]);
    rQ[n] += phi[m]*cos(exQ);
    iQ[n] += phi[m]*sin(exQ);
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

15



Tiling k-space data to fit into constant memory

```
__constant__ float    kx_c[CHUNK_SIZE],
                      ky_c[CHUNK_SIZE],
                      kz_c[CHUNK_SIZE];

... void main() {

    int i;
    for (i = 0; i < M/CHUNK_SIZE; i++){
        cudaMemcpyToSymbol(kx_c, &kx[i*CHUNK_SIZE], 4*CHUNK_SIZE);
        cudaMemcpyToSymbol(ky_c, &ky[i*CHUNK_SIZE], 4*CHUNK_SIZE);
        cudaMemcpyToSymbol(kz_c, &kz[i*CHUNK_SIZE], 4*CHUNK_SIZE);
        ...
        cmpFHD<<<N/FHD_THREADS_PER_BLOCK, FHD_THREADS_PER_BLOCK>>>
            (rPhi, iPhi, phiMag, x, y, z, rMu, iMu, int M);
    }
    /* Need to call kernel one more time if M is not */
    /* perfect multiple of CHUNK SIZE */
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu, 2007-2009
University of Illinois, Urbana-Champaign

16



Revised Kernel for Constant Memory

```
__global__ void cmpFHD(float* x, y, z, rMu, iMu, int M) {

    int n = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    float xn_r = x[n]; float yn_r = y[n]; float zn_r = z[n];
    float rFhDn_r = rFhD[n]; float iFhDn_r = iFhD[n];

    for (m = 0; m < M; m++) {
        float expFhD = 2*PI*(kx_c[m]*xn_r+ky_c[m]*yn_r+kz_c[m]*zn_r);

        float cArg = cos(expFhD);
        float sArg = sin(expFhD);

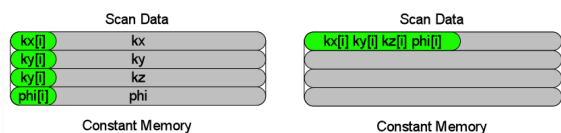
        rFhDn_r += rMu[m]*cArg - iMu[m]*sArg;
        iFhDn_r += iMu[m]*cArg + rMu[m]*sArg;
    }
    rFhD[n] = rFhD_r; iFhD[n] = iFhD_r;
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007-2009
University of Illinois, Urbana-Champaign

17



Sidebar: Cache-Conscious Data Layout



- kx, ky, kz, and phi components of same scan point have spatial and temporal locality
- Prefetching
- Caching
- Old layout does not fully leverage that locality
- New layout does fully leverage that locality

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007–2009
University of Illinois, Urbana-Champaign

18



Adjusting K-space Data Layout

```
struct kdata {
    float x, float y, float z;
} k;

__constant__ struct kdata k_c[CHUNK_SIZE];

...

void main() {
    int i;

    for (i = 0; i < M/CHUNK_SIZE; i++);
        cudaMemcpyToSymbol(k_c, k, 12*CHUNK_SIZE);

    cmpFHD<<<FHD_THREADS_PER_BLOCK,N/FHD_THREADS_PER_BLOCK>>>
        (...);
}
```

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007–2009
University of Illinois, Urbana-Champaign

19



```
__global__ void cmpFHD(float* rPhi, iPhi, phiMag,
    x, y, z, rMu, iMu, int M) {

    int n = blockIdx.x * FHD_THREADS_PER_BLOCK + threadIdx.x;

    float xn_r = x[n]; float yn_r = y[n]; float zn_r = z[n];
    float rFhDn_r = rFhD[n]; float iFhDn_r = iFhD[n];

    for (m = 0; m < M; m++) {
        float expFhD = 2*PI*(k[m].x*xn_r+k[m].y*yn_r+k[m].z*zn_r);

        float cArg = cos(expFhD);
        float sArg = sin(expFhD);

        rFhDn_r += rMu[m]*cArg - iMu[m]*sArg;
        iFhDn_r += iMu[m]*cArg + rMu[m]*sArg;
    }
    rFhD[n] = rFhDn_r; iFhD[n] = iFhDn_r;
}
```

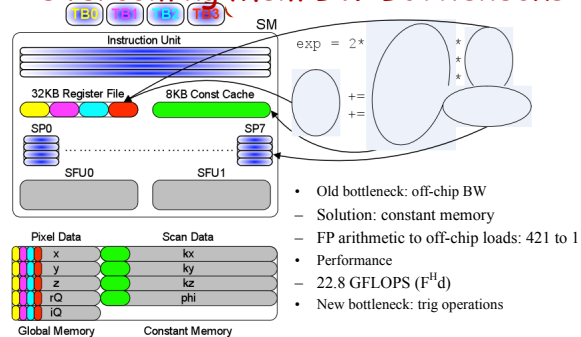
Figure 7.16 Adjusting the k-space data memory layout in the F^Hd kernel

© David Kirk/NVIDIA and Wen-mei W. Hwu,
2007–2009
University of Illinois, Urbana-Champaign

20



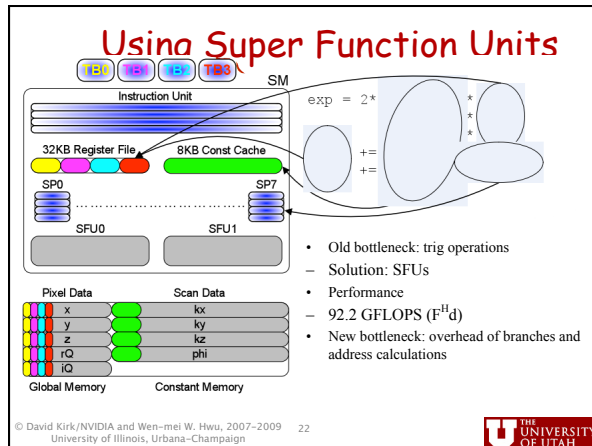
Overcoming Mem BW Bottlenecks



© David Kirk/NVIDIA and Wen-mei W. Hwu, 2007–2009
University of Illinois, Urbana-Champaign

21





Sidebar: Effects of Approximations

- Avoid temptation to measure only absolute error ($I_0 - I$)
- Can be deceptively large or small
- Metrics
 - PSNR: Peak signal-to-noise ratio
 - SNR: Signal-to-noise ratio
- Avoid temptation to consider only the error in the computed value
- Some apps are resistant to approximations; others are very sensitive

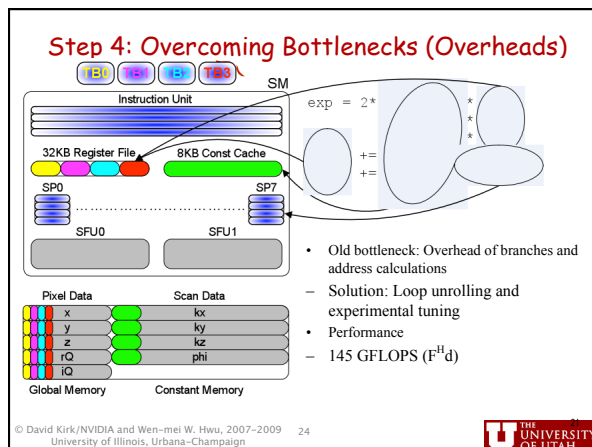
$$MSE = \frac{1}{mn} \sum_i \sum_j (I(i, j) - I_0(i, j))^2 \quad A_i = \frac{1}{mn} \sum_j I_0(i, j)^2$$

$$PSNR = 20 \log_{10} \left(\frac{\max(I_0(i, j))}{\sqrt{MSE}} \right) \quad SNR = 20 \log_{10} \left(\frac{\sqrt{A_i}}{\sqrt{MSE}} \right)$$

A.N. Netravali and B.G. Haskell, Digital Pictures: Representation, Compression, and Standards (2nd Ed), Plenum Press, New York, NY (1995).

© David Kirk/NVIDIA and Wen-mei W. Hwu, 2007-2009
University of Illinois, Urbana-Champaign

23



Summary of Results

Reconstruction	Q		F ^H d		Total	
	Run Time (m)	GFLOP	Run Time (m)	GFLOP	Linear Solver (m)	Recon. Time (m)
Gridding + FFT (CPU, DP)	N/A	N/A	N/A	N/A	N/A	0.39
LS (CPU, DP)	4009.0	0.3	518.0	0.4	1.59	519.59
LS (CPU, SP)	2678.7	0.5	342.3	0.7	1.61	343.91
LS (GPU, Naive)	260.2	5.1	41.0	5.4	1.65	42.65
LS (GPU, CMem)	72.0	18.6	9.8	22.8	1.57	11.37
LS (GPU, CMem, SFU)	13.6	98.2	2.4	92.2	1.60	4.00
LS (GPU, CMem, SFU, Exp)	7.5	178.9	1.5	144.5	1.69	3.19
	357X		228X			108X

© David Kirk/NVIDIA and Wen-mei W. Hwu, 2007-2009
University of Illinois, Urbana-Champaign

THE UNIVERSITY OF UTAH