



PROGRAMMING AND OPTIMIZATION FOR INTEL[®] ARCHITECTURE

The Hands-On Workshop (HOW) Series
Session 1

Colfax International — colfaxresearch.com

September 2016

While best efforts have been used in preparing this training, Colfax International makes no representations or warranties of any kind and assumes no liabilities of any kind with respect to the accuracy or completeness of the contents and specifically disclaims any implied warranties of merchantability or fitness of use for a particular purpose. The publisher shall not be held liable or responsible to any person or entity with respect to any loss or incidental or consequential damages caused, or alleged to have been caused, directly or indirectly, by the information or programs contained herein. No warranty may be created or extended by sales representatives or written sales materials.

▶ HOW to Program Intel Architecture

- 01. Parallelism, specialization, guided tour – Sep 26
- 02. Programming Intel Xeon Phi (KNC, KNL) – Sep 27

▶ HOW to Express Parallelism

- 03. Automatic vectorization – Sep 28
- 04. Multi-threading with OpenMP – Sep 29

▶ HOW to Get Performance

- 05. Comprehensive demo – Sep 30
- 06. Scalar & vectorization tuning – Oct 3
- 07. Multi-threading: common issues – Oct 4
- 08. Multi-threading: memory aspect – Oct 5
- 09. Memory traffic – Oct 6

▶ HOW to Scale

- 10. Distributed Computing: MPI – Oct 7

September 2016

S	M	T	W	H	F	S
				1	2	3
4	5	6	7	8	9	10
11	12	13	14	15	16	17
18	19	20	21	22	23	24
25	26	27	28	29	30	

October 2016

S	M	T	W	H	F	S
						1
2	3	4	5	6	7	8
9	10	11	12	13	14	15
16	17	18	19	20	21	22
23	24	25	26	27	28	29
30	31					

— Webinar+remote access

Course page: colfaxresearch.com/how-16-09

- ▶ Slides (including this one), code downloads
- ▶ Video of recorded sessions
- ▶ Chat (during webinars or offline)



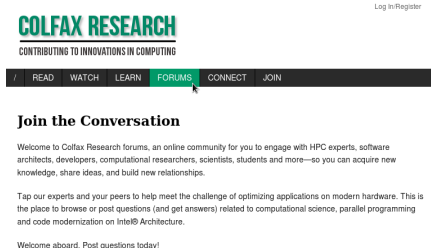
Additional resources:

- ▶ More workshops like this one: colfaxresearch.com/training
- ▶ Video courses: colfaxresearch.com/video-courses

Chat (current):
colfaxresearch.com/how-16-09



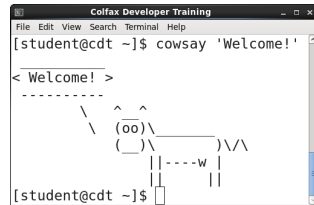
Forums (technical):
colfaxresearch.com/discussion



Email (organizational):
training@colfax-intl.com

HANDS-ON EXERCISES AND REMOTE ACCESS

- ▶ 96 people receive a remote access token
- ▶ Virtualized Intel Xeon CPU, real Intel Xeon Phi coprocessor (1st gen, KNC), SW tools
- ▶ Can access the system the entire 2 weeks of the workshop



- ▶ Not among the 96? Stay tuned: follow along with instructor, use own system, or wait for a seat
- ▶ Use it or lose it: if you do not log in for a while, remote access token goes to next student on the list

A solid green vertical bar is positioned on the left side of the image, extending from the top to the bottom.

LEARN MORE

HOW SERIES "KNIGHTS LANDING":

PROGRAMMING AND OPTIMIZATION FOR
INTEL XEON PHI X200 FAMILY

Free 2-hour video course



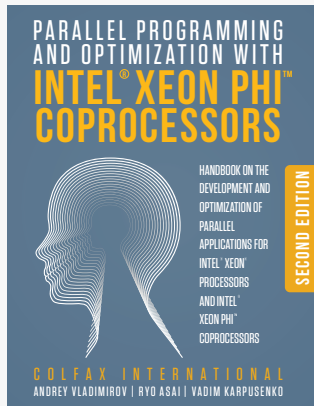
colfaxresearch.com/how-knl/

ISBN: 978-0-9885234-0-1 (508 pages, Electronic or Print)

Parallel Programming and Optimization with Intel® Xeon Phi™ Coproprocessors

Handbook on the Development and
Optimization of Parallel Applications
for Intel® Xeon® Processors
and Intel® Xeon Phi™ Coprocessors

© Colfax International, 2015



<http://xeonphi.com/book>

COLFAX RESEARCH
CONTRIBUTING TO INNOVATIONS IN COMPUTING

[/](#) [READ](#) [WATCH](#) [LEARN](#) [CONNECT](#) [JOIN](#)

Introduction to Intel DAAL, Part 1: Polynomial Regression with Batch Mode Computation

To search, type and hit enter

Popular

The Hands-On Tutorials (HOT) webinars: details on efficient programming for Intel architecture

The Hands-On Workshop (HOW) Series

Introduction to Intel DAAL, Part 1: Polynomial Regression with Batch Mode Computation

Parallel Programming Book

Introduction to parallel programming, deep discussion of optimization techniques, exercises.

© 2015, Colfax International, 508 pages.

Featured Video

Use Research material over reconstruction to a bleeding code

Optimization Techniques for the Intel MIC Architecture, Part 1 of 3: Strip-Mining for Vectorization

Performance to Power and Performance to Cost Ratios with Intel Xeon Phi Coprocessors (and why its Acceleration May Be Enough)

Research and Educational Publications

Introduction to Intel DAAL, Part 1: Polynomial Regression with Batch Mode Computation

Optimization Techniques for the Intel MIC Architecture, Part 2 of 3: False Sharing and Padding

Software Developer's Introduction to the HGST Ultrastar Archive H800 SMR Drives

Optimization Techniques for the Intel MIC Architecture, Part 2 of 3: Strip-Mining for Vectorization

Performance to Power and Performance to Cost Ratios with Intel Xeon Phi Coprocessors (and why its Acceleration May Be Enough)

Examples

Documentation

Cardinal

[Log In/Out](#) or [Register](#)

Consulting

[t](#) [w](#) [v](#) [+](#) [Share](#)

Colfax offers consulting services for enterprises, research help to:

- Optimize your existing application to take advantage of parallelism, from vectors to cores to clusters and
- Future-proof your application for upcoming innovations
- Accelerate your application using coprocessor technology
- Investigate the potential system configurations that satisfy your cost, power, performance requirements.
- Take a clean sheet to develop a custom architecture to achieve your computing performance goals

40 Videos (Current) • 107 VHS • 1 Chapter • 1 Episode • 1

Episode 2.1 — Purpose of the MIC architecture

Parallel Computing in the Search for New Physics at LHC

Fluid Dynamics with Fortran on Intel Xeon Phi coprocessors

At this demonstration, a Fortran program is executed on a cluster of Intel Xeon Phi coprocessors. The program is a Fortran program that simulates the flow of a fluid around a cylinder. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors.

Configuration and Benchmarks of Peer-to-Peer Communication over Gigabit Ethernet and InfiniBand in a Cluster with Intel Xeon Phi Coprocessors

At this demonstration, a Fortran program is executed on a cluster of Intel Xeon Phi coprocessors. The program is a Fortran program that simulates the flow of a fluid around a cylinder. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors.

Interview with James Reinders: future of Intel MIC architecture, parallel programming, education

At this demonstration, a Fortran program is executed on a cluster of Intel Xeon Phi coprocessors. The program is a Fortran program that simulates the flow of a fluid around a cylinder. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors. The program is executed on a cluster of Intel Xeon Phi coprocessors.

<http://colfaxresearch.com/>

LEARN MORE

colfaxresearch.com/how-16-09

© Colfax International, 2013–2016

§2. INTEL ARCHITECTURE

A solid green vertical bar is positioned on the left side of the slide.

COMPUTING PLATFORMS

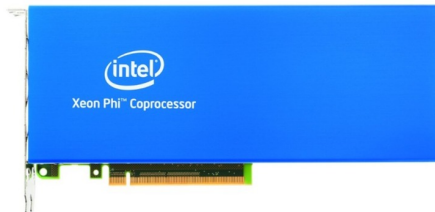
Intel Xeon Processor



Current: Broadwell
Upcoming: Skylake

Multi-Core Architecture

Intel Xeon Phi Coprorocessor, 1st generation



Knights Corner (KNC)

Intel Xeon Phi Processor, 2nd generation*



* socket and coprocessor versions

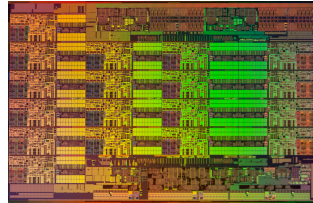
Knights Landing (KNL)

Intel Many Integrated Core (MIC) Architecture

INTEL XEON CPU: PURPOSE AND SPECIFICATIONS

General-purpose platform for demanding computing applications.

- ▶ Up to ~ 1 TFLOP/s in DP*
- ▶ Up to ~ 2 TFLOP/s in SP*
- ▶ Up to 3072 GiB DDR4 RAM*
- ▶ ~ 154 GB/s bandwidth*
- ▶ Hardware-rich: forgiving of sub-optimal code

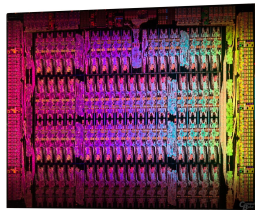
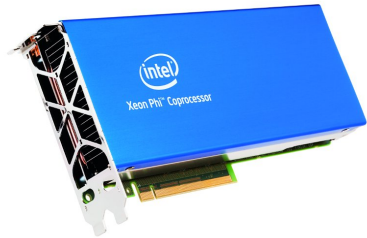


* 2-way Intel Xeon processor, Skylake architecture, top-of-the-line (e.g., E5-2699 V4)

INTEL XEON PHI PROCESSORS (1ST GEN)

Specialized platform for demanding computing applications.

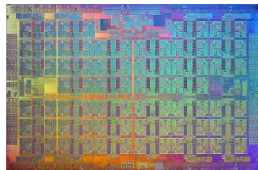
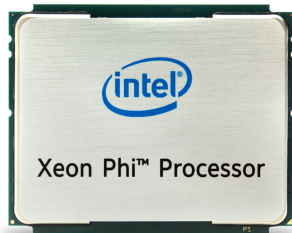
- ▶ PCIe add-in card
- ▶ ~ 1.2 TFLOP/s in DP
- ▶ ~ 2.4 TFLOP/s in SP
- ▶ Up to 16 GiB GDDR5 RAM
- ▶ ~ 176 GB/s bandwidth
- ▶ Heterogeneous clustering
- ▶ Runs special Linux distribution



INTEL XEON PHI PROCESSORS (2ND GEN)

Specialized platform for demanding computing applications.

- ▶ Socket version or coprocessor
- ▶ 64-72 cores $\times$ 4 HT at 1.3-1.5 GHz
- ▶ 3+ TFLOP/s in DP (FMA)
- ▶ 6+ TFLOP/s in SP (FMA)
- ▶ ≤ 384 GiB DDR4 (> 90 GB/s)
- ▶ 16 GiB HBM (MCDRAM, > 400 GB/s)
- ▶ Binary-compatible with Xeon
- ▶ Common OS
(RHEL/CentOS/Fedora/SUSE/Windows)

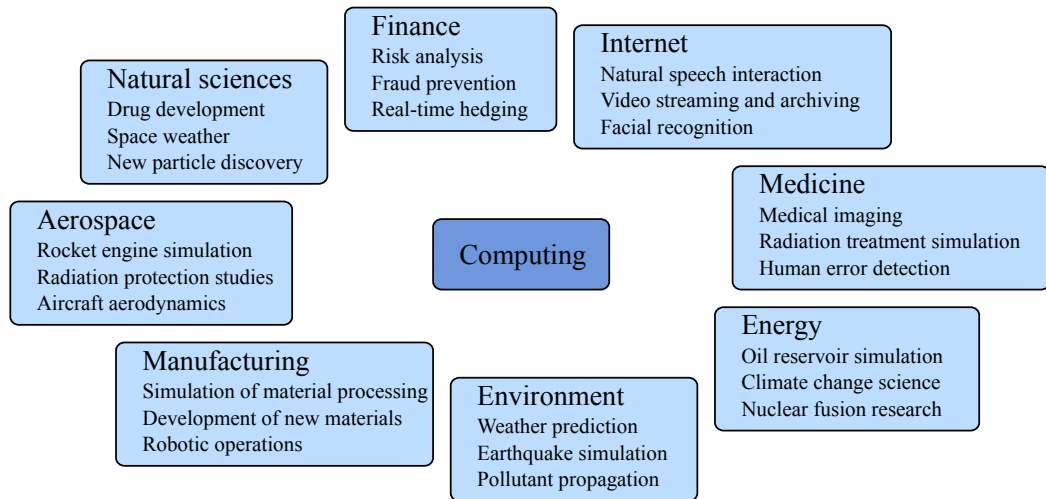


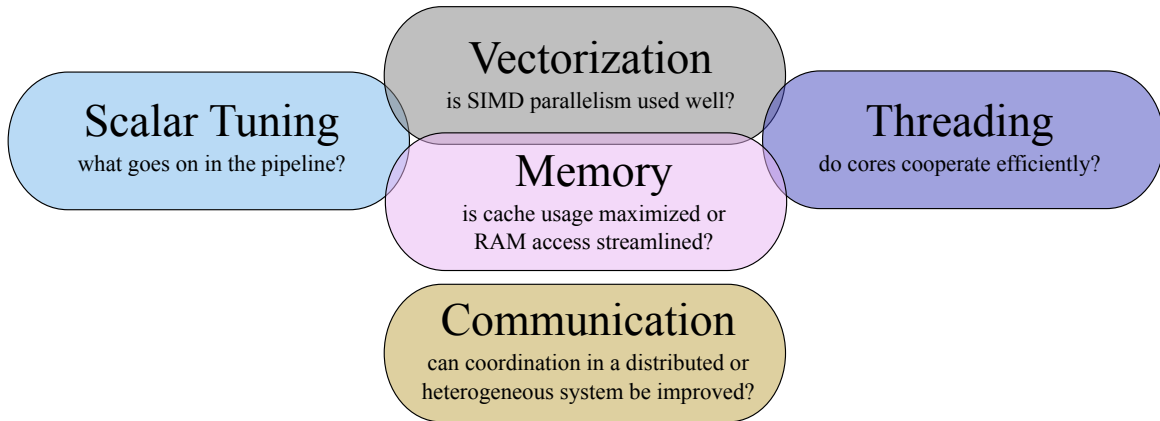


IT IS ALL ABOUT PERFORMANCE

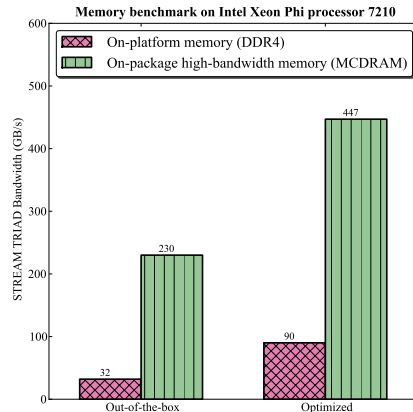
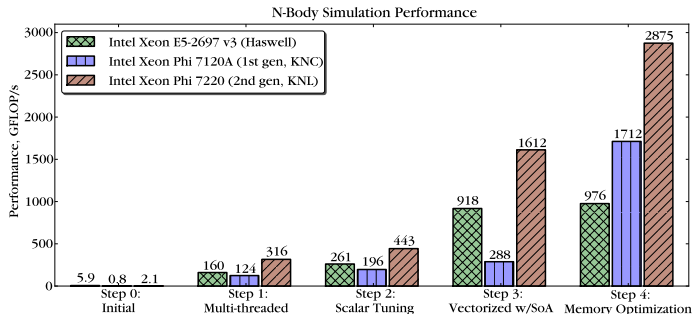
COMPUTING APPLICATIONS

Just some examples

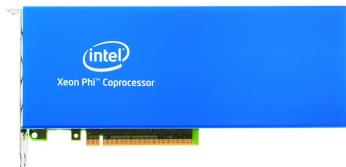




IT TAKES GOOD SOFTWARE TO UNLOCK THE PERFORMANCE!



Details on N-body simulation in Chapter 23 of this book: lotsofcores.com/KNLbook



One Intel Xeon Phi 7120P
coprocessor

vs.



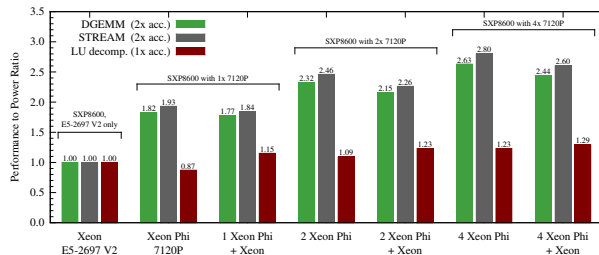
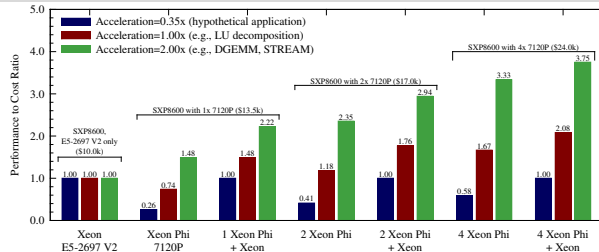
Two Intel Xeon E5-2697 v2
CPUs

- ▶ Why compare 1 coprocessor against 2 processors?
Same thermal design power (TDP).

See also “[Intel Xeon Product Family: Performance Brief](#)”

WHAT IS YOUR PERFORMANCE METRIC?

- ▶ Performance per System
- ▶ Performance per Watt
- ▶ Performance/Cost Ratio



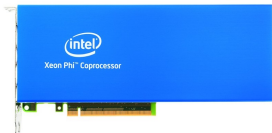
See [this paper](#) for details

A solid green vertical bar is positioned on the left side of the slide.

COMMON ARCHITECTURE



- C/C++/Fortran
- Linux/Windows
- ≤ 3 TiB DDR4
- ≤ 44 cores (2-way)
- ≈ 3 GHz
- 2 HT/core
- 256-bit AVX



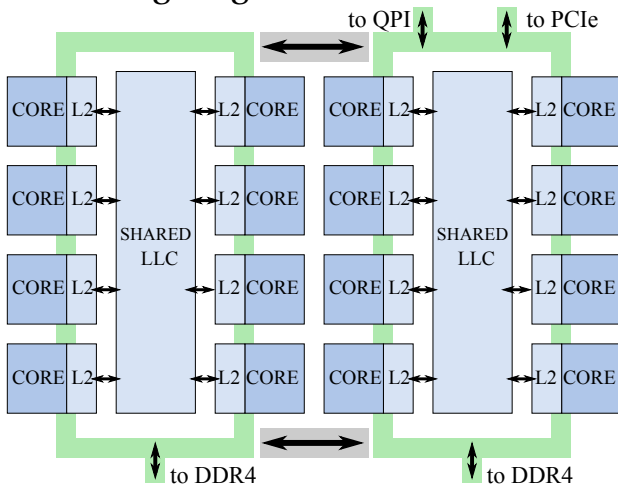
- C/C++/Fortran
- Special Linux
- 3–16 GiB GDDR
- 57–61 cores
- ≈ 1.2 GHz
- 4 HW THR/core
- 512-bit IMCI



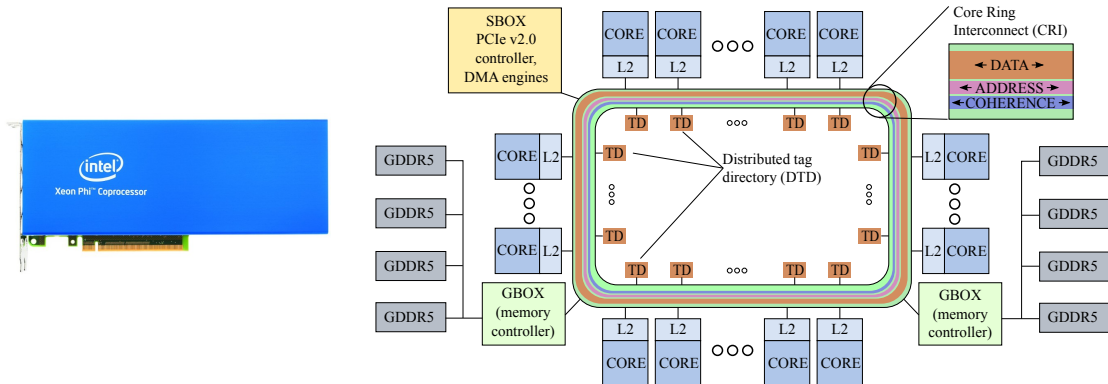
- C/C++/Fortran
- Linux/Windows
- 16 GiB MCDRAM
- 64–72 cores
- 1.3–1.5 GHz
- 4 HT/core
- 512-bit AVX-512

INTEL XEON CPU: DIE ORGANIZATION

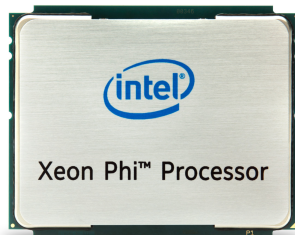
Likes data locality, but large LLC is forgiving.



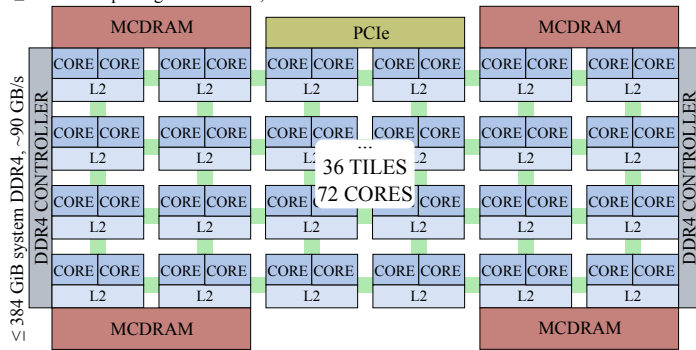
In a ring bus with distributed cache, data access locality is key.



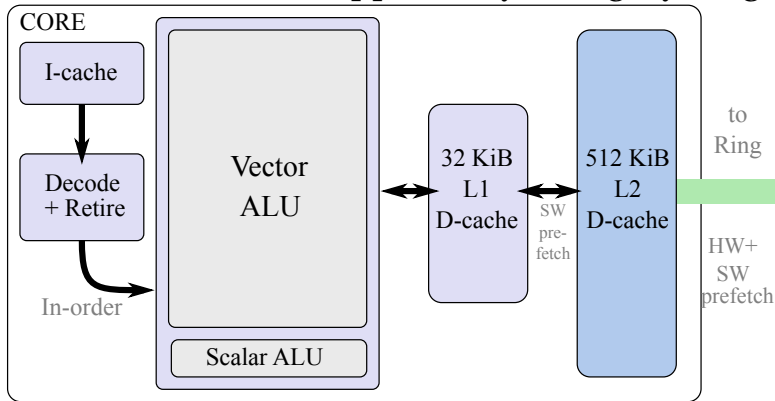
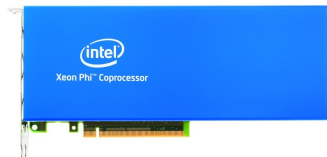
- Mesh interconnect relaxes data locality requirement [somewhat]
- All-to-all, quadrant or sub-numa domain communication in mesh



≤ 16 GiB on-package MCDRAM, ~ 400 GB/s

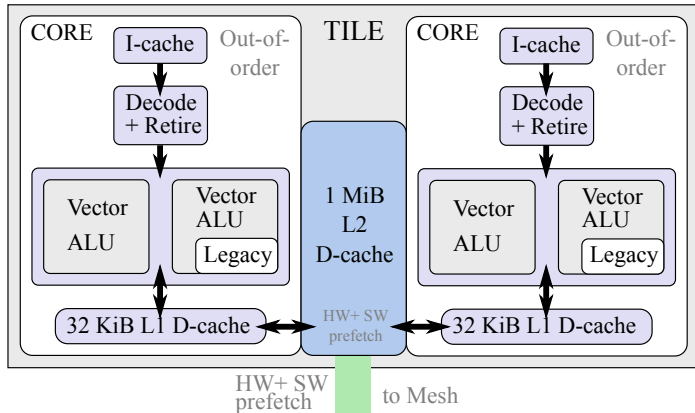
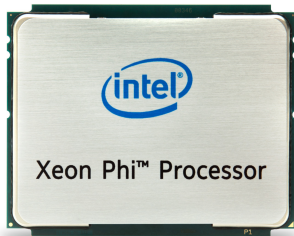


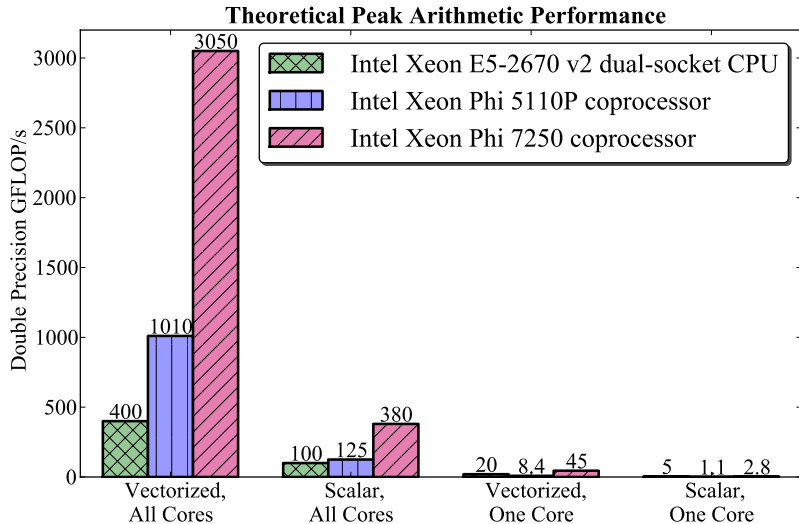
Computing power is in vector units. Scalar support only for legacy usage.



KNL CORES

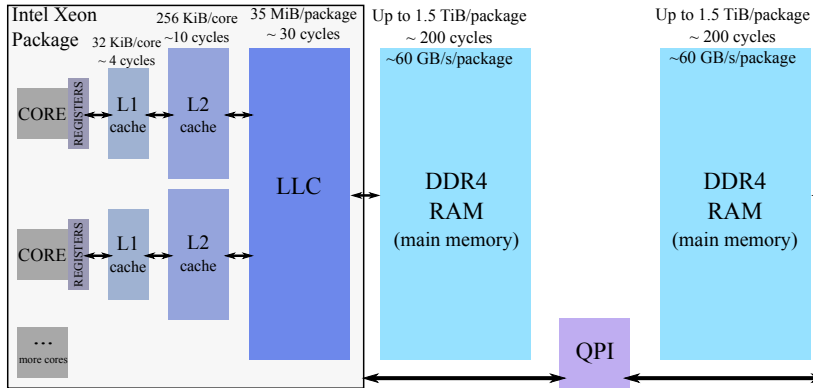
- Even more power in vector units
- Binary compatible with Xeon, but in legacy mode





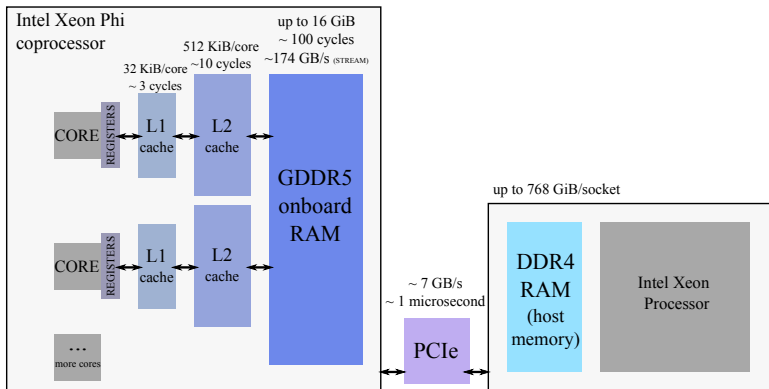
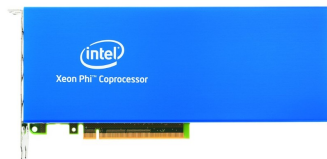
INTEL XEON CPU: MEMORY ORGANIZATION

- ▶ Hierarchical cache structure
- ▶ Two-way processors have NUMA architecture



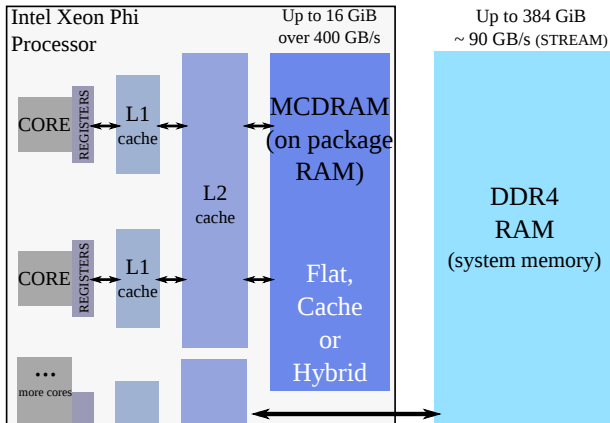
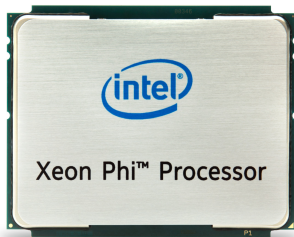
KNC MEMORY ORGANIZATION

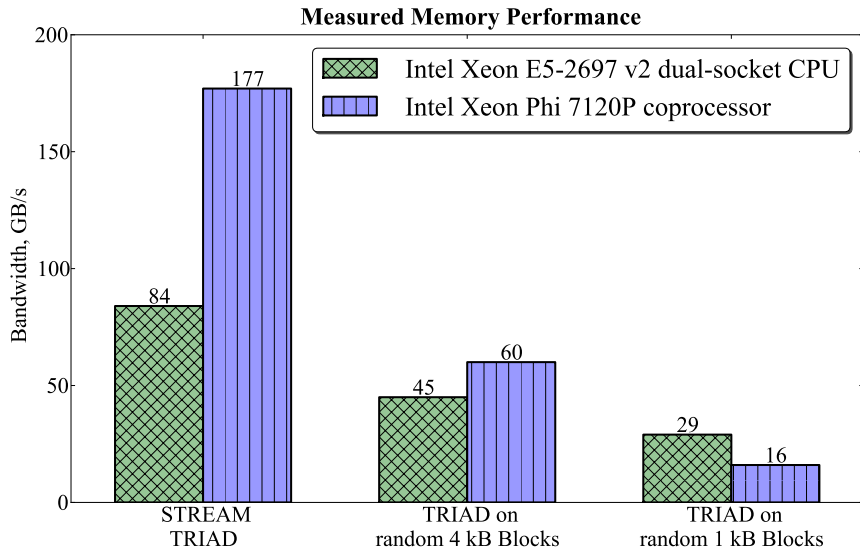
- ▶ Direct access to ≤ 16 GiB of cached GDDR5 memory on board
- ▶ No access to system DDR4, connected to host via PCIe



KNL MEMORY ORGANIZATION

- Direct access to on-package MCDRAM *and* system DDR4 (socket)
- Use MCDRAM as cache, in flat mode, or as hybrid



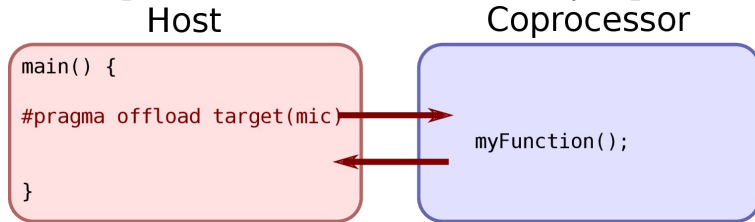


A solid green vertical bar is located on the left side of the slide.

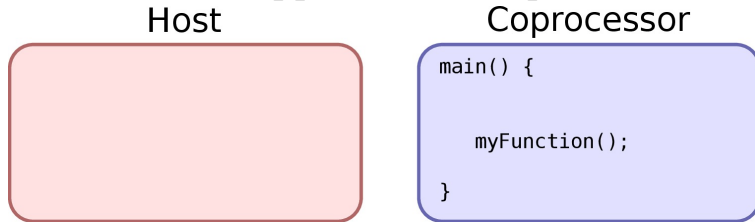
PROGRAMMING COPROCESSORS

OFFLOAD AND NATIVE MODELS

- ▶ Offload model (explicit/virtual-shared memory/OpenMP 4.0):



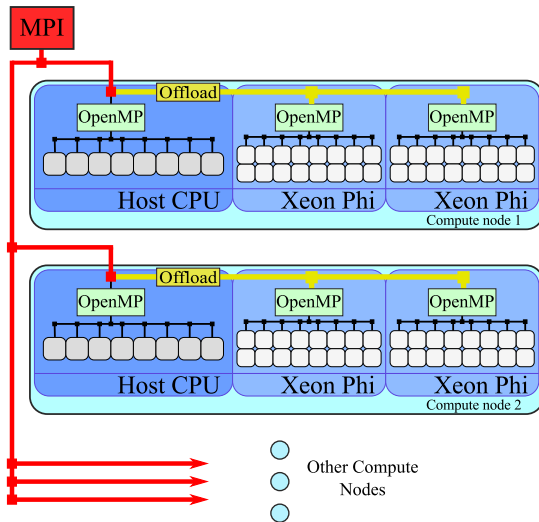
- ▶ Native model (standalone application/MPI process):

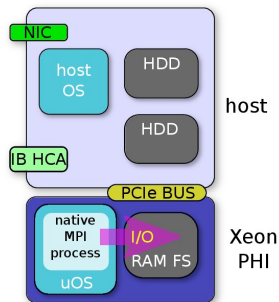


HETEROGENEOUS DISTRIBUTED COMPUTING WITH XEON PHI

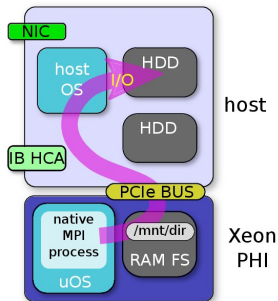
Option 1: MPI+OpenMP with Offload.

- MPI processes are multi-threaded with OpenMP.
- MPI runs only on CPUs.
- MPI processes offload to coprocessor(s).
- OpenMP in offload regions.

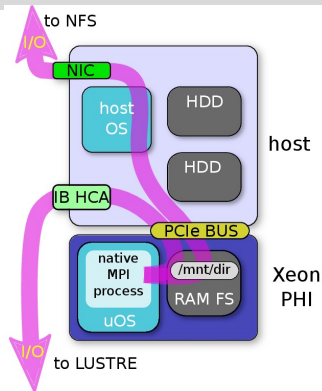




RAM Filesystem



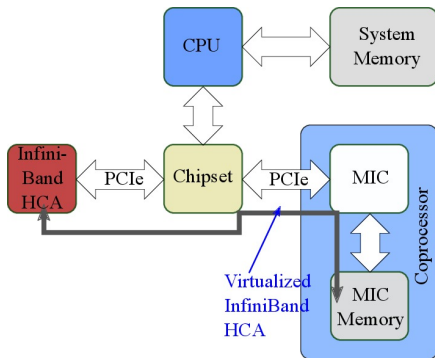
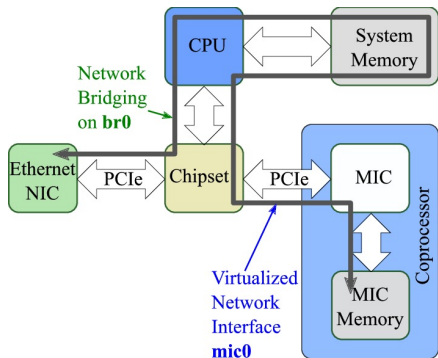
VirtIO Transfer



Network Storage

Details: <http://xeonphi.com/papers/io>

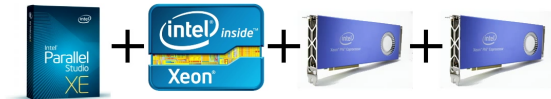
BRIDGED CONFIGURATION FOR PEER-TO-PEER COMMUNICATION



- Left: Gigabit Ethernet bridging and TCP/IP virtualization
- Right: InfiniBand + Coprocessor Communication Link (CCL)
- Details: <http://xeonphi.com/papers/p2p>

A solid green vertical bar is positioned on the left side of the slide.

COLFAX'S CASE STUDIES

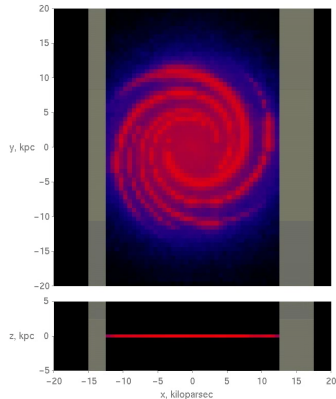


Porting to Intel® Xeon Phi™ coprocessors

- We ported Frankie code using explicit offload model
- Same code & optimization methods for Xeon Phi™
- Simultaneous calculations on CPU and coprocessors with automatic load balancing was easy to implement
- With two Intel® Xeon Phi™ coprocessors, performance for high-res calculations is 3.2x better than with two Intel® Xeon® E5 processors alone.
- **RESULT:** estimated target project calculation time is now 2 weeks (down from 6+ years)

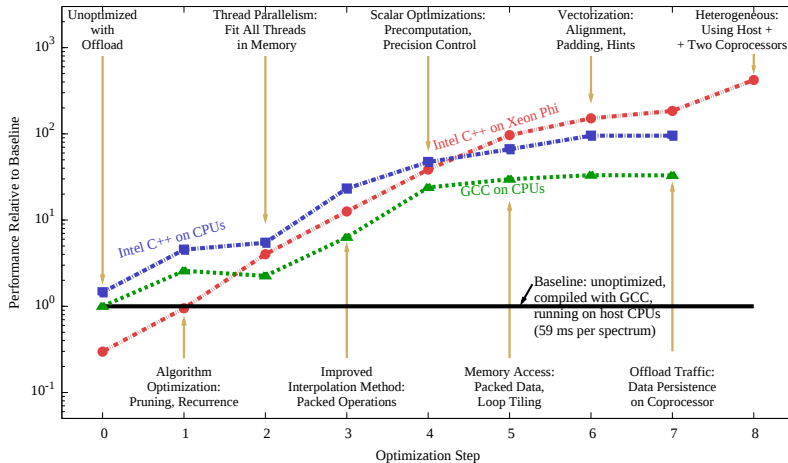
Goal achieved!

Transient Emission of Cosmic Dust Grains
in the Milky Way Galaxy,
Simulation with Frankie Code



<http://xeonphi.com/papers/heatcode>

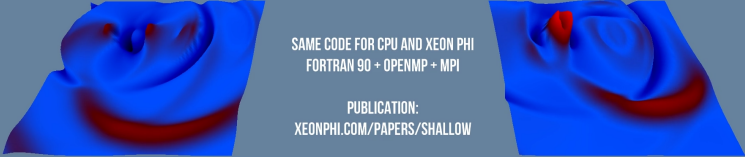
ASTROPHYSICAL CODE HEATCODE: AN OFFLOAD STORY



<http://xeonphi.com/papers/heatcode>

FLUID DYNAMICS WITH FORTRAN ON INTEL® XEON PHI™ COPROCESSORS

SHALLOW WATER EQUATION SOLVER



SAME CODE FOR CPU AND XEON PHI
FORTRAN 90 + OPENMP + MPI

PUBLICATION:
XEONPHI.COM/PAPERS/SHALLOW

PERFORMANCE ON CPU: 19.5 GFLOP/S

PERFORMANCE WITH COPROCESSORS: 52.5 GFLOP/S

SIMULATION SIZE: 9600X9600

ACCELERATION: 2.7X



INTEL XEON E5-2697 V3 PROCESSOR



INTEL XEON E5-2697 V3 PROCESSOR +
TWO INTEL XEON PHI 7120A COPROCESSORS

 COLFAX

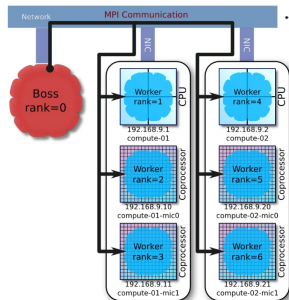
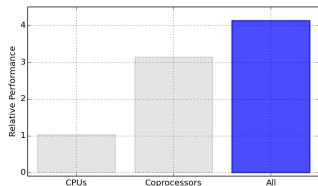
SERVICES WORKSTATIONS TRAINING CONSULTING RESEARCH

WWW.COLFAX-INTL.COM

<http://xeonphi.com/papers/shallow>

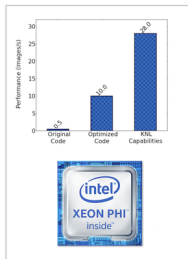
ASIAN OPTION PRICING: HETEROGENEOUS CLUSTERING

Heterogeneous Clustering with Homogeneous Code:
Asian Option Pricing



<http://xeonphi.com/papers/heterogeneous>

INTEL® XEON PHI™ PROCESSORS — MACHINE LEARNING



NEURALTALK2 — OPEN SOURCE IMAGE TAGGING CODE (KARPATY & FEI-FEI, STANFORD)



<http://colfaxresearch.com/isc16-neuraltalk>

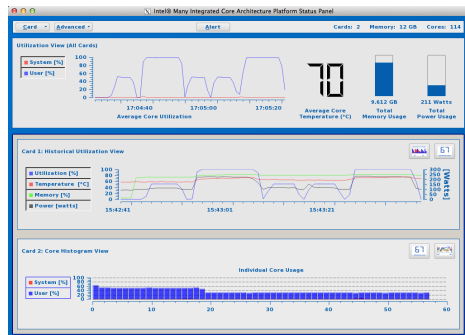
§3. HANDS-ON PART: GUIDED TOUR

LINUX ENVIRONMENT ON INTEL XEON PHI COPROCESSORS

```
vega@lyra% lspci | grep -i "co-processor"
06:00.0 Co-processor: Intel Corporation Xeon Phi coprocessor 7120 series (rev 20)
82:00.0 Co-processor: Intel Corporation Xeon Phi coprocessor 7120 series (rev 20)
vega@lyra% sudo service mpss status
mpss is running
vega@lyra% cat /etc/hosts | grep mic
172.31.1.1  lyra-mic0 mic0
172.31.2.1  lyra-mic1 mic1
vega@lyra% ssh mic0
vega@mic0% cat /proc/cpuinfo | grep proc | tail -n 3
processor: 241
processor: 242
processor: 243
vega@mic0% ls /
amplxe  dev    home   lib64   oldroot  proc    sbin    sys     usr
bin     etc    lib    linuxrc opt      root    sep3.10 tmp     var
```

INTEL MANYCORE PLATFORM SOFTWARE STACK

- micinfo – system information
- micsmc – monitor and modify the physical parameters: temperature, power modes, core utilization, etc.
- micctrl – configure the Intel Xeon Phi coprocessor operating system
- miccheck – verify the Intel Xeon Phi coprocessor configuration
- micrasd – log of hardware errors reported by Intel Xeon Phi coprocessors
- micflash – flash memory agent



Monitoring MIC activity with
micsmc (an **MPSS** tool)

SOFTWARE NECESSARY TO BUILD XEON PHI APPLICATIONS:

Compilers : Intel C Compiler, Intel C++ Compiler, and Intel Fortran Compiler — mandatory

Optimization tools : Intel VTune Amplifier XE and Intel Trace Analyzer and Collector (ITAC) — highly recommended

Mathematics support : Intel Math Kernel Library (MKL) — highly recommended

Cluster Development : Intel MPI — industry standard parallel framework

Development : Intel Inspector XE, Intel Advisor XE — optional



All-in-One bundle,
Intel Parallel Studio XE

“Hello World” application:

```
1 #include <stdio>
2 #include <unistd.h>
3 int main(){
4     printf("Hello world! I have %ld logical processors.\n",
5         sysconf(_SC_NPROCESSORS_ONLN ));
6 }
```

Compile and run on host CPU:

```
vega@lyra% icpc hello.cc -xhost
vega@lyra% ./a.out
Hello world! I have 48 logical processors.
vega@lyra%
```

NATIVE EXECUTION ON AN INTEL XEON PHI COPROCESSOR (KNC)

Compile and run the same code on the coprocessor in the native mode:

```
vega@lyra% icpc hello.cc -mmic # Cross-compile
vega@lyra% scp a.out mic0:~/ # Put executable on coprocessor
a.out 100% 10KB 10.4KB/s 00:00
vega@lyra% ssh mic0 # Log in to coprocessor
vega@mic0% pwd
/home/lyra
vega@mic0% ls
a.out
vega@mic0% ./a.out # Launch application
Hello world! I have 244 logical processors.
vega@mic0%
```

- ▶ Use `-mmic` to produce executable for MIC architecture
- ▶ Must transfer executable to coprocessor (or NFS-share) and run from shell
- ▶ Native MPI applications work the same way (need Intel MPI library)

COMPILING FOR AN INTEL XEON PHI PROCESSOR (KNL)

“Hello World” application:

```
1 #include <stdio>
2 #include <unistd.h>
3 int main(){
4     printf("Hello world! I have %ld logical processors.\n",
5         sysconf(_SC_NPROCESSORS_ONLN ));
6 }
```

Compile and run on host CPU:

```
vega@lyra% icpc hello.cc -xMIC-AVX512
vega@lyra% ./a.out
Hello world! I have 256 logical processors.
vega@lyra%
```

From Intel compilers 17.0 on, `-xMIC-AVX512` is not necessary when compiling on KNL.

- ▶ Intel Xeon and Intel Xeon Phi – parallel processors
- ▶ Xeon Phi (MIC architecture) – specialized for highly parallel workloads without complex memory access
- ▶ Coprocessor – either offload device or an additional compute node
- ▶ Native+offload programming allow for a range of design options

Next session: details of Intel Xeon Phi processor and coprocessor programming.