

Area, and Power Performance Analysis of a Floating-point based Application on FPGAs

Gokul Govindu, Ling Zhuo, Seonil Choi, Padma Gundala, Viktor K Prasanna

Department of Electrical Engineering, Systems

University of Southern California

Los Angeles, CA 90089

{govindu, lzhuo, seonilch, pgundala, prasanna}@usc.edu

Introduction

Almost all signal processing algorithms are initially represented as double precision floating-point in languages such as Matlab. For hardware implementations, these algorithms have to be converted to large precision fixed-point to have a sufficiently large dynamic range. However the inevitable quantization effects and the complexity of converting the floating-point algorithm into a fixed point one, limit the use of fixed-point arithmetic for high precision embedded computing.

FPGAs have become an attractive option for implementing computationally intensive applications. However, the common conception has been that efficient FPGA implementations of floating-point arithmetic have a lot of performance, area and power overheads compared to fixed-point arithmetic. With recent technology advances, FPGA densities are increasing at a rate at which area considerations are becoming less significant. These advances have also reduced the performance and power overhead of floating-point arithmetic. With appropriate designs, floating-point applications can even be more efficient than fixed-point ones for large bitwidths. The overheads in the context of the overall application can be quite low. In this paper, we present a preliminary area, and power performance analysis of double precision matrix multiplication, an extensively used kernel in embedded computing and also show that FPGAs are good candidates for implementing high precision floating-point based applications when compared to a general-purpose processor.

Currently many FPGA based floating-point units, both open source [2] and commercial [1], are available. However, most of them consider only single precision floating-point operations, and do not make use of the recent advances in FPGAs. Moreover, an area, and power performance analysis of the floating-point units in the context of a common application is lacking.

Description of our Floating point units and the Matrix Multiply architecture

For matrix multiplication, we require add and multiply floating-point units. Our floating-point units follow the IEEE 754 single and double precision (64-bit) format. We developed both deeply pipelined and moderately pipelined units. The units essentially consist of three stages: denormalization, the add/multiply, and normalization/rounding/renormalization. Exception handling at all stages is done and enable/done signals are provided for easy integration into a pipelined architecture. The implementation of floating-point units involves extensive use of fast fixed point adder/subtractors, multiplier units, and large bus multiplexers (for shifting operations). Recent FPGAs, such as Virtex-II Pro [4], provide a large number of embedded multipliers as well as fast carry chains for addition. Similarly, large multiplexers used in shifting can make use of the MUXCY, MUXF attributes on the FPGAs. Recent FPGA fabrics also contain a lot of registers, which can be utilized for extensive pipelining between stages.

We used the block matrix multiplication architecture from [3] in which a linear array of n processing elements is used for an $n \times n$ matrix multiplication. Each processing elements essentially consists of an adder, a multiplier, storage elements, and related control logic. Since the matrix multiply architecture (see [3] for more details) is modular, multiple chips can be used in an array for large n . Here we use the GFLOPS per device for a given n as the performance metric.

Analysis of the Floating-point units

Table 1a and 1b show a comparison of the fixed and floating-point units for a bitwidth of 32 and 64. We see that the overhead for double precision is less than that for single precision. Note that, for the fixed-point designs, truncation to make the output bitwidth equal to the input bitwidth results in a lot of quantization error. Moreover the fixed-point multiplier unit takes up more embedded multipliers than the floating-point unit. We also show a comparison between an extensively pipelined and a moderately pipelined version of the floating-point units. We see that extensive pipelining to increase the clock frequency requires a lot of area for the registers in between the pipeline stages. The pipelining done to split the adder/multiplier, the large priority encoder and the shift registers for the normalizing unit shows an immediate improvement in frequency, without much increase in area. Further pipelining, shows diminishing returns in frequency and the area increases significantly. Hence a design trade-off will be the frequency required which influences the number of pipelining stages and area. Here, for the double precision matrix multiply, we decided to use the moderately pipelined units since we can achieve higher GFLOPs. From synthesis results, we saw that normalization takes up a lot of area (560 slices for the deeply pipelined and 200 slices for the moderately pipelined units, for double precision) and can also be the critical path for timing (because of a large priority encoder and shift registers). Hence a design trade-off would be the use of custom formats in the architecture, with conversion from and back to the IEEE754 standard at the interface to say, a processor. Considering power, the 64bit fixed-point multiplier unit with more embedded multipliers consumes a lot more power. Note that, for the power values of individual units, only clocks, logic and signal powers were included.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 20 AUG 2004		2. REPORT TYPE N/A		3. DATES COVERED -	
4. TITLE AND SUBTITLE Area, and Power Performance Analysis of a Floating-point based Application on FPGAs				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Department of Electrical Engineering, Systems University of Southern California Los Angeles, CA 90089				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release, distribution unlimited					
13. SUPPLEMENTARY NOTES See also ADM001694, HPEC-6-Vol 1 ESC-TR-2003-081; High Performance Embedded Computing (HPEC) Workshop (7th)., The original document contains color images.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT UU	18. NUMBER OF PAGES 34	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

	32, 64bit Fixed-point (with 2, 4 pipeline stages)		32, 64bit Floating-point (with 6, 8 pipeline stages)		32, 64bit Floating-point (with 18, 23 pipeline stages)	
Area (slices)	36	139	293	693	504	1383
Max Freq. (MHz) achievable	250	250	140	130	250	200
Power (mW) at 100MHz	23.48	104	148.7	329	-	-

Table 1a: A comparison of addition units (Virtex2Pro-c2vp125-7)

	32, 64bit Fixed-point (with 5, 7 pipeline stages)		32, 64bit Floating-point (with 9, 11 pipeline stages)		64bit Floating-point (with 18, 23 pipeline stages)	
Area (slices)/Embedded multipliers	190 / 4	1024 / 16	249 / 3	775 / 10	492 / 3	1558 / 10
Max Freq. (MHz) achievable	200	130	140	130	200	200
Power (mW) at 100MHz	136.3	804	164.7	424	-	-

Table 1b: A comparison of multiplication units (Virtex2Pro-c2vp125-7)

Analysis of the Matrix Multiply

The area and power performance overhead of the floating-point units has to be seen in the context of an application. Table 2 shows the area and power performance of both fixed and floating-point implementations of a double precision, n -point matrix multiply on a FPGA. The double precision implementation shows us an interesting result of the floating-point unit having a better performance than the fixed-point implementation. The maximum number of fixed-point processing elements that the device can accommodate when block RAMs are used for storage, is smaller than the number of processing elements when slice based RAM is used. This is probably because of more routing resources used up due to the fixed locations of the block RAMs and the embedded multipliers. Moreover, the number of slices on a given device being constant, the device will accommodate fewer processing elements if deeply pipelined units occupy a large area. Hence, the performance of the device might be lower even if the frequency of the units is high. Also, the overall application's architecture's operating frequency should be considered. Performance was measured as one multiplication and one addition happening every clock cycle in each processing element. The total power for the matrix multiply takes into account output, input, quiescent, logic, signals and the clocks power. We see that floating-point unit overheads in terms of area, and power performance are not too drastic. Table 3 shows the performance comparison of a floating-point based n -point matrix multiplication both on an FPGA and a Pentium4 SSE2, 1.5GHz processor. The performance of the design on FPGAs shows a 3.48x improvement over that of the processor. Moreover the power per GFLOP of the FPGA is much lower than that of the processor.

	Fixed-point based		Floating point based (moderately pipelined) using block RAM	Floating point based (deeply pipelined) (estimated)
	using block RAM	using slice based RAM		
Area (slices) / BRAM / multipliers of each Processing element of matrix multiply	1344 / 4 / 16	1626 / 0 / 16	1872 / 4 / 10	3441 / 4 / 10
Maximum number of processing elements on the device	28	32	29	16
Frequency (MHz) of each element	130	130	130	200
Power of each PE (mW) at 100MHz	843	894	762	-
Frequency (MHz) achieved for the matrix multiply	110	110	120	200
Performance of matrix multiply, per device-Virtex2Pro xc2vp125-7	6.16 GOPS	7.04 GOPS	6.96 GFLOPS	6.4 GFLOPS
Total Power (W) per GOP or GFLOP for matrix multiply,	27.2 / 6.16 = 4.41	33.04 / 7.04 = 4.68	26.4 / 6.96 = 3.79	-

Table 2: A comparison of double precision, fixed and floating-point, n -point Matrix Multiply, requiring n processing elements on FPGAs (Virtex2Pro-c2vp125-7)

	FPGA (Virtex2Pro xc2vp125-7)	Pentium4 with SSE2 (1.5GHz)
GFLOPS	6.96	2
Power (W) per GFLOP	$26.4 / 6.96 = 3.79$	$57.9 / 2 = 28.95$

Table 3: A comparison of the performance of Matrix Multiply

All the above results were obtained after the VHDL code was synthesized and placed and routed using the Xilinx ISE5.2i, on a Virtex2Pro XC2VP125-7f1696 device. Power values were obtained from Xpower. The Pentium4 SSE2 results were from [5]. Better results can be obtained after the units have been optimized more, by manually placing them.

Conclusion and Future Work

We have presented a preliminary analysis of a floating-point implementation of a computationally intensive application on FPGAs. We show that when the floating-point units are considered in the context of an application, their overheads in terms of area, and power performance are not too drastic. We also show that a significant increase in performance can be obtained on FPGAs over general-purpose processors with much lower power expended. Future work will involve extensive analysis of the floating-point units to identify more design trade-offs. We will also provide a documented and extensively tested, open source library of the floating-point units, shortly.

References

1. Nallatech, www.nallatech.com
2. P. Belanović, M. Leeser, *Library of Parameterized Floating-point Modules and Their Use*, International Conference on Field Programmable Logic (ICFPL), September 2002.
3. J. Jang, S. Choi, and V. K. Prasanna, *Area and Time Efficient Implementation of Matrix Multiplication on FPGAs*, International Conference on Field Programmable Technology (ICFPT), December 2002.
4. Xilinx, www.xilinx.com
5. J. Dongarra. *An Update of a Couple of Tools: ATLAS and PAPI*. Technical report, DOE Salishan Meeting, April 2001.

Area and Power Performance Analysis of Floating-point based Applications on FPGAs

**Gokul Govindu, Ling Zhuo, Seonil Choi, Padma Gundala,
and Viktor K. Prasanna**

**Dept. of Electrical Engineering
University of Southern California
September 24, 2003**

<http://ceng.usc.edu/~prasanna>

Outline

- **Floating-point based Applications on FPGAs**
- **Floating-point Units**
 - Area/Power Analysis
- **Floating-point based Algorithm/Architecture Design**
- **Area, Power, Performance analysis for example kernels:**
 - FFT
 - Matrix Multiply
- **Conclusion**

Floating-point based Applications on FPGAs

Applications requiring

- High numerical stability, faster numerical convergence
- Large dynamic range

Examples:

- Audio/Image processing, Radar/Sonar/Communication, etc.

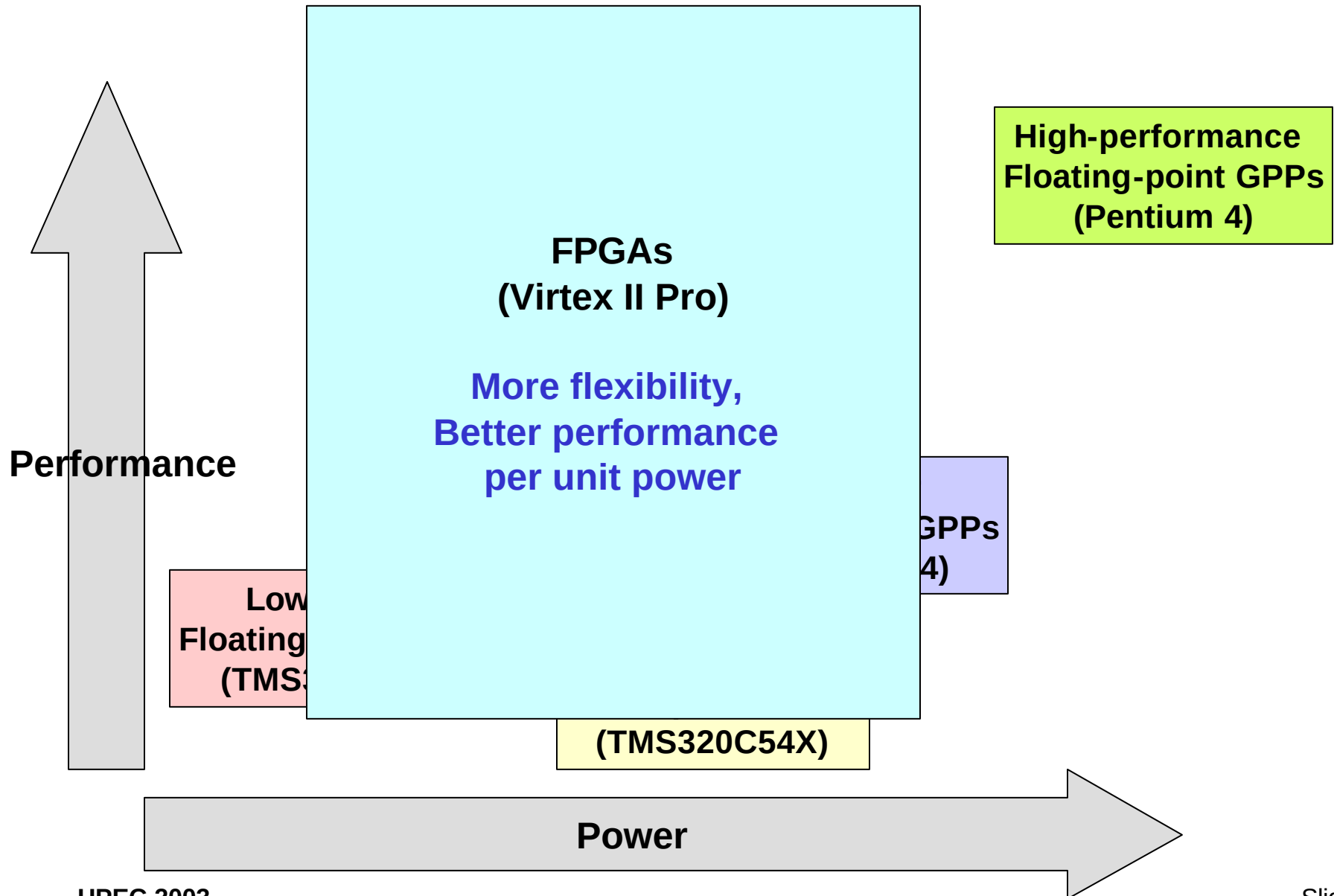
Fixed-point vs. Floating-point

- Resources
 - Slices
- Latency/Throughput
 - Pipeline stages
 - Frequency
- Precision
- Design complexity of fixed/floating-point units

Energy – Area – Performance

Tradeoffs

Floating-point Device Options



Need for FPU Design in the Context of the Kernel

Integration

- **Latency**
 - Number of pipeline stages as a parameter
- **Frequency**
 - FPU frequency should match the frequency of the kernel/application's logic
- **Area/Frequency/Latency tradeoffs**

Optimal Kernel Performance

- **High throughput**
 - Maximize frequency
- **Minimize Energy**
 - Architectural tradeoffs - FPUs parameterized in terms of latency/ throughput/ area
- **Optimize F/A for FPU**
 - Maximize the performance of the kernel

Algorithm/Architecture Design

- **Re-evaluation of the algorithm/architecture**
 - Tolerate latencies of FPU - low area vs. high frequency tradeoffs
 - Re-scheduling

Outline

- Floating-point based Applications on FPGAs
- **Floating-point Units**
 - Area/Power Analysis
- Floating-point based Algorithm/Architecture Design
- Area, Power, Performance analysis for example kernels:
 - FFT
 - Matrix Multiply
- Conclusion

Our Floating-point Units

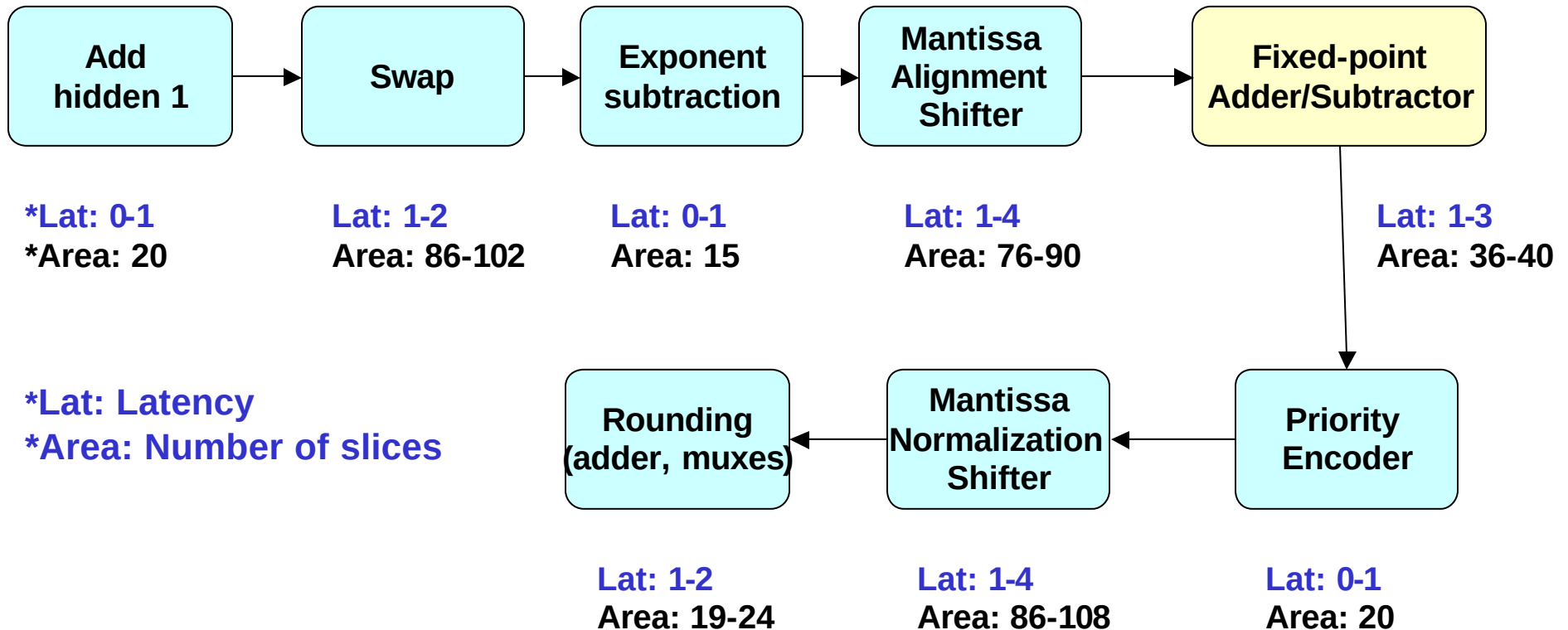
- **Now, easier to implement floating-point units on FPGAs**
 - Optimized IP cores for fixed-point adders and multipliers
 - Fast priority encoders, comparators, shift registers, fast carry chains....

Our floating-point units

- **Precision**
 - Optimized for 32, 48 and 64 bits
- **IEEE 754 format**
- **Number of pipeline stages**
 - Number of pipeline stages parameterized
 - For easy integration of the units into the kernel
 - For a given kernel frequency, units with optimal pipelining and thus optimal resources, can be used
- **Metrics**
 - Frequency/Area
 - Overall performance of the kernel (using floating-point units)
 - Energy

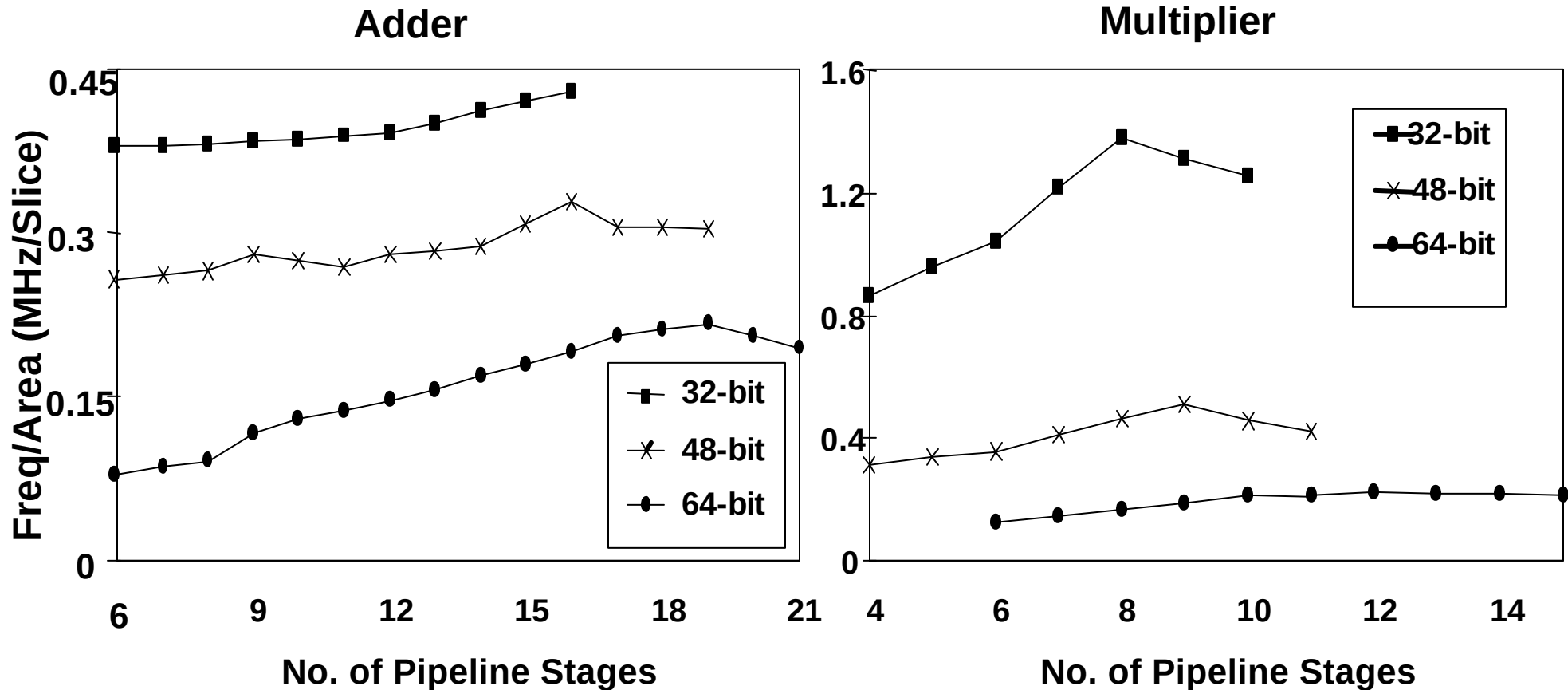
Floating-point Adder/Subtractor

32 bits Precision



- Pipeline stages: 6-18
- Area: 390- 550; Achievable frequency: 150-250MHz
- Xilinx XC2VP125 –7

Frequency/ Area vs. Number of Pipeline Stages



- Diminishing returns beyond optimal F/A
- Tools' optimization set as "balanced - area and speed"
 - Area and Speed optimization give different results in terms of area and speed

Addition Units: Some Trade-offs

	Fixed-point		Floating-point		Floating-point	
	32 bits with 2 stages	64 bits with 4 stages	32 bits with 14 stages	64 bits with 19 stages	32 bits with 19 stages	64 bits with 21 stages
Area(slices)	36	139	485	933	551	1133
Max. Freq. (MHz) achievable	250	230	230	200	250	220
Power(mW) at 100MHz	23.48	102	200	463	254	529

Floating-point vs. Fixed-point

- Area : 7x-15x
- Speed: 0.8x-1x
- Power: 5x-10x

Multiplier Units: Some Trade-offs

	Fixed-point		Floating-point		Floating-point	
	32 bits with 5 stages	64 bits with 7 stages	32 bits with 7 stages	64 bits with 10 stages	32 bits with 10 stages	64 bits with 15 stages
Area(slices)/Embed ded Multipliers	190/4	1024/16	180/3	838/10	220/3	1019/10
Max. Freq. (MHz) Achievable	200	130	220	175	220	215
Power(mW) at 100MHz	136.3	414	227	390	263	419

Floating-point vs. Fixed-point

- Area : 0.9x-1.2x
- Speed: 1.1x-1.4x
- Power: 1x-1.6x

A Comparison of Floating-point units

Our units vs. the units from the **NEU** library*

	USC 32 bits			NEU 32 bits			USC 64 bits			NEU 64 bits		
	F	A	F/A	F	A	F/A	F	A	F/A	F	A	F/A
Adder	250	551	.45	120	391	.35	200	933	.22	50	770	.07
Multiplier	250	182	1.4	95	124	0.6	205	910	.23	90	477	.18

F: Frequency

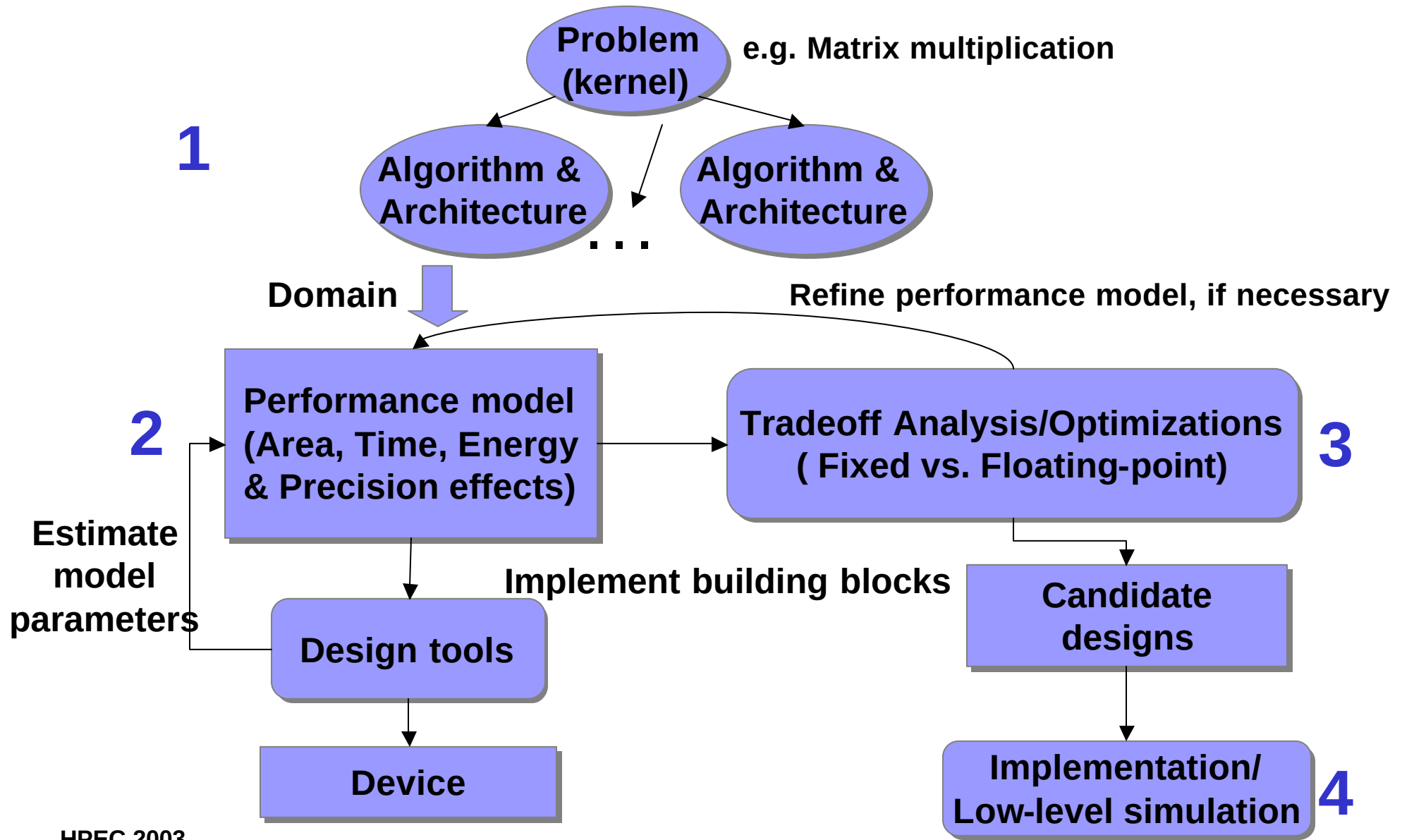
A: Slices

* P. Belanovic, M. Leeser, *Library of Parameterized Floating-point Modules and Their Use*, International Conference on Field Programmable Logic (ICFPL), Sept., 2002

Outline

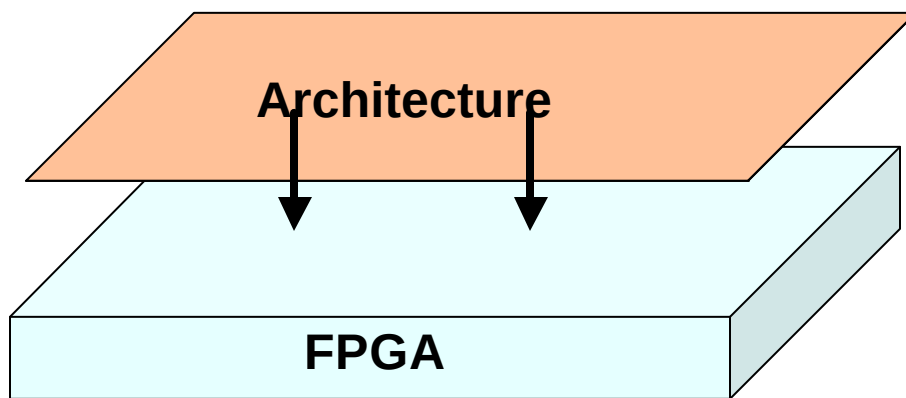
- Floating-point based Applications on FPGAs
- Floating-point Units
 - Area/Power Analysis
- Floating-point based Algorithm/Architecture Design
- Area, Power, Performance analysis for example kernels:
 - FFT
 - Matrix Multiply
- Conclusion

The Approach: Overview



1. Domain

- **FPGA is too fine-grained to model at high level**
 - No fixed structure comparable to that of a general purpose processor
 - Difficult to model at high level
- **A family of architectures and algorithms for a given kernel or application**
 - E.g. matrix multiplication on a linear array
- **Imposes an architecture on FPGAs**
 - Facilitates high-level modeling and high-level performance analysis



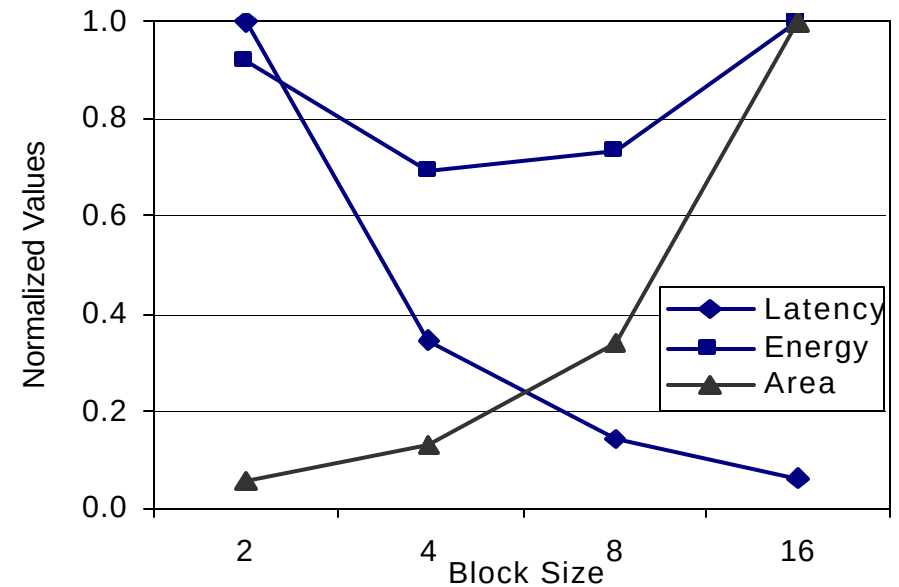
- **Choose domains by analyzing algorithms and architectures for a given kernel**
 - Tradeoffs in Area, Energy, Latency

2. Performance Modeling

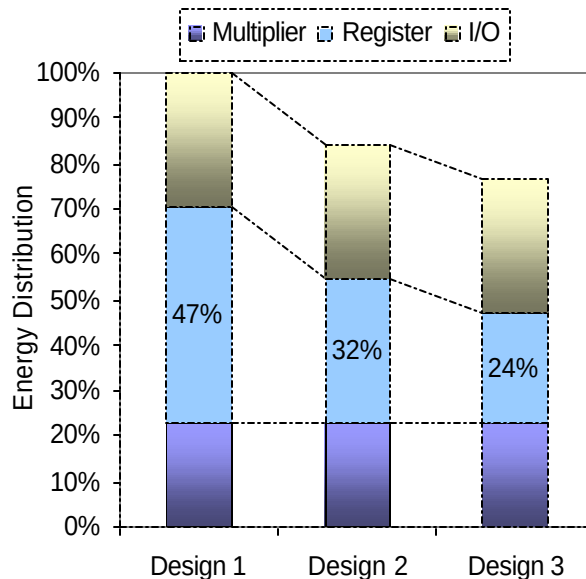
- **Domain Specific Modeling**
- **High-level model**
 - Model parameters are specific to the domain
 - Design is composed based on the parameters
 - Design is abstracted to allow easier (but coarse) tradeoff analysis and design space exploration
 - Precision effects are studied
 - Only those parameters that make a significant impact on area and energy dissipation are identified
- **Benefit:** Rapid evaluation of architectures and algorithms without low-level simulation
 - Identify candidate designs that meet requirements

3. Tradeoff Analysis and Manual Design Space Exploration

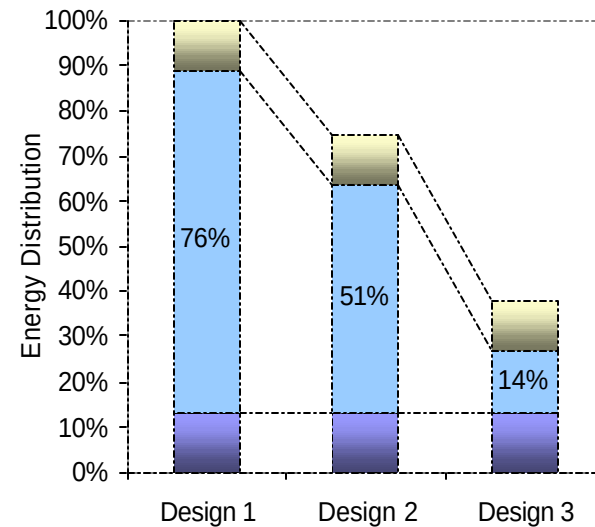
- Vary model parameters to see the effect on performance
- Analyze tradeoffs
- Weed out designs that are not promising



Example: Energy Tradeoffs



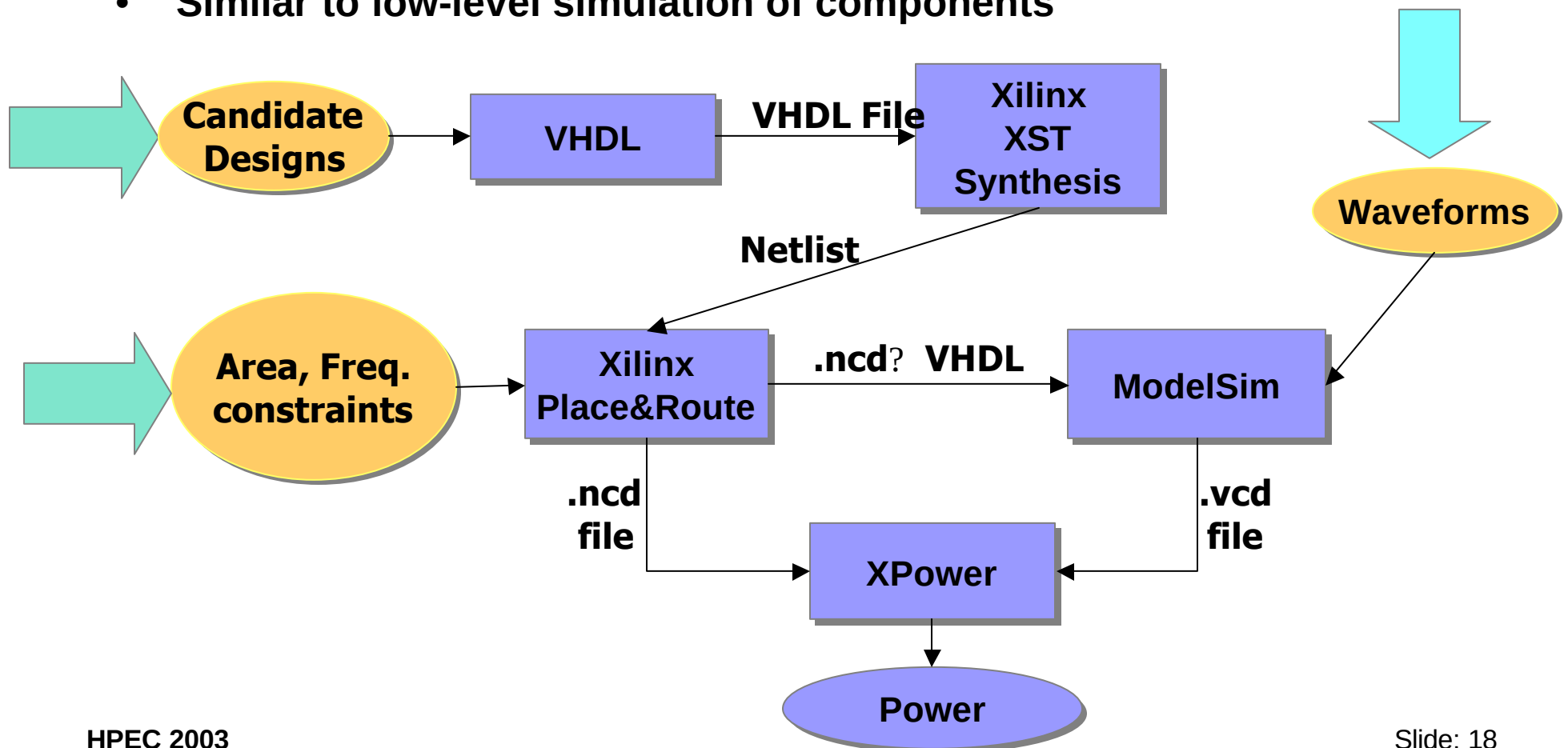
(a) 3x3



(b) 12x12

4. Low Level Simulation of Candidate Designs

- Verify high-level estimation of area and energy for a design
- Select the best design within the range of the estimation error among candidate designs
- Similar to low-level simulation of components

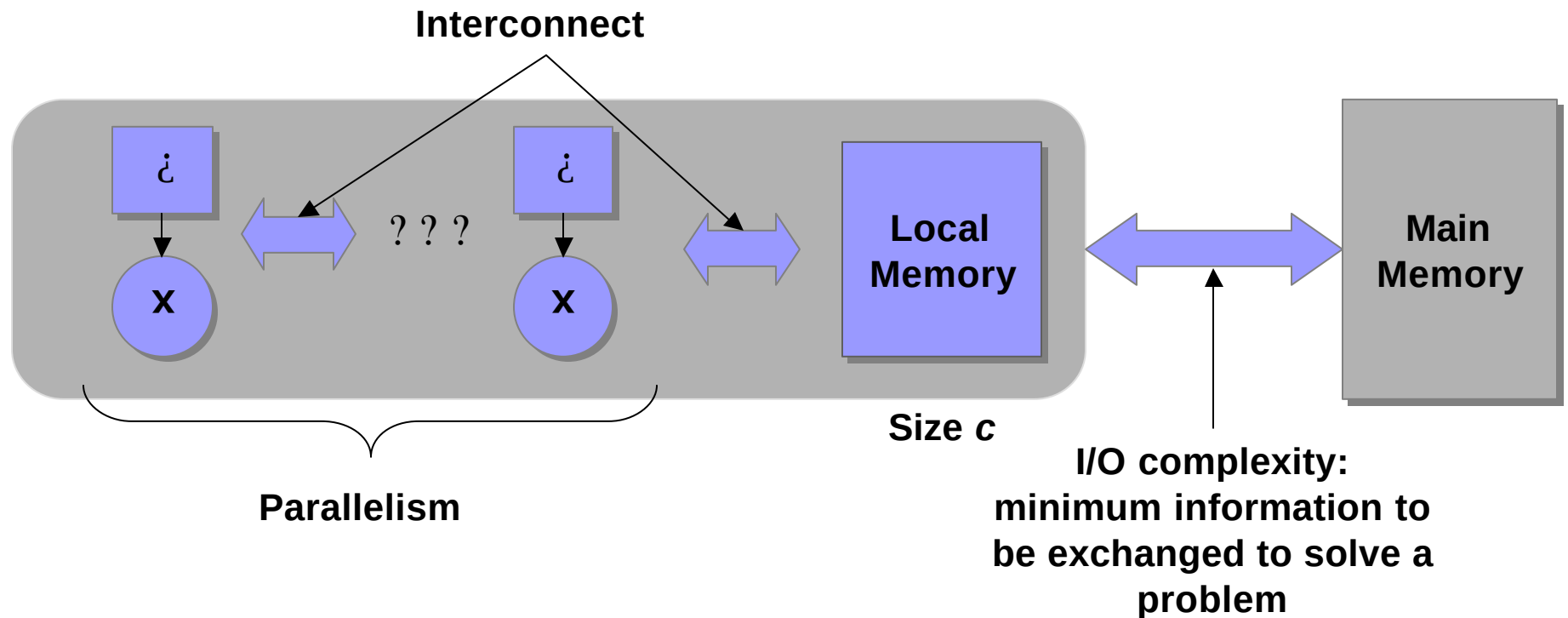


Outline

- **Floating-point based Applications on FPGAs**
- **Floating-point Units**
 - **Area/Power Analysis**
- **Floating-point based Algorithm/Architecture Design**
- **Area, Power, Performance analysis for example kernels:**
 - **FFT**
 - **Matrix Multiply**
- **Conclusion**

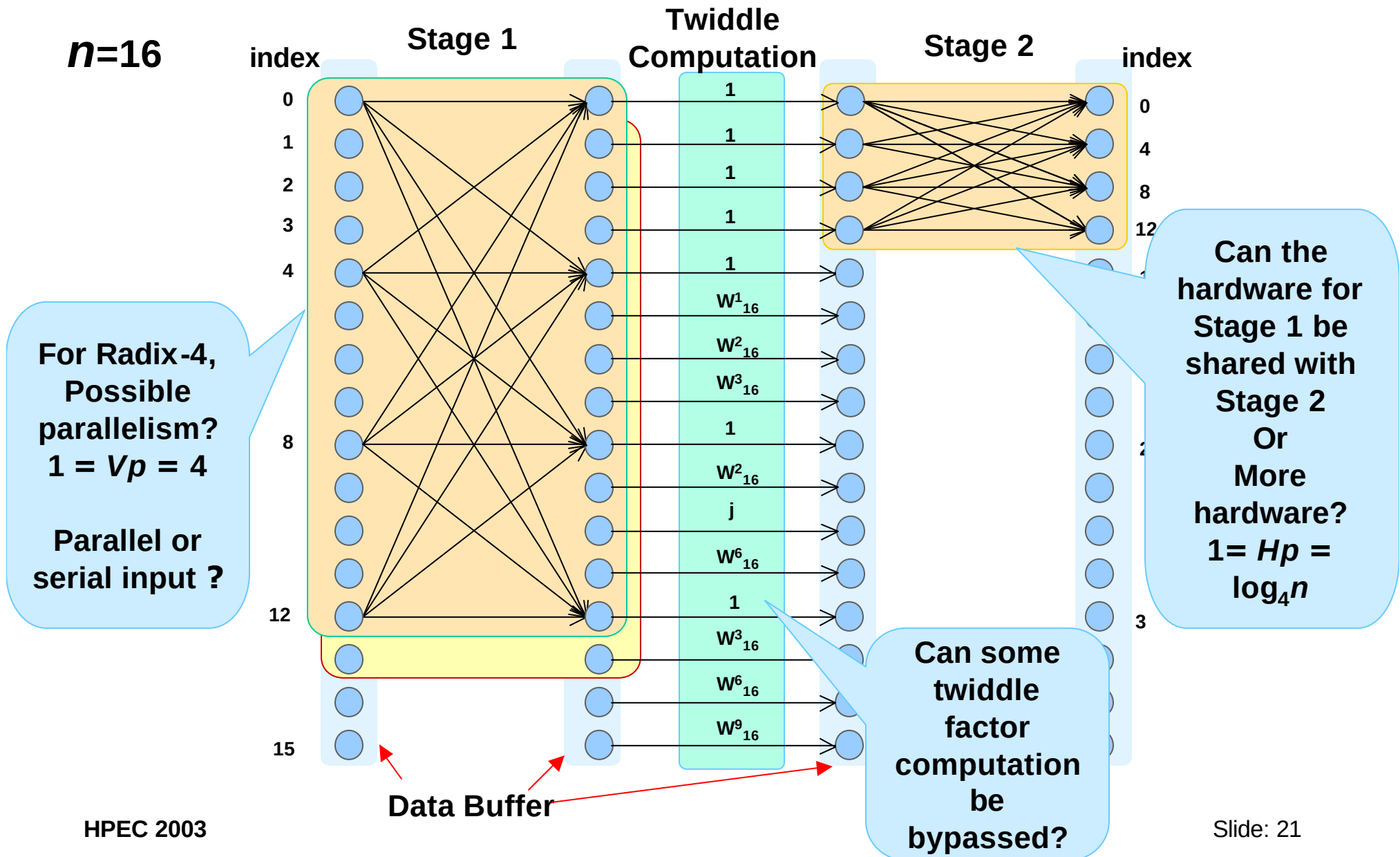
Example 1: FFT Architecture Design Tradeoffs

n -point FFT



For n -point FFT, I/O complexity = ? ($n \log n / \log c$)

FFT Architecture Design Tradeoffs (2)

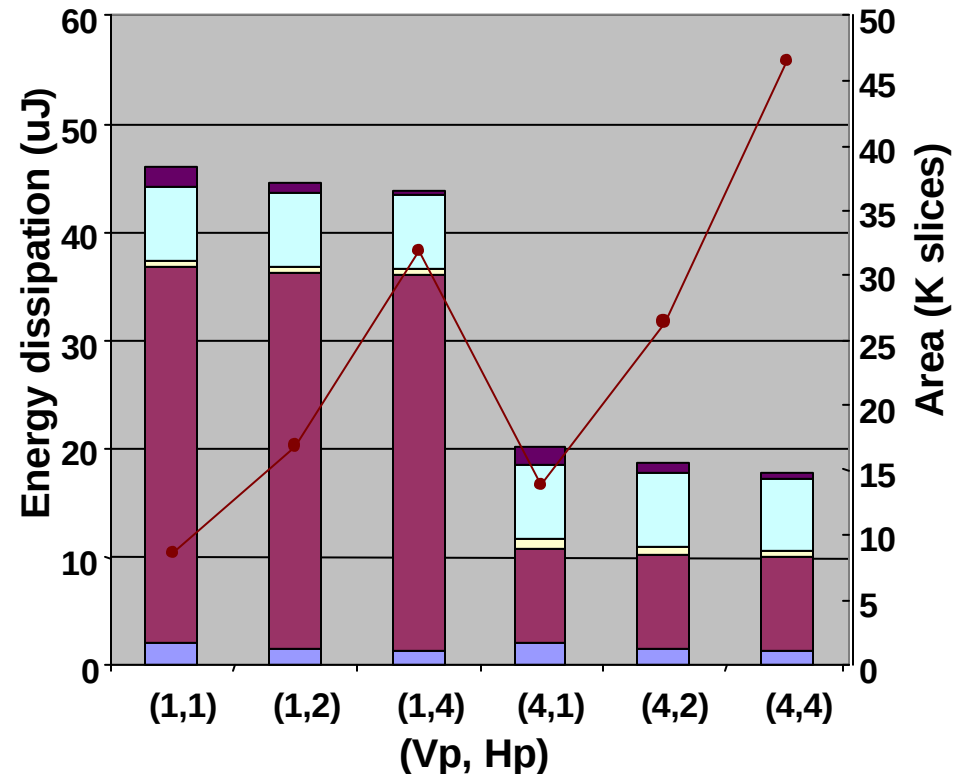
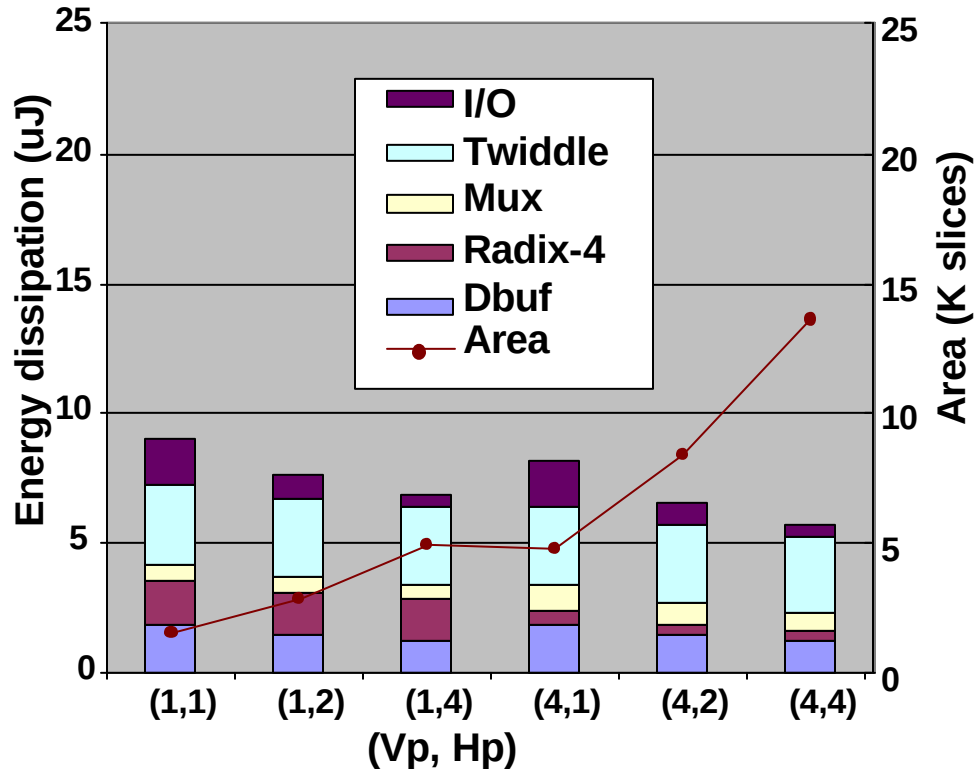


FFT Architecture Design Trade-offs (3)

Fixed-point

256 Point FFT (32 bits)

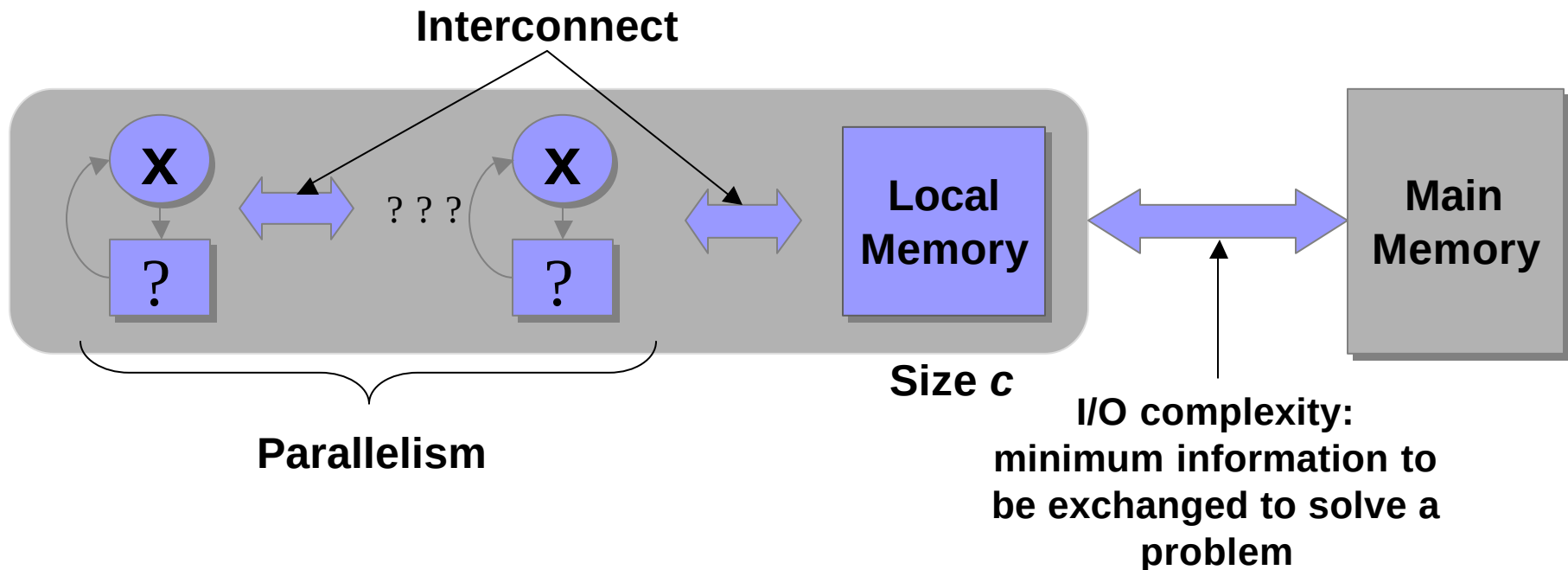
Floating-point



- Optimal FFT architectures with respect to EAT
 - Fixed-point: (Vp, Hp) = (1,4)
 - Floating-point: (Vp, Hp) = (4,1)

Example 2: Matrix Multiplication Architecture Design (1)

I/O Complexity of Matrix Multiplication

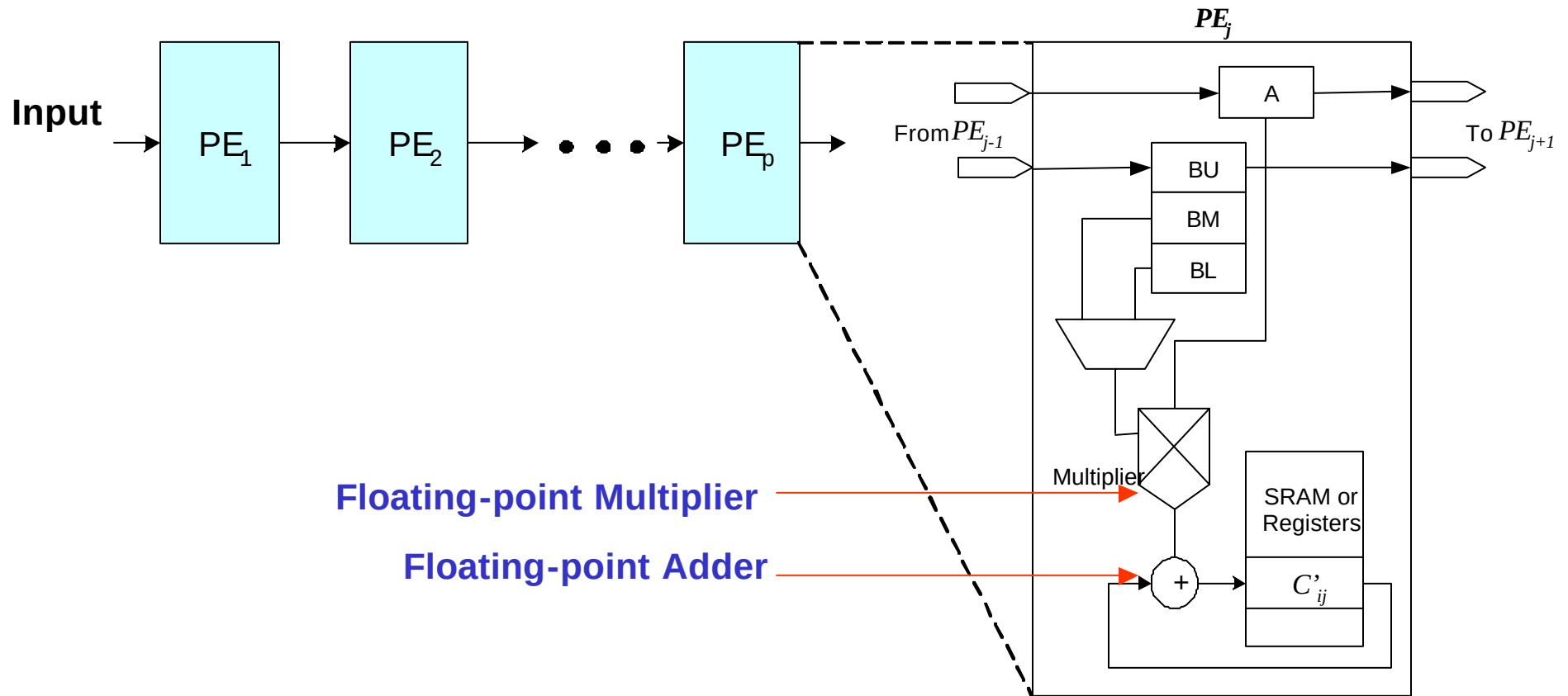


Theorem (Hong and Kung): For $n \times n$ matrix multiplication

$$\text{I/O complexity} = \Theta(n^3/c)$$

Matrix Multiplication Architecture Design (2)

Processing Element Architecture*



* J. W. Jang, S. Choi, and V. K. Prasanna, *Area and Time Efficient Implementation of Matrix Multiplication on FPGAs*, ICFPT 2002.

Matrix Multiplication Architecture Design (3)

- **Our design**
 - **Number of PEs = n**
 - **Storage = ? ($n \times n$)**
 - **Latency = ? (n^2)**
- **For $n \times n$ matrix multiplication, I/O complexity = ? (n^3/c)**
- **Our design has optimal I/O complexity**

Performance of 32, 64 bits Floating-point Matrix Multiplication (4)

Pipeline stages	32 bits XC2VP125 –7			64 bits XC2VP125 –7		
	Min	Max	Optimal	Min	Max	Optimal
Area(slices) of each Processing Element	718	991	933	1524	2575	2256
Max. No. PEs	77	56	59	36	21	24
Achievable Frequency (MHz)	90	215	210	50	190	180
Sustained Performance (GFLOPS)	13.8	24.1	24.7	3.6	8.0	8.6

The performance (in GFLOPS) is maximum for the design with floating-point units with maximum frequency/area.

FPGA vs. Processor

32 bits floating-point matrix multiplication on FPGA using our FPU and architecture

	FPGA XC2VP125 –7 230MHz	TI TMS320 C6713* 225 MHz	Analog TigerSharc * 500 MHz	Pentium 4 SSE2 * 2.53 GHz	PowerPC G4 * 1.25 GHz
GFLOPS	24.7 (sustained)	1.325 (peak)	1.0 (peak)	6.56 (peak)	6.22 (peak)
Power(W)	26	1.8 (core power)	2.4 (core power)	59.3	30
GFLOPS/W	0.95	0.7	0.4166	0.11	0.2

FPGA vs. Processor

- Performance (in GFLOPS): up to 24.7x
- Performance/Power (in GFLOPS/W): up to 8.6x

* From data sheets

FPGA vs. Processor

64 bits floating-point matrix multiplication on FPGA using our FPU and architecture

	FPGA XC2VP125 –7 200MHz	Pentium 4 SSE2 1.5 GHz*	AMD Athlon 1 GHz*
GFLOPS	8.6 (sustained)	2.0 (peak)	1.1 (peak)
Power(W)	26	54.7	60
GFLOPS/W	0.33	0.036	0.018

FPGA vs. Processor

- Performance (in GFLOPS): up to 7.8x
- Performance/Power (in GFLOPS/W): up to 18.3x

* From data sheets

Conclusion and Future Work

Conclusion

- Floating-point based implementations are not prohibitively expensive either in terms of area or latency or power
- High performance kernels can be designed with appropriate FPU's
- In terms of GFLOPS and GFLOPS/W, FPGAs offer significant over general purpose processors and DSPs

Future Work

- Floating-point based beamforming....
- Tool for automatic integration of FPU's into kernels

<http://ceng.usc.edu/~prasanna>

MILAN for System-Level Design: Design Flow

Model PARIS kernels,
end-to-end
application, hardware
choices, mission
parameters, etc.

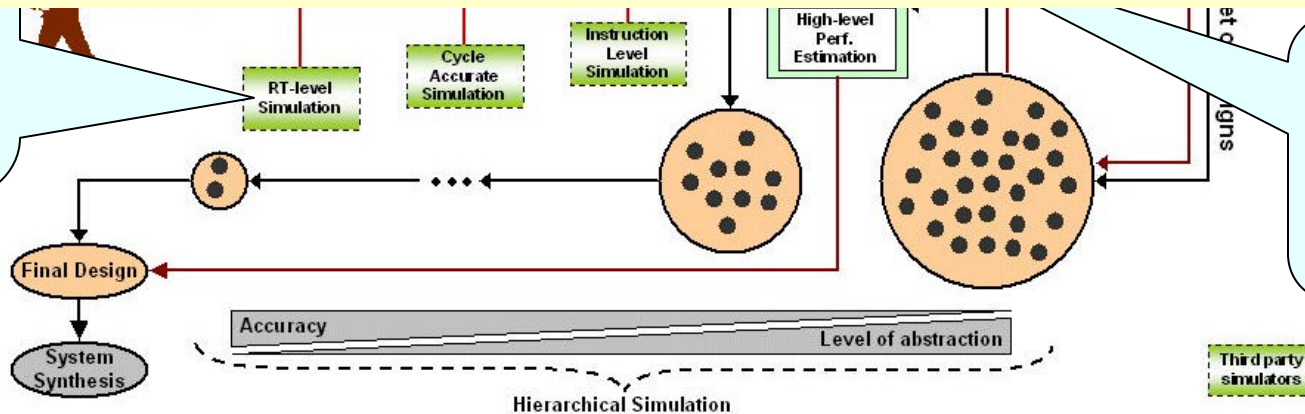
PARIS design space

Dynamic programming
based heuristics
Multi-rate application
optimization
Interval arithmetic



Download-<http://www.isis.vanderbilt.edu/Projects/milan/>

VHDL and C
implementations
Energy, latency,
and area estimates



Enhanced
HiPerE
High-level
estimator
for FPGAs

Questions?

<http://ceng.usc.edu/~prasanna>