



Matrix factorization routines on heterogeneous architectures

Nikita Shustrov, Nadya Mozartova, Tamara Kashevarova, Konstantin Arturov. Intel Corporation

Introduction

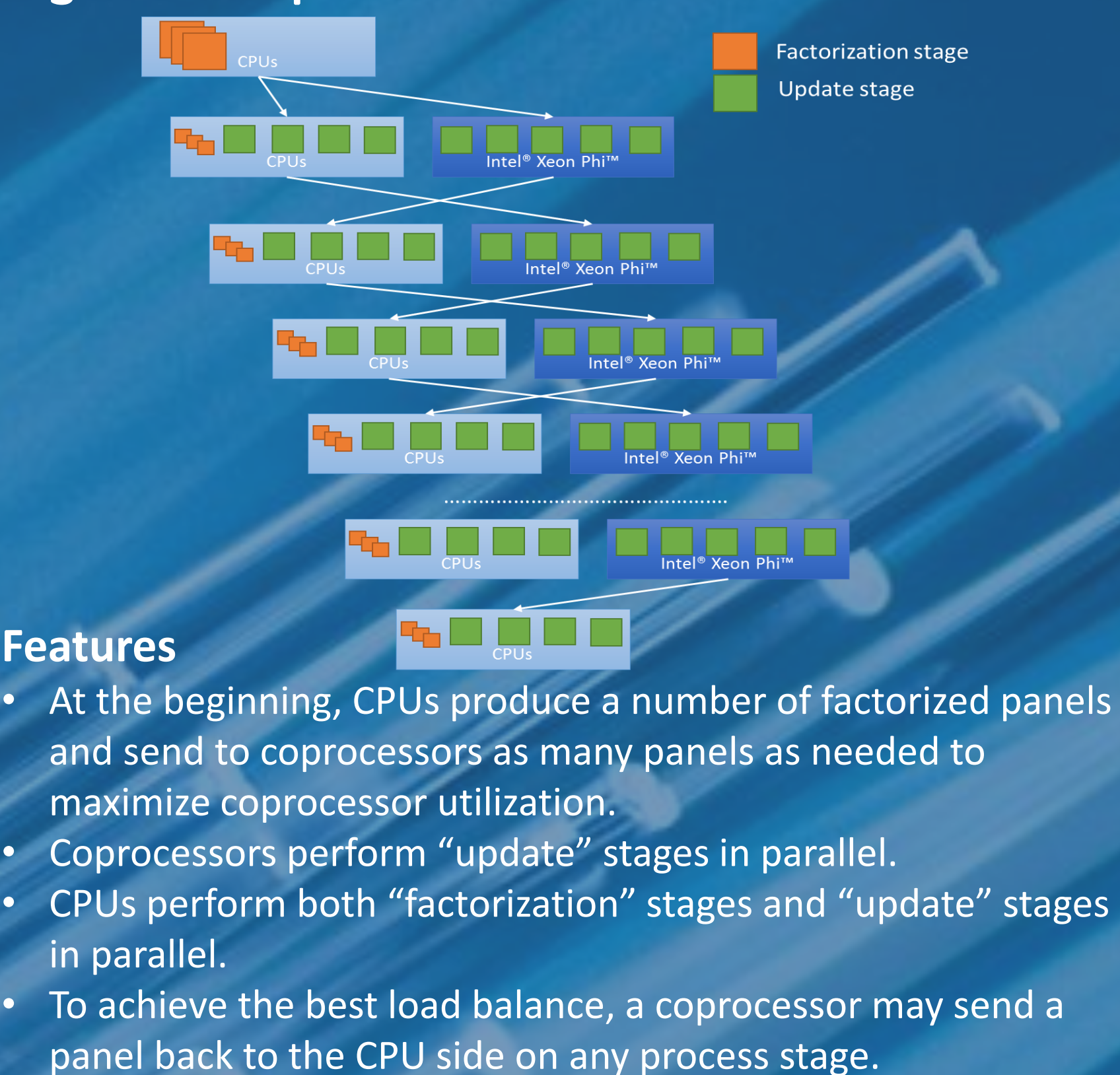
- New heterogeneous systems consisting of a multicore CPU with coprocessors introduce new challenges to designing efficient parallel algorithms.
- Simultaneous usage of all computational resources of such systems for solving one large problem requires uneven distribution of data and computations that leads to more complex parallelization methods.
- We present a new parallelization method for matrix factorization that efficiently utilizes all computational resources of a heterogeneous system consisting of a multicore CPU with coprocessors.

We show how this method can be applied for parallelization of the key linear algebra factorization algorithms (QR, LU and Cholesky) on systems with Intel® Xeon Phi™ coprocessors.

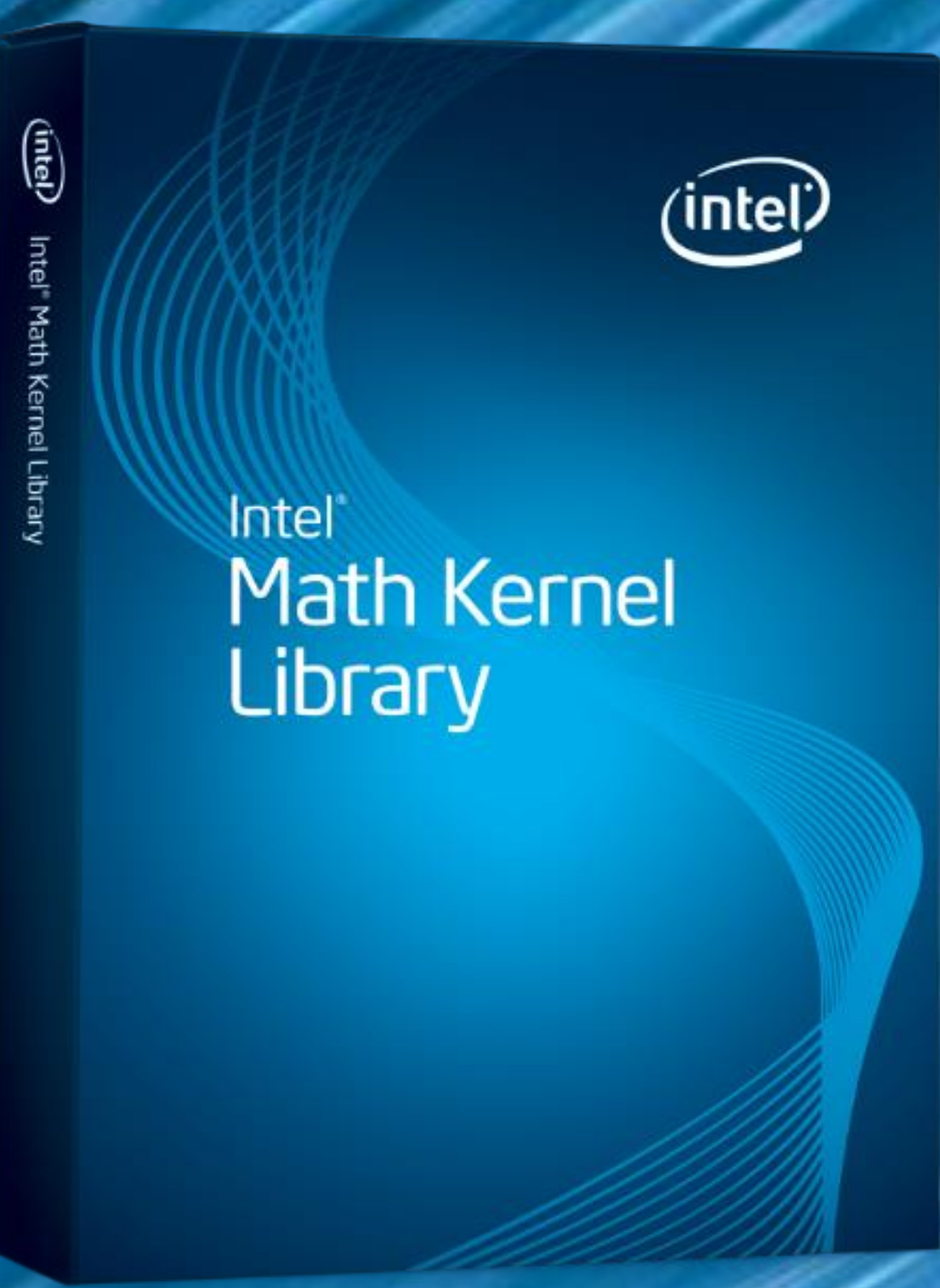
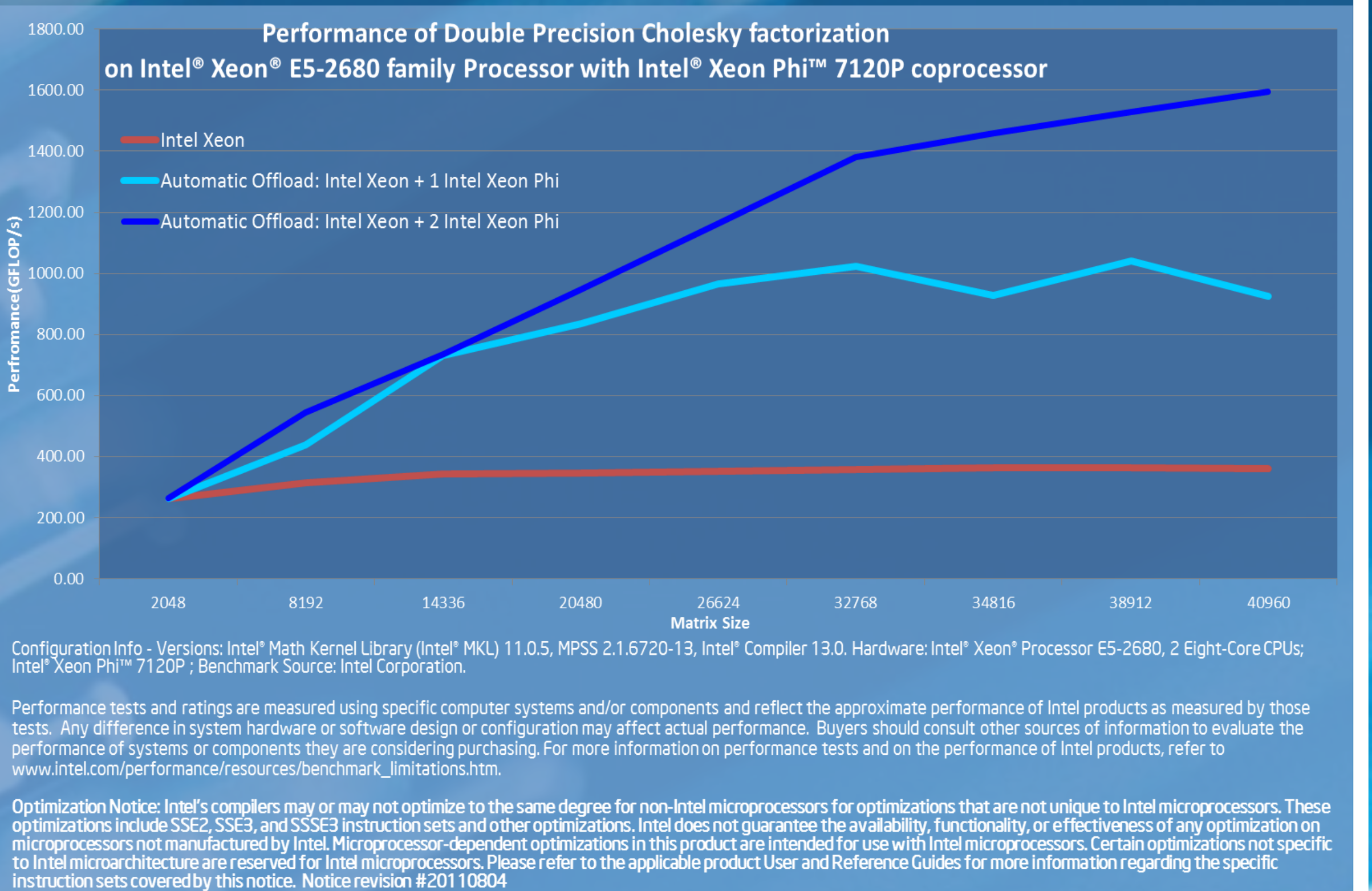
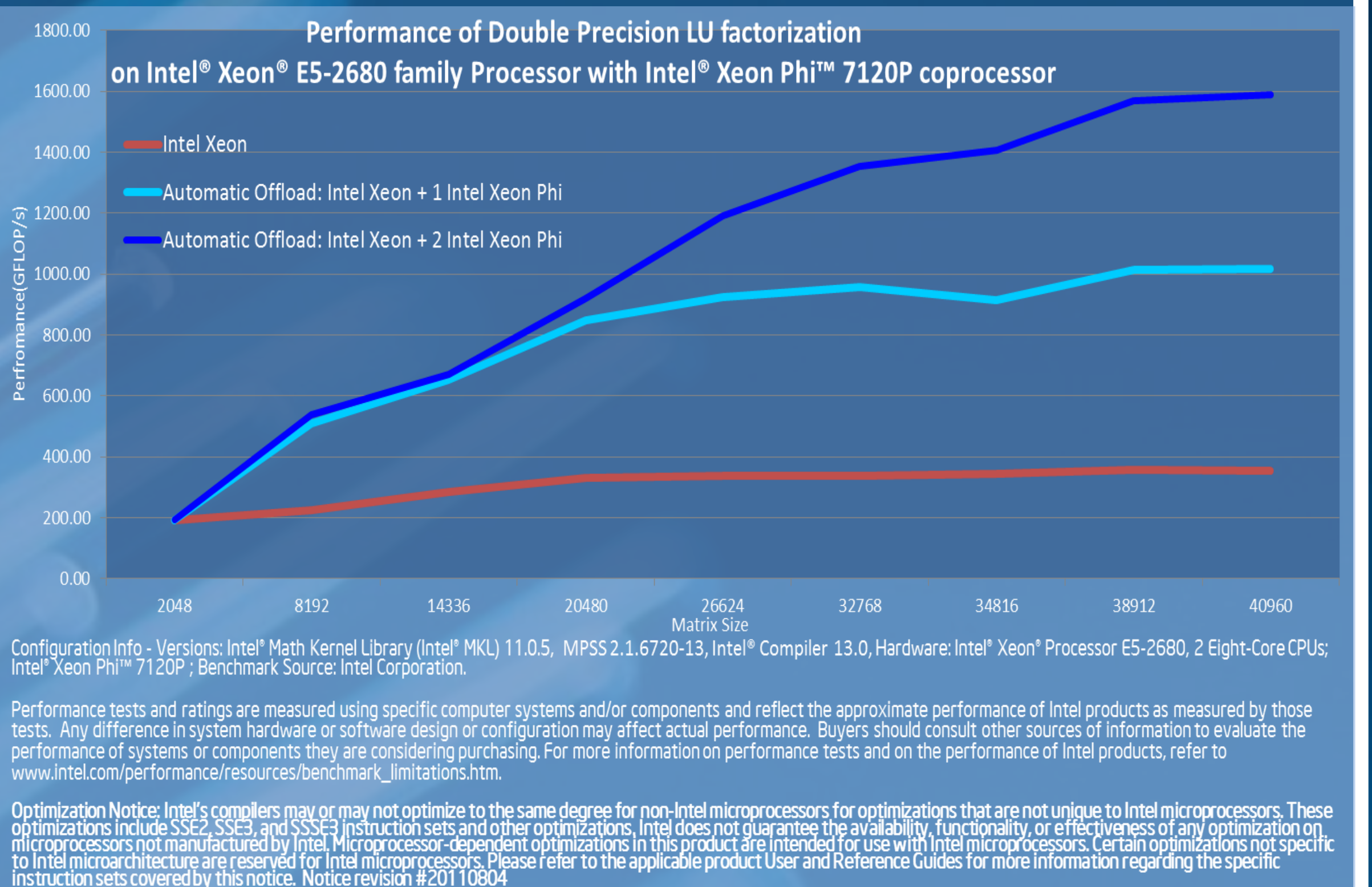
Algorithm - Overview

- Algorithm is based on the panel factorization method. It has advantages over the communication-avoiding method, the tile method, and their combination because it does not require:
 - additional computational cost
 - additional memory consumption
- Panel factorization has the same computational cost and memory usage as classic LAPACK algorithms.
- Implementation preserves LAPACK standard interfaces and data layout.
- Algorithm can be applied to any matrices.
- Implementation is DAG-based and uses panel factorization kernels that were redesigned and rewritten for new Intel Xeon Phi architectures.

Algorithm represented as DAG



Performance results



Algorithm - Features

- Adaptive data/task distribution on the fly between CPUs and coprocessors to improve load balancing.
- Efficient utilization of all available computational units in heterogeneous systems.
- No limit on number of coprocessors on heterogenous systems.
- Scalability: A system with CPUs and 1 coprocessor shows >3x speed up. A system with CPUs and 2 coprocessors shows >5x speed up.

The algorithm provides a high degree of parallelism with minimal requirements for synchronization and communication

Algorithm - Implementation

- The algorithm is implemented within the framework of Intel® Math Kernel Library (Intel® MKL) LU, QR, and Cholesky factorization routines.
- The implementation detects the presence of Intel Xeon Phi coprocessors and automatically offloads the computations that can benefit from additional computational resources.
- The usage model hides the complexity of heterogeneous systems from the user, providing ease of use and providing the same API as usual Intel MKL routines.

Users do not need to change application code.

<http://www.intel.com>
<http://software.intel.com/en-us/intel-mkl>

Intel, Xeon, and Intel Xeon Phi are trademarks of Intel Corporation in the U.S. and/or other countries.

