

Specialized Sparse Matrix Formats and SpMV Kernel Tuning for GPUs

Alexander Monakov, amonakov@ispras.ru

Institute for System Programming of Russian Academy of Sciences

May 16, 2012

- Most elements are zero
→ Need specialized formats: store only non-zero values and positions
- Appear in many scientific codes, most notably CFD
- Usually need to optimize only matrix-vector multiplication:

$$y = Ax : y_i = \sum_{j:A_{ij} \neq 0} A_{ij}x_j$$

- Only two flops per non-zero element
→ Usually limited by memory bandwidth
- We consider matrices with unknown complex sparsity patterns
→ Otherwise specialized methods can do better

Develop sparse matrix storage formats with these requirements:

- Optimize only matrix-vector multiplication on GPU
 - When multiplication by transpose is required, store the transpose explicitly
- Reasonably versatile: matrices for arbitrary CFD meshes, preconditioners
 - Not purely diagonal, or with fixed number of non-zeros per row
 - Not only for blocked sparse matrices
- Not relying on ability to reorder the matrix
- For matrices that are reused many times:
 - May need complex conversion
 - May rely on run-time auto-tuning

Primary limiting factor: global memory bandwidth

- Matrix data: streaming, too large for caches
- Source vector: non-streaming random access, but usually good locality
- Destination vector: need coalesced updates

Need to design the format accordingly:

- Facilitate efficient implementation:
 - Minimize synchronization
 - Maximize coalescing
- Reduce memory footprint
 - But make space-saving tricks GPU-friendly
- Facilitate coalesced updates to the destination vector

Common Formats: CSR

$$\begin{matrix} & 0 & 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ \begin{matrix} 0 \\ 1 \\ 2 \\ 3 \\ 4 \\ 5 \\ 6 \\ 7 \end{matrix} & \left(\begin{array}{cccccccc} a & b & & & & & & \\ & c & d & e & & & & \\ & & f & & g & h & & \\ i & & & j & & k & & \\ & & & & l & & m & n \\ & o & & & & p & & q \\ & & & & & & r & \\ & & s & & & & & t \end{array} \right) \end{matrix}$$

rowptr	0	2			5			8			11			14			17	18	20	
value	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	<i>f</i>	<i>g</i>	<i>h</i>	<i>i</i>	<i>j</i>	<i>k</i>	<i>l</i>	<i>m</i>	<i>n</i>	<i>o</i>	<i>p</i>	<i>q</i>	<i>r</i>	<i>s</i>	<i>t</i>
column	0	1	1	2	3	2	4	5	0	3	5	4	6	7	1	5	7	6	2	7

Common Formats: ELLPACK

$$\begin{array}{c} 0 \\ 1 \\ 2 \\ 3 \\ 4 \\ 5 \\ 6 \\ 7 \end{array} \begin{pmatrix} 0 & 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ a & b & & & & & & \\ & c & d & e & & & & \\ & & f & & g & h & & \\ i & & & j & & k & & \\ & & & & l & & m & n \\ & o & & & & p & & q \\ & & & & & & r & \\ & & s & & & & & t \end{pmatrix}$$

$$\begin{pmatrix} a & b & 0 \\ c & d & e \\ f & g & h \\ i & j & k \\ l & m & n \\ o & p & q \\ r & 0 & 0 \\ s & t & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 & * \\ 1 & 2 & 3 \\ 2 & 4 & 5 \\ 0 & 3 & 5 \\ 4 & 6 & 7 \\ 1 & 5 & 7 \\ 6 & * & * \\ 2 & 7 & * \end{pmatrix}$$

Basic Sliced ELLPACK

$$\begin{array}{c}
 0 \\ 1 \\ 2 \\ 3 \\ 4 \\ 5 \\ 6 \\ 7
 \end{array}
 \left(
 \begin{array}{cccccccc}
 0 & 1 & 2 & 3 & 4 & 5 & 6 & 7 \\
 a & b & & & & & & \\
 & c & d & e & & & & \\
 \hline
 & & f & & g & h & & \\
 i & & & j & & k & & \\
 \hline
 & & & & l & & m & n \\
 o & & & & & p & & q \\
 \hline
 & & & & & & r & \\
 & & s & & & & & t
 \end{array}
 \right)
 \begin{array}{c}
 a \ b \ 0 \\
 c \ d \ e \\
 f \ g \ h \\
 i \ j \ k \\
 l \ m \ n \\
 o \ p \ q \\
 r \ 0 \\
 s \ t
 \end{array}$$

sliceptr	0					6						12						18			22	
value	a	c	b	d	0	e	f	i	g	j	h	k	l	o	m	p	n	q	r	s	0	t
column	0	1	1	2	*	3	2	0	4	3	5	5	4	1	6	5	7	7	6	2	*	7

A hybrid between CSR and ELLPACK:

- Exactly like CSR for slice size = 1
- Exactly like ELLPACK for infinite slice size

Allows for run-time auto-tuning:

- Slice height
- Threads for slice
- May use variable-height slices

Can do even better:

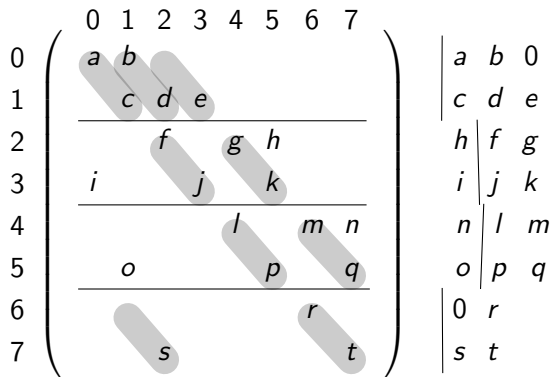
- Compact encoding for 1×2 , 1×4 dense sub-blocks
- Compact encoding for diagonal regions in slices

Sliced ELLPACK with Short Blocks

	0	1	2	3	4	5	6	7	
0	<i>a</i>	<i>b</i>							0 <i>a</i> <i>b</i>
1		<i>c</i>	<i>d</i>	<i>e</i>					<i>c</i> <i>d</i> <i>e</i>
2			<i>f</i>		<i>g</i>	<i>h</i>			<i>f</i> <i>g</i> <i>h</i>
3	<i>i</i>			<i>j</i>		<i>k</i>			<i>i</i> <i>j</i> <i>k</i>
4					<i>l</i>		<i>m</i>	<i>n</i>	<i>l</i> <i>m</i> <i>n</i>
5		<i>o</i>				<i>p</i>		<i>q</i>	<i>o</i> <i>p</i> <i>q</i>
6							<i>r</i>		<i>r</i> 0
7			<i>s</i>					<i>t</i>	<i>s</i> <i>t</i>

sliceptr	(0, 0)	(6, 4)	(12, 10)	(18, 16)	(22, 20)
value	0	<i>c</i> <i>a</i> <i>b</i> <i>d</i> <i>e</i>	<i>f</i> <i>i</i> <i>g</i> <i>j</i> <i>h</i> <i>k</i>	<i>l</i> <i>o</i> <i>m</i> <i>p</i> <i>n</i> <i>q</i>	<i>r</i> <i>s</i> 0 <i>t</i>
column	*	1 0 2 2 0 4 3 5 5 4 1 6 5 7 7 6 2 *	7		

Sliced ELLPACK with Diagonals



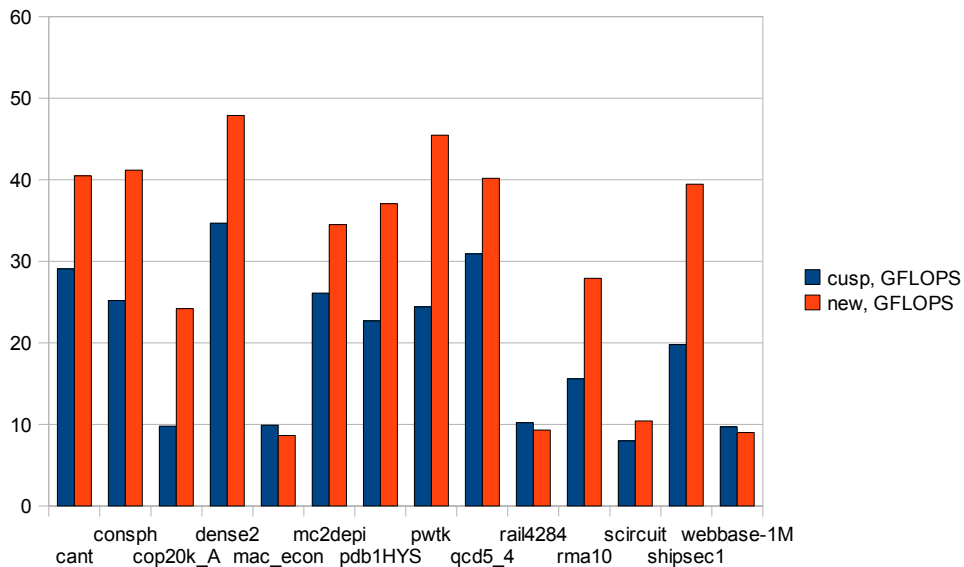
sliceptr	(0, 0, 0)	(6, 0, 3)	(12, 2, 5)	(18, 4, 7)	(22, 4, 9)
value	a c b d 0 e h i f j g k n o l p m q	0 s r t			
column		5 0	7 1		
diagoffs	0	1	2	0	2
				-5	0

Conversion time — memory footprint tradeoff

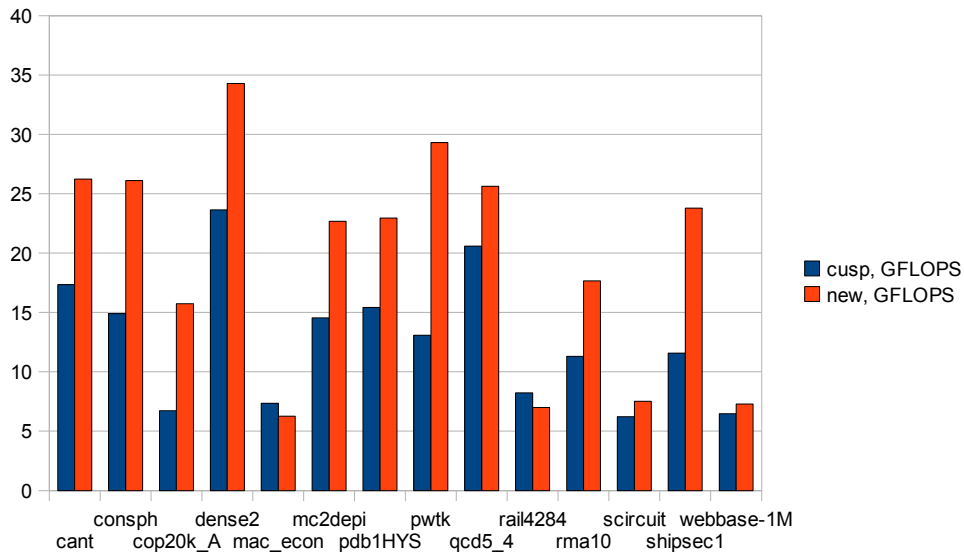
- Require additional pass over matrix data to choose blocks/diagonals
- ≈ 1 GB/s conversion speed for non-diagonal variant
- ≈ 0.2 GB/s for diagonal variant
- If the matrix is reused, cost is amortized

Auto-tuning will have to try multiple formats: can do that in a long-running CFD simulation

Performance Results: Single Precision (Tesla M2090)



Performance Results: Double Precision (Tesla M2090)



Thanks!

Memory Footprint Stats

Testcase	SpMV bytes, CSR	SpMV bytes, Proposed
cant.mtx	32M	24M
consph.mtx	49M	43M
cop20k	22M	22M
dense2.mtx	32M	24M
mac	12M	39M
mc2depi.mtx	23M	13M
pdb1HYS.mtx	35M	33M
pwtck.mtx	95M	76M
qcd5	15M	13M
rail4284.mtx	90M	140M
rma10.mtx	19M	24M
scircuit.mtx	9M	23M
shipsec1.mtx	64M	56M
webbase-1M.mtx	36M	64M