



Intel® oneAPI Math Kernel Library for Windows*

Developer Guide

Intel® oneAPI Math Kernel Library- Windows*

Notices and Disclaimers

2022.2

Contents

Notices and Disclaimers	6
Getting Help and Support	7
Introducing the Intel® oneAPI Math Kernel Library	8
Notational Conventions	10
Related Information	11
Chapter 1: Getting Started	
Shared Library Versioning	12
CMake Config for oneMKL	12
Checking Your Installation	13
Setting Environment Variables	14
Compiler Support	14
Using Code Examples	15
Before You Begin Using the Intel® oneAPI Math Kernel Library	15
Chapter 2: Structure of the Intel® oneAPI Math Kernel Library	
Architecture Support	18
High-level Directory Structure	18
Layered Model Concept	19
Chapter 3: Linking Your Application with the Intel® oneAPI Math Kernel Library	
Linking Quick Start	21
Using the /Qmkl Compiler Option	21
Automatically Linking a Project in the Visual Studio* Integrated Development Environment with Intel® oneAPI Math Kernel Library	22
Automatically Linking Your Microsoft Visual C/C++* Project with oneMKL	22
Automatically Linking Your Intel® Visual Fortran Project with oneMKL	22
Using the Single Dynamic Library	23
Selecting Libraries to Link with	23
Using the Link-line Advisor	24
Using the Command-line Link Tool	24
Linking Examples	26
Linking on IA-32 Architecture Systems	26
Linking on Intel(R) 64 Architecture Systems	26
Linking in Detail	27
Dynamically Selecting the Interface and Threading Layer	28
Linking with Interface Libraries	29
Using the ILP64 Interface vs. LP64 Interface	29
Linking with Fortran 95 Interface Libraries	31
Linking with Threading Libraries	31
Linking with Computational Libraries	33
Linking with Compiler Run-time Libraries	34
Linking with System Libraries	34
Building Custom Dynamic-link Libraries	34
Using the Custom Dynamic-link Library Builder in the Command-line Mode	35

Composing a List of Functions	37
Specifying Function Names	38
Building a Custom Dynamic-link Library in the Visual Studio* Development System	38
Distributing Your Custom Dynamic-link Library	39
Building a Universal Windows Driver	39
Chapter 4: Managing Performance and Memory	
Improving Performance with Threading	43
OpenMP* Threaded Functions and Problems	43
Functions Threaded with Intel® Threading Building Blocks	45
Avoiding Conflicts in the Execution Environment	46
Techniques to Set the Number of Threads	47
Setting the Number of Threads Using an OpenMP* Environment Variable	47
Changing the Number of OpenMP* Threads at Run Time	48
Using Additional Threading Control	50
oneMKL-specific Environment Variables for OpenMP Threading Control	50
MKL_DYNAMIC	51
MKL_DOMAIN_NUM_THREADS	52
MKL_NUM_STRIPES	53
Setting the Environment Variables for Threading Control	54
Calling oneMKL Functions from Multi-threaded Applications	55
Using Intel® Hyper-Threading Technology	56
Managing Multi-core Performance	57
Managing Performance with Heterogeneous Cores	58
Improving Performance for Small Size Problems	58
Using MKL_DIRECT_CALL in C Applications	59
Using MKL_DIRECT_CALL in Fortran Applications	59
Limitations of the Direct Call	60
Other Tips and Techniques to Improve Performance	60
Coding Techniques	60
Improving oneMKL Performance on Specific Processors	61
Operating on Denormals	61
Using Memory Functions	62
Memory Leaks in Intel® oneAPI Math Kernel Library	62
Redefining Memory Functions	62
Chapter 5: Language-specific Usage Options	
Using Language-Specific Interfaces with Intel® oneAPI Math Kernel Library	64
Interface Libraries and Modules	64
Fortran 95 Interfaces to LAPACK and BLAS	65
Compiler-dependent Functions and Fortran 90 Modules	66
Mixed-language Programming with the Intel Math Kernel Library	66
Calling LAPACK, BLAS, and CBLAS Routines from C/C++ Language Environments	67
Using Complex Types in C/C++	68
Calling BLAS Functions that Return the Complex Values in C/C++ Code ..	69
Chapter 6: Obtaining Numerically Reproducible Results	
Getting Started with Conditional Numerical Reproducibility	73
Specifying Code Branches	74
Reproducibility Conditions	76
Setting the Environment Variable for Conditional Numerical Reproducibility	77

Code Examples	77
Chapter 7: Coding Tips	
Example of Data Alignment.....	80
Using Predefined Preprocessor Symbols for Intel® MKL Version-Dependent Compilation	81
Chapter 8: Managing Output	
Using oneMKL Verbose Mode.....	83
Version Information Line	84
Call Description Line	86
Chapter 9: Working with the Intel® oneAPI Math Kernel Library Cluster Software	
Message-Passing Interface Support.....	90
Linking with oneMKL Cluster Software	91
Determining the Number of OpenMP* Threads	92
Using DLLs	92
Setting Environment Variables on a Cluster.....	93
Interaction with the Message-passing Interface	93
Using a Custom Message-Passing Interface.....	94
Examples of Linking for Clusters.....	95
Examples for Linking a C Application.....	95
Examples for Linking a Fortran Application.....	96
Chapter 10: Managing Behavior of the Intel® oneAPI Math Kernel Library with Environment Variables	
Managing Behavior of Function Domains with Environment Variables	98
Setting the Default Mode of Vector Math with an Environment Variable...	98
Managing Performance of the Cluster Fourier Transform Functions	99
Managing Invalid Input Checking in LAPACKE Functions.....	100
Instruction Set Specific Dispatching on Intel® Architectures	101
Chapter 11: Programming with Intel® oneAPI Math Kernel Library in the Visual Studio* Integrated Development Environment	
Configuring Your Integrated Development Environment to link with Intel® oneAPI Math Kernel Library	103
Configuring the Microsoft Visual C/C++* Development System to Link with Intel® oneAPI Math Kernel Library	103
Configuring Intel(R) Visual Fortran to Link with Intel® oneAPI Math Kernel Library	104
Getting Assistance for Programming in the Microsoft Visual Studio* IDE	104
Using Context-Sensitive Help	104
Using the IntelliSense* Capability.....	104
Chapter 12: Intel® oneAPI Math Kernel Library Benchmarks	
Intel Optimized LINPACK Benchmark for Windows*	107
Contents	107
Running the Software	108
Known Limitations.....	109
Intel® Distribution for LINPACK* Benchmark.....	109
Overview	109
Contents	110

Building the Intel® Distribution for LINPACK* Benchmark for a Customized MPI Implementation	110
Building the Netlib HPL from Source Code.....	111
Configuring Parameters.....	111
Ease-of-use Command-line Parameters	112
Running the Intel® Distribution for LINPACK* Benchmark.....	112
Heterogeneous Support in the Intel® Distribution for LINPACK* Benchmark.....	113
Environment Variables	115
Improving Performance of Your Cluster	116
 Appendix A: Appendix A: Intel® oneAPI Math Kernel Library Language Interfaces Support	
Language Interfaces Support, by Function Domain.....	117
Include Files	118
 Appendix B: Support for Third-Party Interfaces	
FFTW Interface Support	120
 Appendix C: Appendix C: Directory Structure In Detail	
Detailed Structure of the IA-32 Architecture Directories.....	121
Static Libraries in the lib\ia32_win Directory	121
Dynamic Libraries in the lib\ia32_win Directory	122
Contents of the redist\ia32 Directory	122
Detailed Structure of the Intel(R) 64 Architecture Directories	125
Static Libraries in the lib\intel64_win Directory.....	125
Dynamic Libraries in the lib\intel64_win Directory.....	126
Contents of the redist\intel64 Directory	128

Notices and Disclaimers

Intel technologies may require enabled hardware, software or service activation.

No product or component can be absolutely secure.

Your costs and results may vary.

© Intel Corporation. Intel, the Intel logo, and other Intel marks are trademarks of Intel Corporation or its subsidiaries. Other names and brands may be claimed as the property of others.

Product and Performance Information
Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex . Notice revision #20201201

No license (express or implied, by estoppel or otherwise) to any intellectual property rights is granted by this document.

The products described may contain design defects or errors known as errata which may cause the product to deviate from published specifications. Current characterized errata are available on request.

Intel disclaims all express and implied warranties, including without limitation, the implied warranties of merchantability, fitness for a particular purpose, and non-infringement, as well as any warranty arising from course of performance, course of dealing, or usage in trade.

Microsoft, Windows, and the Windows logo are trademarks, or registered trademarks of Microsoft Corporation in the United States and/or other countries.

Java is a registered trademark of Oracle and/or its affiliates.

Getting Help and Support

Intel provides a support web site that contains a rich repository of self help information, including getting started tips, known product issues, product errata, license information, user forums, and more. Visit the Intel® oneAPI Math Kernel Library support website at <http://www.intel.com/software/products/support/>.

You can get context-sensitive help when editing your code in the Microsoft Visual Studio* integrated development environment (IDE). See [Getting Assistance for Programming in the Microsoft Visual Studio* IDE](#) for details.

Introducing the Intel® oneAPI Math Kernel Library

Intel® oneAPI Math Kernel Library (oneMKL) is a computing math library of highly optimized, extensively threaded routines for applications that require maximum performance. The library provides Fortran and C programming language interfaces. oneMKL C language interfaces can be called from applications written in either C or C++, as well as in any other language that can reference a C interface.

oneMKL provides comprehensive functionality support in these major areas of computation:

- BLAS (level 1, 2, and 3) and LAPACK linear algebra routines, offering vector, vector-matrix, and matrix-matrix operations.
- ScaLAPACK distributed processing linear algebra routines, as well as the Basic Linear Algebra Communications Subprograms (BLACS) and the Parallel Basic Linear Algebra Subprograms (PBLAS).
- oneMKL PARDISO (a direct sparse solver based on Parallel Direct Sparse Solver PARDISO*), an iterative sparse solver, and supporting sparse BLAS (level 1, 2, and 3) routines for solving sparse systems of equations, as well as a distributed version of oneMKL PARDISO solver provided for use on clusters.
- Fast Fourier transform (FFT) functions in one, two, or three dimensions with support for mixed radices (not limited to sizes that are powers of 2), as well as distributed versions of these functions provided for use on clusters.
- Vector Mathematics (VM) routines for optimized mathematical operations on vectors.
- Vector Statistics (VS) routines, which offer high-performance vectorized random number generators (RNG) for several probability distributions, convolution and correlation routines, and summary statistics functions.
- Data Fitting Library, which provides capabilities for spline-based approximation of functions, derivatives and integrals of functions, and search.
- Extended Eigensolver, a shared memory programming (SMP) version of an eigensolver based on the Feast Eigenvalue Solver.

For details see the *Intel® oneAPI Math Kernel Library Developer Reference*.

Intel® oneAPI Math Kernel Library (oneMKL) is optimized for performance on Intel processors. oneMKL also runs on non-Intel x86-compatible processors.

For Windows* and Linux* systems based on Intel® 64 Architecture, oneMKL also includes support for the Intel® Many Integrated Core Architecture (Intel® MIC Architecture) and provides libraries to help you port your applications to Intel MIC Architecture.

NOTE

oneMKL provides limited input validation to minimize the performance overheads. It is your responsibility when using oneMKL to ensure that input data has the required format and does not contain invalid characters. These can cause unexpected behavior of the library. Examples of the inputs that may result in unexpected behavior:

- Not-a-number (NaN) and other special floating point values
- Large inputs may lead to accumulator overflow

As the oneMKL API accepts raw pointers, it is your application's responsibility to validate the buffer sizes before passing them to the library. The library requires subroutine and function parameters to be valid before being passed. While some oneMKL routines do limited checking of parameter errors, your application should check for NULL pointers, for example.

Product and Performance Information
Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex . Notice revision #20201201

Notational Conventions

The following term is used in reference to the operating system.

Windows* This term refers to information that is valid on all supported Windows* operating systems.

The following notations are used to refer to Intel® oneAPI Math Kernel Library directories.

<parent directory> The installation directory that includes Intel® oneAPI Math Kernel Library directory; for example, the directory for Intel® Parallel Studio XE Composer Edition.

<mkl directory> The main directory where Intel® oneAPI Math Kernel Library is installed:
<mkl directory>=<parent directory>\mkl.
Replace this placeholder with the specific pathname in the configuring, linking, and building instructions.

The following font conventions are used in this document.

Italic Italic is used for emphasis and also indicates document names in body text, for example:
see *Intel® oneAPI Math Kernel Library Developer Reference*.

Monospace lowercase mixed with uppercase Indicates:

- Commands and command-line options, for example,

```
ifort myprog.f mkl_blas95.lib mkl_c.lib libiomp5md.lib
```
- Filenames, directory names, and pathnames, for example,

```
C:\Program Files\Java\jdk1.5.0_09
```
- C/C++ code fragments, for example,

```
a = new double [SIZE*SIZE];
```

UPPERCASE MONOSPACE Indicates system variables, for example, \$MKLPATH.

Monospace italic Indicates a parameter in discussions, for example, *lda*.
When enclosed in angle brackets, indicates a placeholder for an identifier, an expression, a string, a symbol, or a value, for example, *<mkl directory>*.
Substitute one of these items for the placeholder.

[items] Square brackets indicate that the items enclosed in brackets are optional.

{ item | item } Braces indicate that only one of the items listed between braces should be selected. A vertical bar (|) separates the items.

Related Information

To reference how to use the library in your application, use this guide in conjunction with the following documents:

- The *Intel® oneAPI Math Kernel Library Developer Reference*, which provides *reference* information on routine functionalities, parameter descriptions, interfaces, calling syntaxes, and return values.
- The *Intel® oneAPI Math Kernel Library forWindows* OS Release Notes*.

Getting Started

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Shared Library Versioning

Intel® oneAPI Math Kernel Library (oneMKL) adds shared library versioning for all operating systems and platforms, as opposed to not using any library versioning up to Intel® Math Kernel Library (Intel® MKL) 2020 Update 4.

This new feature:

- Allows applications to work correctly in an environment with multiple oneMKL and/or Intel® MKL packages installed
- Communicates clearly when incompatible changes are made, and an application should be rebuilt
- Allows you to link against a specific version of shared libraries

The starting version for shared libraries is "1" and any change that break backward compatibility will result in an increment to this number. Intel expects to make this change as seldom as possible and inform customers about it at least 24 months in advance.

The product version "2021.1" is now decoupled from the library version, meaning that "2021.2" can ship with shared libraries versioned as "1". This means that the libraries shipped in "2021.2" are backward compatible with libraries shipped in "2021.1".

Changes to the link-line:

- No changes are required to the link-line because symbolic links are provided with the old names, which point to the new library that contains the version information on Linux* and MacOS*. The symbolic link name is also the `soname` and `install_name` of that library on Linux and MacOS, respectively.
 - For example, `libmkl_core.so` -> `libmkl_core.<version>.so`
 - For example, `libmkl_core.dylib` -> `libmkl_core.<version>.dylib`
 - Using `-lmkl_core` will still work as before, ensuring backward compatibility with Intel® MKL 2020 line-up (including Intel® Parallel Studio and Intel® System Studio distributions).
- On Windows*, import libraries used in the link-line do not contain any version information, as before, but point to the new DLL, which contains the version information.
 - For example, `mkl_core_dll.lib` has the same name as before and requires no change to the link-line. The linker, however, resolves this to the new `mkl_core.<version>.dll` instead of the older `mkl_core.dll`.

CMake Config for oneMKL

If you want to integrate oneMKL into your CMake projects, starting with the Intel® oneAPI Math Kernel Library (oneMKL) 2021.3 release, `MKLConfig.cmake` is provided as part of the package and installation. `MKLConfig.cmake` supports all oneMKL configurations, compilers, and runtimes, as the oneMKL product itself. Help/usage is provided in the top section of `MKLConfig.cmake`.

Example

```
my_test/
|_ build/          <-- Out-of-source build directory
|_ CMakeLists.txt  <-- User side project's CMakeLists.txt
|_ app.c           <-- Source file that uses oneMKL API
```

CMakeLists.txt

```
cmake_minimum_required(VERSION 3.13)
enable_testing()
project(oneMKL_Example LANGUAGES C)
find_package(MKL CONFIG REQUIRED)
#message(STATUS "${MKL_IMPORTED_TARGETS}") #Provides available list of targets based on input
add_executable(myapp app.c)

target_compile_options(myapp PUBLIC $<TARGET_PROPERTY:MKL::MKL,INTERFACE_COMPILE_OPTIONS>)
target_include_directories(myapp PUBLIC
$<TARGET_PROPERTY:MKL::MKL,INTERFACE_INCLUDE_DIRECTORIES>)
target_link_libraries(myapp PUBLIC $<LINK_ONLY:MKL::MKL>)

add_test(NAME mytest COMMAND myapp)
if(MKL_ENV)
    set_tests_properties(mytest PROPERTIES ENVIRONMENT "${MKL_ENV}")
endif()
```

Command line

```
# Source the compiler and runtime beforehand
build$ cmake .. -G Ninja -DCMAKE_C_COMPILER=icl
build$ cmake --build . && ctest
# If MKLConfig.cmake is not located by CMake automatically, its path can be manually specified
by MKL_DIR:
build$ cmake .. -G Ninja -DCMAKE_C_COMPILER=icl -DMKL_DIR=<Full path to MKLConfig.cmake>
```

NOTE

Ninja and NMake Makefiles are supported on Windows. Other CMake Generators may work but are not guaranteed.

When the Ninja build system is in use, Ninja 1.10.2+ is required for Fortran support.

Checking Your Installation

After installing the , verify that the library is properly installed and configured:

1. Intel® oneAPI Math Kernel Library installs in the *<parent directory>* directory.

Check that the subdirectory of *<parent directory>* referred to as *<mkl directory>* was created.

Check that subdirectories for Intel® oneAPI Math Kernel Library redistributable DLLs *redist\ia32* and *redist\intel64* were created in the *<parent directory>* directory (See *redist.txt* in the Intel® oneAPI Math Kernel Library documentation directory for a list of files that can be redistributed.)

2. If you want to keep multiple versions of Intel® oneAPI Math Kernel Library installed on your system, update your build scripts to point to the correct Intel® oneAPI Math Kernel Library version.
3. Check that the *vars.bat* file appears in the *<mkl directory>\env* directory.

Use this file to assign Intel® oneAPI Math Kernel Library-specific values to several environment variables, as explained in [Scripts to Set Environment Variables Setting Environment Variables](#) .

4. To understand how the Intel® oneAPI Math Kernel Library directories are structured, see [Structure of the Intel® Math Kernel Library](#).

5. To make sure that Intel® oneAPI Math Kernel Library runs on your system, launch an Intel® oneAPI Math Kernel Library example, as explained in [Using Code Examples](#).

See Also

[Notational Conventions](#)

Setting Environment Variables

When the installation of Intel® oneAPI Math Kernel Library for Windows* is complete, set environment variables using [oneAPI setvars.sh](#). The script will source the `INCLUDE`, `MKLROOT`, `LD_LIBRARY_PATH`, `LIBRARY_PATH`, `CPATH`, `NLSPATH`, and `PKG_CONFIG_PATH`.

The script accepts the oneMKL-specific parameters, explained in the following table:

Setting Specified	Required (Yes/No)	Possible Values	Comment
Architecture	Yes, when applicable	intel64	
Use of Intel® oneAPI Math Kernel Library Fortran modules precompiled with the Intel® Visual Fortran compiler	No	mod	Supply this parameter only if you are using this compiler.
Programming interface (LP64 or ILP64)	No	lp64, default ilp64	

For example:

- The command `setvars ia32` sets the environment for Intel® oneAPI Math Kernel Library to use the Intel 32 architecture.
- The command `setvars intel64 mod ilp64` sets the environment for Intel® oneAPI Math Kernel Library to use the Intel 64 architecture, ILP64 programming interface, and Fortran modules.
- The command `setvars intel64 mod` sets the environment for Intel® oneAPI Math Kernel Library to use the Intel 64 architecture, LP64 interface, and Fortran modules.

NOTE

Supply the parameter specifying the architecture first, if it is needed. Values of the other two parameters can be listed in any order.

See Also

[High-level Directory Structure](#)

[Intel® oneAPI Math Kernel Library Interface Libraries and Modules](#)

[Fortran 95 Interfaces to LAPACK and BLAS](#)

[Setting the Number of Threads Using an OpenMP* Environment Variable](#)

Compiler Support

Intel® oneAPI Math Kernel Library supports compilers identified in the *Release Notes*. However, the library has been successfully used with other compilers as well.

Intel® oneAPI Math Kernel Library provides the `cdecl` interface (default interface of the Microsoft Visual C* application) for the IA-32 architecture.

When building Intel® oneAPI Math Kernel Library code examples, you can select a compiler:

- For Fortran examples: Intel® or PGI* compiler
- For C examples: Intel, Microsoft Visual C++*, or PGI compiler

Intel® oneAPI Math Kernel Library provides a set of include files to simplify program development by specifying enumerated values and prototypes for the respective functions. Calling Intel® oneAPI Math Kernel Library functions from your application without an appropriate include file may lead to incorrect behavior of the functions.

See Also

[Intel® oneAPI Math Kernel Library Include Files](#)

Using Code Examples

The Intel® oneAPI Math Kernel Library package includes code examples, located in the `examples` subdirectory of the installation directory. Use the examples to determine:

- Whether Intel® oneAPI Math Kernel Library is working on your system
- How you should call the library
- How to link the library

If an Intel® oneAPI Math Kernel Library component that you selected during installation includes code examples, these examples are provided in a separate archive. Extract the examples from the archives before use.

For each component, the examples are grouped in subdirectories mainly by Intel® oneAPI Math Kernel Library function domains and programming languages. For instance, the `blas` subdirectory (extracted from the `examples_core` archive) contains a makefile to build the BLAS examples and the `vmlc` subdirectory contains the makefile to build the C examples for Vector Mathematics functions. Source code for the examples is in the next-level `sources` subdirectory.

See Also

[High-level Directory Structure](#)

What You Need to Know Before You Begin Using the Intel® oneAPI Math Kernel Library

Target platform	Identify the architecture of your target machine: • IA-32 or compatible • Intel® 64 or compatible Reason: Because Intel® oneAPI Math Kernel Library libraries are located in directories corresponding to your particular architecture (see Architecture Support), you should provide proper paths on your link lines (see Linking Examples). To configure your development environment for the use with Intel® oneAPI Math Kernel Library, set your environment variables using the script corresponding to your architecture (see Scripts to Set Environment Variables Setting Environment Variables for details).
Mathematical problem	Identify all Intel® oneAPI Math Kernel Library function domains that you require: • BLAS • Sparse BLAS • LAPACK

- PBLAS
- ScaLAPACK
- Sparse Solver routines
- Parallel Direct Sparse Solvers for Clusters
- Vector Mathematics functions (VM)
- Vector Statistics functions (VS)
- Fourier Transform functions (FFT)
- Cluster FFT
- Trigonometric Transform routines
- Poisson, Laplace, and Helmholtz Solver routines
- Optimization (Trust-Region) Solver routines
- Data Fitting Functions
- Extended Eigensolver Functions

Reason: The function domain you intend to use narrows the search in the *Intel® oneAPI Math Kernel Library Developer Reference* for specific routines you need. Additionally, if you are using the Intel® oneAPI Math Kernel Library cluster software, your link line is function-domain specific (see [Working with the Intel® oneAPI Math Kernel Library Cluster Software](#)). Coding tips may also depend on the function domain (see [Other Tips and Techniques to Improve Performance](#)).

Programming language

Intel® oneAPI Math Kernel Library provides support for both Fortran and C/C++ programming. Identify the language interfaces that your function domains support (see [Appendix A: Intel® oneAPI Math Kernel Library Language Interfaces Support](#)).

Reason: Intel® oneAPI Math Kernel Library provides language-specific include files for each function domain to simplify program development (see [Language Interfaces Support_ by Function Domain](#)).

For a list of language-specific interface libraries and modules and an example how to generate them, see also [Using Language-Specific Interfaces with Intel® oneAPI Math Kernel Library](#).

Range of integer data

If your system is based on the Intel 64 architecture, identify whether your application performs calculations with large data arrays (of more than $2^{31}-1$ elements).

Reason: To operate on large data arrays, you need to select the ILP64 interface, where integers are 64-bit; otherwise, use the default, LP64, interface, where integers are 32-bit (see [Using the ILP64 Interface vs.](#)).

Threading model

Identify whether and how your application is threaded:

- Threaded with the Intel compiler
- Threaded with a third-party compiler
- Not threaded

Reason: The compiler you use to thread your application determines which threading library you should link with your application. For applications threaded with a third-party compiler you may need to use Intel® oneAPI Math Kernel Library in the sequential mode (for more information, see [Linking with Threading Libraries](#)).

Number of threads

If your application uses an OpenMP* threading run-time library, determine the number of threads you want Intel® oneAPI Math Kernel Library to use.

Reason: By default, the OpenMP* run-time library sets the number of threads for Intel® oneAPI Math Kernel Library. If you need a different number, you have to set it yourself using one of the available mechanisms. For more information, see [Improving Performance with Threading](#).

Linking model

Decide which linking model is appropriate for linking your application with Intel® oneAPI Math Kernel Library libraries:

- Static
- Dynamic

Reason: The link libraries for static and dynamic linking are different. For the list of link libraries for static and dynamic models, linking examples, and other relevant topics, like how to save disk space by creating a custom dynamic library, see [Linking Your Application with the Intel® oneAPI Math Kernel Library](#).

MPI used

Decide what MPI you will use with the Intel® oneAPI Math Kernel Library cluster software. You are strongly encouraged to use the latest available version of Intel® MPI.

Reason: To link your application with ScaLAPACK and/or Cluster FFT, the libraries corresponding to your particular MPI should be listed on the link line (see [Working with the Intel® oneAPI Math Kernel Library Cluster Software](#)).

Structure of the Intel® oneAPI Math Kernel Library

2

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Architecture Support

forWindows* provides architecture-specific implementations for supported platforms. The following table lists the supported architectures and directories where each architecture-specific implementation is located.

Architecture	Location
IA-32 or compatible	<code><mkl directory>\lib\ia32</code> <code><mkl directory>\redist\ia32</code>
Intel® 64 or compatible	<code><mkl directory>\lib\intel64</code> <code><mkl directory>\redist\intel64 (DLLs)</code>

See Also

[High-level Directory Structure](#)

[Notational Conventions](#)

[Detailed Structure of the IA-32 Architecture Directories](#)

[Detailed Structure of the Intel® 64 Architecture Directories](#)

High-level Directory Structure

Directory	Contents
<code><mkl directory></code>	Installation directory of the
Subdirectories of <code><mkl directory></code>	
<code>env\vars.bat</code>	Source script to set environment variables
<code>benchmarks\linpack</code>	Shared-Memory (SMP) version of the LINPACK benchmark
<code>benchmarks \mp_linpack</code>	Message-passing interface (MPI) version of the LINPACK benchmark
<code>lib\ia32</code>	Static libraries and static interfaces to DLLs for the IA-32 architecture
<code>lib\intel64</code>	Static libraries and static interfaces to DLLs for the Intel® 64 architecture
<code>examples</code>	Source and data files for Intel® oneAPI Math Kernel Library examples. Provided in archives corresponding to Intel® oneAPI Math Kernel Library components selected during installation.

Directory	Contents
include	Include files for the library routines and examples
include\ia32	Fortran 95 .mod files for the IA-32 architecture and Intel Fortran compiler
include\intel64\lp64	Fortran 95 .mod files for the Intel® 64 architecture, Intel® Fortran compiler, and LP64 interface
include\intel64\ilp64	Fortran 95 .mod files for the Intel® 64 architecture, Intel Fortran compiler, and ILP64 interface
include\fftw	Header files for the FFTW2 and FFTW3 interfaces
include\oneapi	Header files for DPC++ interfaces.
interfaces\blas95	Fortran 95 interfaces to BLAS and a makefile to build the library
interfaces\fftw2x_cdft	MPI FFTW 2.x interfaces to Intel® oneAPI Math Kernel Library Cluster FFT
interfaces\fftw3x_cdft	MPI FFTW 3.x interfaces to Intel® oneAPI Math Kernel Library Cluster FFT
interfaces\fftw2xc	FFTW 2.x interfaces to the Intel® oneAPI Math Kernel Library FFT (C interface)
interfaces\fftw2xf	FFTW 2.x interfaces to the Intel® oneAPI Math Kernel Library FFT (Fortran interface)
interfaces\fftw3xc	FFTW 3.x interfaces to the Intel® oneAPI Math Kernel Library FFT (C interface)
interfaces\fftw3xf	FFTW 3.x interfaces to the Intel® oneAPI Math Kernel Library FFT (Fortran interface)
interfaces\lapack95	Fortran 95 interfaces to LAPACK and a makefile to build the library
interfaces\mklmpi	Tool to create a custom MKLMPI wrapper library (BLACS) for use in MKL MPI-based applications like Cluster Sparse Solver and Scalapack.
tools	Command-line link tool and tools for creating custom dynamically linkable libraries
tools\builder	Tools for creating custom dynamically linkable libraries
redist\ia32	DLLs for applications running on processors with the IA-32 architecture
redist\intel64	DLLs for applications running on processors with Intel® 64 architecture

See Also

[Notational Conventions](#)

[Using Code Examples](#)

Layered Model Concept

Intel® oneAPI Math Kernel Library is structured to support multiple compilers and interfaces, both serial and multi-threaded modes, different implementations of threading run-time libraries, and a wide range of processors. Conceptually Intel® oneAPI Math Kernel Library can be divided into distinct parts to support different interfaces, threading models, and core computations:

1. Interface Layer
2. Threading Layer

3. Computational Layer

You can combine Intel® oneAPI Math Kernel Library libraries to meet your needs by linking with one library in each part layer-by-layer.

To support threading with different compilers, you also need to use an appropriate threading run-time library (RTL). These libraries are provided by compilers and are not included in Intel® oneAPI Math Kernel Library.

The following table provides more details of each layer.

Layer	Description
Interface Layer	This layer matches compiled code of your application with the threading and/or computational parts of the library. This layer provides: • cdecl interface. • LP64 and ILP64 interfaces. • Compatibility with compilers that return function values differently.
Threading Layer	This layer: • Provides a way to link threaded Intel® oneAPI Math Kernel Library with supported compilers. • Enables you to link with a threaded or sequential mode of the library. This layer is compiled for different environments (threaded or sequential) and compilers (from Intel).
Computational Layer	This layer accommodates multiple architectures through identification of architecture features and chooses the appropriate binary code at run time.

See Also

[Using the ILP64 Interface vs. LP64 Interface](#)

[Linking Your Application with the Intel® oneAPI Math Kernel Library](#)

[Linking with Threading Libraries](#)

Linking Your Application with the Intel® oneAPI Math Kernel Library

3

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Linking Quick Start

provides several options for quick linking of your application. The simplest options depend on your development environment:

Intel® Parallel Studio XE Composer Edition compiler

see [Using the /Qmkl Compiler Option](#).

Microsoft Visual Studio* Integrated Development Environment (IDE)

see [Automatically Linking a Project in the Visual Studio* Integrated Development Environment with Intel® MKL](#).

Other options are independent of your development environment, but depend on the way you link:

Explicit dynamic linking

see [Using the Single Dynamic Library](#) for how to simplify your link line.

Explicitly listing libraries on your link line

see [Selecting Libraries to Link with](#) for a summary of the libraries.

Using pkg-config tool to get compilation and link lines

see [Using pkg-config metadata files](#) for a summary on how to use Intel® oneAPI Math Kernel Library pkg-config metadata files.

Using an interactive interface

see [Using the Link-line Advisor](#) to determine libraries and options to specify on your link or compilation line.

Using an internally provided tool

see [Using the Command-line Link Tool](#) to determine libraries, options, and environment variables or even compile and build your application.

Using the /Qmkl Compiler Option

The Intel® Parallel Studio XE Composer Edition compiler supports the following variants of the `/Qmkl` compiler option:

`/Qmkl` or
`/Qmkl:parallel`

to link with a certain Intel® oneAPI Math Kernel Library threading layer depending on the threading option provided:

- For `-qopenmp` the OpenMP threading layer for Intel compilers
- For `-tbb` the Intel® Threading Building Blocks (Intel® TBB) threading layer

<code>/Qmkl:sequential</code>	to link with sequential version of Intel® oneAPI Math Kernel Library.
<code>/Qmkl:cluster</code>	to link with Intel® oneAPI Math Kernel Library cluster components (sequential) that use Intel MPI.

NOTE

The `-qopenmp` option has higher priority than `-tbb` in choosing the Intel® oneAPI Math Kernel Library threading layer for linking.

For more information on the `/Qmkl` compiler option, see the Intel Compiler User and Reference Guides.

For each variant of the `/Qmkl` option, the compiler links your application using the following conventions:

- `cdecl` for the IA-32 architecture
- `LP64` for the Intel® 64 architecture

If you specify any variant of the `/Qmkl` compiler option, the compiler automatically includes the Intel® oneAPI Math Kernel Library libraries. In cases not covered by the option, use the Link-line Advisor or see [Linking in Detail](#).

See Also

[Intel® oneAPI DPC++/C++ Compiler Developer Guide and Reference](#)

[Intel® Fortran Compiler Classic and Intel® Fortran Compiler Developer Guide and Reference](#)

[Using the ILP64 Interface vs. LP64 Interface](#)

[Using the Link-line Advisor](#)

[Intel® Software Documentation Library](#) for Intel® compiler documentation
for Intel® compiler documentation

Automatically Linking a Project in the Visual Studio* Integrated Development Environment with Intel® oneAPI Math Kernel Library

After a default installation of the or Intel® Parallel Studio XE Composer Edition, you can easily configure your project to automatically link with Intel® oneAPI Math Kernel Library.

Automatically Linking Your Microsoft Visual C/C++* Project with oneMKL

Configure your Microsoft Visual C/C++* project for automatic linking with Intel® oneAPI Math Kernel Library as follows:

1. Go to **Project>Properties>Configuration Properties>Intel Performance Libraries**.
2. Change the **Use MKL** property setting by selecting **Parallel**, **Sequential**, or **Cluster** as appropriate.

Specific Intel® oneAPI Math Kernel Library libraries that link with your application may depend on more project settings. For details, see the documentation for Intel® Parallel Studio XE Composer Edition for C++.

See Also

[Intel® Software Documentation Library](#) for the documentation for Intel® Parallel Studio XE Composer Edition
for the documentation for Intel® Parallel Studio XE Composer Edition

Automatically Linking Your Intel® Visual Fortran Project with oneMKL

Configure your Intel® Visual Fortran project for automatic linking with Intel® oneAPI Math Kernel Library as follows:

Go to **Project > Properties > Libraries > Use Intel Math Kernel Library** and select **Parallel**, **Sequential**, or **Cluster** as appropriate.

Specific Intel® oneAPI Math Kernel Library libraries that link with your application may depend on more project settings. For details see the documentation for Intel® Parallel Studio XE Composer Edition for Fortran.

See Also

[Intel® Software Documentation Library](#) for the documentation for Intel® Parallel Studio XE Composer Edition
for the documentation for Intel® Parallel Studio XE Composer Edition

Using the Single Dynamic Library

You can simplify your link line through the use of the Intel® oneAPI Math Kernel Library Single Dynamic Library (SDL).

To use SDL, place `mkl_rt.lib` on your link line. For example:

```
icl.exe application.c mkl_rt.lib
```

`mkl_rt.lib` is the import library for `mkl_rt.dll`.

SDL enables you to select the interface and threading library for Intel® oneAPI Math Kernel Library at run time. By default, linking with SDL provides:

- Intel LP64 interface on systems based on the Intel® 64 architecture
- Intel threading

To use other interfaces or change threading preferences, including use of the sequential version of Intel® oneAPI Math Kernel Library, you need to specify your choices using functions or environment variables as explained in section [Dynamically Selecting the Interface and Threading Layer](#).

NOTE Intel® oneAPI Math Kernel Library SDL (`mkl_rt`) does not support DPC++ APIs. If your application requires support of Intel® oneAPI Math Kernel Library DPC++ APIs, refer to [Intel® oneAPI Math Kernel Library Link-line Advisor](#) to configure your link command.

Selecting Libraries to Link with

To link with Intel® oneAPI Math Kernel Library:

- Choose one library from the Interface layer and one library from the Threading layer
- Add the only library from the Computational layer and run-time libraries (RTL)

The following table lists Intel® oneAPI Math Kernel Library libraries to link with your application.

	Interface layer	Threading layer	Computational layer	RTL
Intel® 64 architecture, static linking	<code>mkl_intel_lp64.lib</code>	<code>mkl_intel_thread.lib</code>	<code>mkl_core.lib</code>	<code>libiomp5md.lib</code>
Intel® 64 architecture, dynamic linking	<code>mkl_intel_lp64_dll.lib</code>	<code>mkl_intel_thread_dll.lib</code>	<code>mkl_core_dll.lib</code>	<code>libiomp5md.lib</code>
Intel® Many Integrated Core	<code>libmkl_intel_lp64.a</code>	<code>libmkl_intel_thread.a</code>	<code>libmkl_core.a</code>	<code>libiomp5.so</code>

	Interface layer	Threading layer	Computational layer	RTL
Architecture (Intel® MIC Architecture), static linking				
Intel MIC Architecture, dynamic linking	libmkl_intel_lp64.so	libmkl_intel_thread.so	libmkl_core.so	libiomp5.so

The Single Dynamic Library (SDL) automatically links interface, threading, and computational libraries and thus simplifies linking. The following table lists Intel® oneAPI Math Kernel Library libraries for dynamic linking using SDL. See [Dynamically Selecting the Interface and Threading Layer](#) for how to set the interface and threading layers at run time through function calls or environment settings.

	SDL	RTL
Intel® 64 architecture	mkl_rt.lib	libiomp5md.lib [†]

[†]Linking with libiomp5md.lib is not required.

For exceptions and alternatives to the libraries listed above, see [Linking in Detail](#).

See Also

[Layered Model Concept](#)

[Using the Link-line Advisor](#)

[Using the /Qmkl Compiler Option](#)

[Working with the Cluster Software](#)

Using the Link-line Advisor

Use the Intel® oneAPI Math Kernel Library Link-line Advisor to determine the libraries and options to specify on your link or compilation line.

The latest version of the tool is available at [Link Line Advisor for Intel® oneAPI Math Kernel Library](#). The tool is also available in the documentation directory of the product.

The Advisor requests information about your system and on how you intend to use Intel® oneAPI Math Kernel Library (link dynamically or statically, use threaded or sequential mode, and so on). The tool automatically generates the appropriate link line for your application.

See Also

[High-level Directory Structure](#)

Using the Command-line Link Tool

Use the command-line Link tool provided by Intel® oneAPI Math Kernel Library to simplify building your application with Intel® oneAPI Math Kernel Library.

The tool not only provides the options, libraries, and environment variables to use, but also performs compilation and building of your application.

The tool `mkl_link_tool.exe` is installed in the `<mkl_directory>\bin<arch>` directory, and supports the modes described in the following table.

Intel oneMKL Command-line Link Tool Modes

Mode	Description	Usage	Example
Inquiry	The tool returns the compiler options, libraries, or environment variables necessary to build and execute the application.	Get Intel® oneAPI Math Kernel Library libraries	<code>mkl_link_tool -libs [Intel oneMKL Link Tool options]</code>
		Get compilation options	<code>mkl_link_tool -opts [Intel oneMKL Link Tool options]</code>
		Get environment variables for application executable	<code>mkl_link_tool -env [Intel oneMKL Link Tool options]</code>
Compilation	The Intel® oneAPI Math Kernel Library Link Tool builds the application.	—	<code>mkl_link_tool [options] <compiler> [options2] file1 [file2 ...]</code> where: • <code>options</code> represents any number of Link Tool options • <code>compiler</code> represents the compiler name: <code>ifort</code>, <code>icc</code> (<code>icpc</code>, <code>icl</code>), <code>cl</code>, <code>gcc</code> (<code>g++</code>), <code>gfortran</code>, <code>pgcc</code> (<code>pgcp</code>), <code>pgf77</code> (<code>pgf90</code>, <code>pgf95</code>, <code>pgfortran</code>), <code>mpiic</code>, <code>mpiifort</code>, <code>mpic</code> (<code>mpic++</code>), <code>mpif77</code> (<code>mpif90</code>, <code>mpif95</code>), <code>dpcpp</code> • <code>options2</code> represents any number of compiler options
Interactive	Allows you to go through all possible Intel® oneAPI Math Kernel Library Link Tool supported options. The output provides libraries, options, or environment variables as in the inquiry mode, or a built application as in the compilation mode (depending on what you specify).	—	<code>mkl_link_tool -interactive</code>

Use the `-help` option for full help with the Intel® oneAPI Math Kernel Library Link Tool and to show the defaults for the current system.

Linking Examples

See Also

[Using the Link-line Advisor](#)

[Examples for Linking with ScaLAPACK and Cluster FFT](#)

Linking on IA-32 Architecture Systems

The following examples illustrate linking that uses Intel(R) compilers.

Most examples use the `.f` Fortran source file. C/C++ users should instead specify a `.cpp` (C++) or `.c` (C) file and replace `ifort` with `icl`:

- Static linking of `myprog.f` and OpenMP* threaded Intel® oneAPI Math Kernel Library supporting the `cdecl` interface:

```
ifort myprog.f mkl_intel_c.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib
```

- Dynamic linking of `myprog.f` and OpenMP* threaded Intel® oneAPI Math Kernel Library supporting the `cdecl` interface:

```
ifort myprog.f mkl_intel_c_dll.lib mkl_intel_thread_dll.lib mkl_core_dll.lib  
libiomp5md.lib
```

- Static linking of `myprog.f` and sequential version of Intel® oneAPI Math Kernel Library supporting the `cdecl` interface:

```
ifort myprog.f mkl_intel_c.lib mkl_sequential.lib mkl_core.lib
```

- Dynamic linking of `myprog.f` and sequential version of Intel® oneAPI Math Kernel Library supporting the `cdecl` interface:

```
ifort myprog.f mkl_intel_c_dll.lib mkl_sequential_dll.lib mkl_core_dll.lib
```

- Static linking of `myprog.f`, Fortran BLAS and 95 LAPACK interfaces, and OpenMP* threaded Intel® oneAPI Math Kernel Library supporting the `cdecl` interface:

```
ifort myprog.f mkl_lapack95.lib mkl_intel_c.lib mkl_intel_thread.lib mkl_core.lib  
libiomp5md.lib
```

- Static linking of `myprog.c` and Intel® oneAPI Math Kernel Library threaded with Intel® Threading Building Blocks (Intel® TBB), provided that the `LIB` environment variable contains the path to Intel TBB library:

```
icl myprog.c /link /libpath:$MKLPATH -I$MKLINCLUDE mkl_intel.lib mkl_tbb_thread.lib  
mkl_core.lib tbb12.lib /MD
```

- Dynamic linking of `myprog.c` and Intel® oneAPI Math Kernel Library threaded with Intel® TBB, provided that the `LIB` environment variable contains the path to Intel® TBB library:

```
icl myprog.c /link /libpath:$MKLPATH -I$MKLINCLUDE mkl_intel_dll.lib  
mkl_tbb_thread_dll.lib mkl_core_dll.lib tbb12.lib /MD
```

See Also

[Fortran 95 Interfaces to LAPACK and BLAS](#)

[Examples for linking a C application using cluster components](#)

[Examples for linking a Fortran application using cluster components](#)

[Using the Single Dynamic Library](#)

Linking on Intel(R) 64 Architecture Systems

The following examples illustrate linking that uses Intel(R) compilers.

Most examples use the `.f` Fortran source file. C/C++ users should instead specify a `.cpp` (C++) or `.c` (C) file and replace `ifort` with `icl`:

- Static linking of `myprog.f` and OpenMP* threadedIntel® oneAPI Math Kernel Library supporting the LP64 interface:

```
ifort myprog.f mkl_intel_lp64.lib mkl_intel_thread.lib mkl_core.lib
libiomp5md.lib
```

- Dynamic linking of `myprog.f` and OpenMP* threadedIntel® oneAPI Math Kernel Library supporting the LP64 interface:

```
ifort myprog.f mkl_intel_lp64_dll.lib mkl_intel_thread_dll.lib mkl_core_dll.lib
libiomp5md.lib
```

- Static linking of `myprog.f` and sequential version of Intel® oneAPI Math Kernel Library supporting the LP64 interface:

```
ifort myprog.f mkl_intel_lp64.lib mkl_sequential.lib mkl_core.lib
```

- Dynamic linking of `myprog.f` and sequential version of Intel® oneAPI Math Kernel Library supporting the LP64 interface:

```
ifort myprog.f mkl_intel_lp64_dll.lib mkl_sequential_dll.lib mkl_core_dll.lib
```

- Static linking of `myprog.f` and OpenMP* threadedIntel® oneAPI Math Kernel Library supporting the ILP64 interface:

```
ifort myprog.f mkl_intel_ilp64.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib
```

- Dynamic linking of `myprog.f` and OpenMP* threadedIntel® oneAPI Math Kernel Library supporting the ILP64 interface:

```
ifort myprog.f mkl_intel_ilp64_dll.lib mkl_intel_thread_dll.lib mkl_core_dll.lib
libiomp5md.lib
```

- Dynamic linking of user code `myprog.f` and OpenMP* threaded or sequential Intel® oneAPI Math Kernel Library supporting the LP64 or ILP64 interface (Call appropriate functions or set environment variables to choose threaded or sequential mode and to set the interface):

```
ifort myprog.f mkl_rt.lib
```

- Static linking of `myprog.f`, Fortran BLAS and 95 LAPACK interfaces, and OpenMP* threadedIntel® oneAPI Math Kernel Library supporting the LP64 interface:

```
ifort myprog.f mkl_lapack95_lp64.lib mkl_intel_lp64.lib mkl_intel_thread.lib
mkl_core.lib libiomp5md.lib
```

- Static linking of `myprog.c` and Intel® oneAPI Math Kernel Library threaded with Intel® Threading Building Blocks (Intel® TBB), provided that the `LIB` environment variable contains the path to Intel TBB library:

```
icl myprog.c /link /libpath:%MKLPATH% -I%MKLINCLUDE% mkl_intel_lp64.lib
mkl_tbb_thread.lib mkl_core.lib tbb12.lib /MD
```

- Dynamic linking of `myprog.c` and Intel® oneAPI Math Kernel Library threaded with Intel® TBB, provided that the `LIB` environment variable contains the path to Intel® TBB library:

```
icl myprog.c /link /libpath:%MKLPATH% -I%MKLINCLUDE% mkl_intel_lp64_dll.lib
mkl_tbb_thread_dll.lib mkl_core_dll.lib tbb12.lib /MD
```

See Also

[Fortran 95 Interfaces to LAPACK and BLAS](#)

[Examples for linking a C application using cluster components](#)

[Examples for linking a Fortran application using cluster components](#)

[Using the Single Dynamic Library](#)

Linking in Detail

This section recommends which libraries to link with depending on your Intel® oneAPI Math Kernel Library usage scenario and provides details of the linking.

Dynamically Selecting the Interface and Threading Layer

The Single Dynamic Library (SDL) enables you to dynamically select the interface and threading layer for Intel® oneAPI Math Kernel Library.

Setting the Interface Layer

To set the interface layer at run time, use the `mkl_set_interface_layer` function or the `MKL_INTERFACE_LAYER` environment variable.

Available interface layers depend on the architecture of your system.

On systems based on the Intel® 64 architecture, LP64 and ILP64 interfaces are available. The following table provides values to be used to set each interface layer.

Specifying the Interface Layer

Interface Layer	Value of <code>MKL_INTERFACE_LAYER</code>	Value of the Parameter of <code>mkl_set_interface_layer</code>
Intel LP64, default	LP64	<code>MKL_INTERFACE_LP64</code>
Intel ILP64	ILP64	<code>MKL_INTERFACE_ILP64</code>

If the `mkl_set_interface_layer` function is called, the environment variable `MKL_INTERFACE_LAYER` is ignored.

See the *Intel® oneAPI Math Kernel Library Developer Reference* for details of the `mkl_set_interface_layer` function.

On systems based on the IA-32 architecture, the `cdecl` interface is available.

Setting the Threading Layer

To set the threading layer at run time, use the `mkl_set_threading_layer` function or the `MKL_THREADING_LAYER` environment variable. The following table lists available threading layers along with the values to be used to set each layer.

Specifying the Threading Layer

Threading Layer	Value of <code>MKL_THREADING_LAYER</code>	Value of the Parameter of <code>mkl_set_threading_layer</code>
Intel threading, default	INTEL	<code>MKL_THREADING_INTEL</code>
Sequential mode of Intel® oneAPI Math Kernel Library	SEQUENTIAL	<code>MKL_THREADING_SEQUENTIAL</code>
PGI threading [†]	PGI	<code>MKL_THREADING_PGI</code>
Intel TBB threading	TBB	<code>MKL_THREADING_TBB</code>

[†] Not supported by the SDL for Intel® Many Integrated Core Architecture.

If the `mkl_set_threading_layer` function is called, the environment variable `MKL_THREADING_LAYER` is ignored.

See the *Intel® oneAPI Math Kernel Library Developer Reference* for details of the `mkl_set_threading_layer` function.

Replacing Error Handling and Progress Information Routines

You can replace the Intel® oneAPI Math Kernel Library error handling routine `xerbla` or progress information routine `mkl_progress` with your own function. If you are using SDL, to replace `xerbla` or `mkl_progress`, call the `mkl_set_xerbla` and `mkl_set_progress` function, respectively. See the *Intel® oneAPI Math Kernel Library Developer Reference* for details.

NOTE

If you are using SDL, you cannot perform the replacement by linking the object file with your implementation of `xerbla` or `mkl_progress`.

See Also

[Using the Single Dynamic Library](#)

[Layered Model Concept](#)

[Directory Structure in Detail](#)

Linking with Interface Libraries

Using the ILP64 Interface vs. LP64 Interface

The Intel® oneAPI Math Kernel Library ILP64 libraries use the 64-bit integer type (necessary for indexing large arrays, with more than $2^{31}-1$ elements), whereas the LP64 libraries index arrays with the 32-bit integer type.

The LP64 and ILP64 interfaces are implemented in the Interface layer. Link with the following interface libraries for the LP64 or ILP64 interface, respectively:

- `mkl_intel_lp64.lib` or `mkl_intel_ilp64.lib` for static linking
- `mkl_intel_lp64_dll.lib` or `mkl_intel_ilp64_dll.lib` for dynamic linking

The ILP64 interface provides for the following:

- Support large data arrays (with more than $2^{31}-1$ elements)
- Enable compiling your Fortran code with the `/4I8` compiler option

The LP64 interface provides compatibility with the previous Intel® oneAPI Math Kernel Library versions because "LP64" is just a new name for the only interface that the Intel® oneAPI Math Kernel Library versions lower than 9.1 provided. Choose the ILP64 interface if your application uses Intel® oneAPI Math Kernel Library for calculations with large data arrays or the library may be used so in the future.

On 64-bit platforms, selected domains provide API extensions with the `_64` suffix (for example, `SGEMM_64`) for supporting large data arrays in the LP64 library. This enables you to mix data types in one application. The selected domains and APIs include the following:

- BLAS: Fortran-style APIs for C applications and CBLAS APIs with integer arguments
- LAPACK: Fortran-style APIs for C applications and LAPACKE APIs with integer arguments

Intel® oneAPI Math Kernel Library provides the same include directory for the ILP64 and LP64 interfaces.

Compiling for LP64/ILP64

The table below shows how to compile for the ILP64 and LP64 interfaces:

Fortran

Compiling for ILP64 `ifort /4I8 /I<mkl directory>\include ...`

Compiling for LP64 `ifort /I<mkl directory>\include ...`

C or C++

Compiling for ILP64 `icl /DMKL_ILP64 /I<mkl directory>\include ...`

Compiling for LP64 `icl /I<mkl directory>\include ...`

Caution

Linking of an application compiled with the `/4I8` or `/DMKL_ILP64` option to the LP64 libraries may result in unpredictable consequences and erroneous output.

Coding for ILP64

You do not need to change existing code if you are not using the ILP64 interface.

To migrate to ILP64 or write new code for ILP64, use appropriate types for parameters of the Intel® oneAPI Math Kernel Library functions and subroutines:

Integer Types	Fortran	C or C++
32-bit integers	<code>INTEGER*4</code> or <code>INTEGER (KIND=4)</code>	<code>int</code>
Universal integers for ILP64/ LP64: 64-bit for ILP64 32-bit otherwise	<code>INTEGER</code> without specifying <code>KIND</code>	<code>MKL_INT</code>
Universal integers for ILP64/ LP64: 64-bit integers	<code>INTEGER*8</code> or <code>INTEGER (KIND=8)</code>	<code>MKL_INT64</code>
FFT interface integers for ILP64/ LP64	<code>INTEGER</code> without specifying <code>KIND</code>	<code>MKL_LONG</code>

To determine the type of an integer parameter of a function, use appropriate include files. For functions that support only a Fortran interface, use the C/C++ include files `*.h`.

The above table explains which integer parameters of functions become 64-bit and which remain 32-bit for ILP64. The table applies to most Intel® oneAPI Math Kernel Library functions except some Vector Mathematics and Vector Statistics functions, which require integer parameters to be 64-bit or 32-bit regardless of the interface:

- **Vector Mathematics:** The `mode` parameter of the functions is 64-bit.
- **Random Number Generators (RNG):**

All discrete RNG except `viRngUniformBits64` are 32-bit.

The `viRngUniformBits64` generator function and `vs1SkipAheadStream` service function are 64-bit.

- **Summary Statistics:** The *estimate* parameter of the `vs1sSSCompute/vs1dSSCompute` function is 64-bit.

Refer to the *Intel® oneAPI Math Kernel Library Developer Reference* for more information.

To better understand ILP64 interface details, see also examples.

Limitations

All Intel® oneAPI Math Kernel Library function domains support ILP64 programming but FFTW interfaces to Intel® oneAPI Math Kernel Library:

- FFTW 2.x wrappers do not support ILP64.
- FFTW 3.x wrappers support ILP64 by a dedicated set of functions `plan_guru64`.

See Also

[High-level Directory Structure](#)

[Intel® oneAPI Math Kernel Library Include Files](#)

[Language Interfaces Support, by Function Domain](#)

[Layered Model Concept](#)

[Directory Structure in Detail](#)

Linking with Fortran 95 Interface Libraries

The `mk1_blas95*.lib` and `mk1_lapack95*.lib` libraries contain Fortran 95 interfaces for BLAS and LAPACK, respectively, which are compiler-dependent. In the Intel® oneAPI Math Kernel Library package, they are prebuilt for the Intel® Fortran compiler. If you are using a different compiler, build these libraries before using the interface.

See Also

[Fortran 95 Interfaces to LAPACK and BLAS](#)

[Compiler-dependent Functions and Fortran 90 Modules](#)

Linking with Threading Libraries

Intel® oneAPI Math Kernel Library threading layer defines how Intel® oneAPI Math Kernel Library functions utilize multiple computing cores of the system that the application runs on. You must link your application with one appropriate Intel® oneAPI Math Kernel Library library in this layer, as explained below. Depending on whether this is a threading or a sequential library, Intel® oneAPI Math Kernel Library runs in a parallel or sequential mode, respectively.

In the *parallel mode*, Intel® oneAPI Math Kernel Library utilizes multiple processor cores available on your system, uses the OpenMP* or Intel TBB threading technology, and requires a proper threading runtime library (RTL) to be linked with your application. Independently of use of Intel® oneAPI Math Kernel Library, the application may also require a threading RTL. You should link not more than one threading RTL to your application. Threading RTLs are provided by your compiler. Intel® oneAPI Math Kernel Library provides several threading libraries, each dependent on the threading RTL of a certain compiler, and your choice of the Intel® oneAPI Math Kernel Library threading library must be consistent with the threading RTL that you use in your application.

The OpenMP RTL of the Intel® compiler is the `libiomp5md.lib` library, located under `<parent directory>\compiler\lib`. You can find additional information about the Intel OpenMP RTL at <https://www.openmpRTL.org>.

The Intel TBB RTL of the Intel® compiler is the `tbb12.lib` library, located under `<parent directory>\tbb\lib`. You can find additional information about the Intel TBB RTL at <https://www.threadingbuildingblocks.org>.

In the *sequential mode*, Intel® oneAPI Math Kernel Library runs unthreaded code, does not require an threading RTL, and does not respond to environment variables and functions controlling the number of threads. Avoid using the library in the sequential mode unless you have a particular reason for that, such as the following:

- Your application needs a threading RTL that none of Intel® oneAPI Math Kernel Library threading libraries is compatible with
- Your application is already threaded at a top level, and using parallel Intel® oneAPI Math Kernel Library only degrades the application performance by interfering with that threading
- Your application is intended to be run on a single thread, like a message-passing interface (MPI) application

It is critical to link the application with the proper RTL. The table below explains what library in the Intel® oneAPI Math Kernel Library threading layer and what threading RTL you should choose under different scenarios:

Application		Intel® oneAPI Math Kernel Library		RTL Required
Uses OpenMP	Compiled with	Execution Mode	Threading Layer	
no	any compiler	parallel	Static linking: mkl_intel_thread.lib Dynamic linking: mkl_intel_thread_dll.lib	libiomp5md.lib
no	any compiler	parallel	Static linking: mkl_tbb_thread.lib Dynamic linking: mkl_tbb_thread_dll.lib	tbb12.lib
no	any compiler	sequential	Static linking: mkl_sequential.lib Dynamic linking: mkl_sequential_dll.lib	none
yes	Intel compiler	parallel	Static linking: mkl_intel_thread.lib Dynamic linking: mkl_intel_thread_dll.lib	libiomp5md.lib
yes	PGI* compiler	parallel	Static linking: mkl_pgi_thread.lib Dynamic linking: mkl_pgi_	PGI OpenMP RTL

Application		Intel® oneAPI Math Kernel Library		RTL Required
Uses OpenMP	Compiled with	Execution Mode	Threading Layer	
yes	any other compiler	parallel	thread_dll.lib Not supported. Use Intel® oneAPI Math Kernel Library in the sequential mode.	

See Also

[Layered Model Concept](#)

[Notational Conventions](#)

Linking with Computational Libraries

If you are not using the Intel® oneAPI Math Kernel Library ScaLAPACK and Cluster Fast Fourier Transforms (FFT), you need to link your application with only one computational library, depending on the linking method:

Static Linking	Dynamic Linking
mk1_core.lib	mk1_core_dll.lib

Computational Libraries for Applications that Use ScaLAPACK or Cluster FFT

ScaLAPACK and Cluster FFT require more computational libraries, which may depend on your architecture.

The following table lists computational libraries for IA -32 architecture applications that use ScaLAPACK or Cluster FFT.

Computational Libraries for IA-32 Architecture

Function domain	Static Linking	Dynamic Linking
ScaLAPACK [†]	mk1_scalapack_core.lib	mk1_scalapack_core_dll.lib
	mk1_core.lib	mk1_core_dll.lib
Cluster Fourier Transform Functions [†]	mk1_cdft_core.lib	mk1_cdft_core_dll.lib
	mk1_core.lib	mk1_core_dll.lib

[†] Also add the library with BLACS routines corresponding to the MPI used.

The following table lists computational libraries for Intel® 64 or Intel® Many Integrated Core Architecture applications that use ScaLAPACK or Cluster FFT.

Computational Libraries for the Intel® 64 or Intel® Many Integrated Core Architecture

Function domain	Static Linking	Dynamic Linking
ScaLAPACK, LP64 interface [†]	mk1_scalapack_lp64.lib	mk1_scalapack_lp64_dll.lib
	mk1_core.lib	mk1_core_dll.lib
ScaLAPACK, ILP64 interface [†]	mk1_scalapack_ilp64.lib	mk1_scalapack_ilp64_dll.lib
	mk1_core.lib	mk1_core_dll.lib

Function domain	Static Linking	Dynamic Linking
Cluster Fourier Transform Functions [‡]	mkl_cdft_core.lib mkl_core.lib	mkl_cdft_core_dll.lib mkl_core_dll.lib

[‡] Also add the library with BLACS routines corresponding to the MPI used.

See Also

[Linking with ScaLAPACK and Cluster FFT](#)

[Using the Link-line Advisor](#)

[Using the ILP64 Interface vs. LP64 Interface](#)

Linking with Compiler Run-time Libraries

Dynamically link `libiomp5` or `tbb` library even if you link other libraries statically.

Linking to the `libiomp5` statically can be problematic because the more complex your operating environment or application, the more likely redundant copies of the library are included. This may result in performance issues (oversubscription of threads) and even incorrect results.

To link `libiomp5` or `tbb` dynamically, be sure the `PATH` environment variable is defined correctly.

Sometimes you may improve performance of your application with threaded Intel® oneAPI Math Kernel Library by using the `/MT` compiler option. The compiler driver will pass the option to the linker and the latter will load multi-thread (MT) static run-time libraries.

However, to link a Vector Mathematics (VM) application that uses the `errno` variable for error reporting, compile and link your code using the option that depends on the linking model:

- `/MT` for linking with static Intel® oneAPI Math Kernel Library libraries
- `/MD` for linking with dynamic Intel® oneAPI Math Kernel Library libraries

See Also

[Setting Environment Variables](#)

[Layered Model Concept](#)

Linking with System Libraries

If your system is based on the Intel® 64 architecture, be aware that Microsoft SDK builds 1289 or higher provide the `bufferoverflowu.lib` library to resolve the `__security_cookie` external references. Makefiles for examples and tests include this library by using the `buf_lib=bufferoverflowu.lib` macro. If you are using older SDKs, leave this macro empty on your command line as follows: `buf_lib=` .

See Also

[Linking Examples](#)

Building Custom Dynamic-link Libraries

Custom dynamic-link libraries (DLL) reduce the collection of functions available in Intel® oneAPI Math Kernel Library libraries to those required to solve your particular problems, which helps to save disk space and build your own dynamic libraries for distribution.

The Intel® oneAPI Math Kernel Library customDLL builder enables you to create a dynamic library containing the selected functions and located in the `tools\builder` directory. The builder contains a makefile and a definition file with the list of functions.

Using the Custom Dynamic-link Library Builder in the Command-line Mode

To build a custom DLL, use the following command:

```
nmake target [<options>]
```

The following table lists possible values of `target` and explains what the command does for each value:

Value	Comment
<code>libia32</code>	The builder uses static Intel® oneAPI Math Kernel Library interface, threading, and core libraries to build a customDLL for the IA-32 architecture.
<code>libintel64</code>	The builder uses static Intel® oneAPI Math Kernel Library interface, threading, and core libraries to build a customDLL for the Intel® 64 architecture.
<code>dllia32</code>	The builder uses the single dynamic library <code>libmkl_rt.dll</code> to build a custom DLL for the IA-32 architecture.
<code>dllintel64</code>	The builder uses the single dynamic library <code>libmkl_rt.dll</code> to build a custom DLL for the Intel® 64 architecture.
<code>help</code>	The command prints Help on the custom DLL builder

The `<options>` placeholder stands for the list of parameters that define macros to be used by the makefile. The following table describes these parameters:

Parameter [Values]	Description
<code>interface</code>	Defines which programming interface to use. Possible values: - For the IA-32 architecture, <code>{cdecl}</code>. The default value is <code>cdecl</code>. - For the Intel 64 architecture, <code>{lp64 ilp64}</code>. The default value is <code>lp64</code>.
<code>threading = {parallel sequential}</code>	Defines whether to use the Intel® oneAPI Math Kernel Library in the threaded or sequential mode. The default value is <code>parallel</code> .
<code>Parallel = {intel tbb}</code>	Specifies whether to use Intel OpenMP or Intel® oneTBB. The default value is <code>intel</code> .
<code>cluster = {yes no}</code>	(For <code>libintel64</code> only) Specifies whether Intel® oneAPI Math Kernel Library cluster components (BLACS, ScaLAPACK and/or CDFT) are needed to build the custom shared object. The default value is <code>no</code> .
<code>blacs_mpi = {intelmpi msmpi}</code>	Specifies the pre-compiled Intel® oneAPI Math Kernel Library BLACS library to use. Ignored if <code>cluster=no</code> . The default value is <code>intelmpi</code> .
<code>blacs_name = <lib name></code>	Specifies the name (without extension) of a custom Intel® oneAPI Math Kernel Library BLACS library to use. Ignored if <code>cluster=no</code> . <code>'blacs_mpi'</code> is ignored if <code>'blacs_name'</code> was explicitly specified. The default value is <code>mkl_blacs_<blacs_mpi>_<interface></code> .
<code>mpi = <lib name></code>	Specifies the name (without extension) of the MPI library used to build the custom DLL. Ignored if <code>'cluster=no'</code> . The default value is <code>impi</code> .
<code>export = <file name></code>	Specifies the full name of the file that contains the list of entry-point functions to be included in the DLL. The default name is <code>user_example_list</code> (no extension).
<code>name = <dll name></code>	Specifies the name of the dll and interface library to be created. By default, the names of the created libraries are <code>mkl_custom.dll</code> and <code>mkl_custom.lib</code> .

Parameter [Values]	Description
<code>xerbla =</code> <code><error handler></code>	Specifies the name of the object file <code><user_xerbla>.obj</code> that contains the error handler of the user. The makefile adds this error handler to the library for use instead of the default Intel® oneAPI Math Kernel Library error handler <code>xerbla</code> . If you omit this parameter, the native Intel® oneAPI Math Kernel Library <code>xerbla</code> is used. See the description of the <code>xerbla</code> function in the Intel® oneAPI Math Kernel Library Developer Reference to develop your own error handler. For the IA-32 architecture, the object file should be in the interface defined by the interface macro (<code>cdecl</code>).
<code>MKLROOT =</code> <code><mkl directory></code>	Specifies the location of Intel® oneAPI Math Kernel Library libraries used to build the custom DLL. By default, the builder uses the Intel® oneAPI Math Kernel Library installation directory.
<code>uwd_compat =</code> <code>{yes no}</code>	Build a Universal Windows Driver (UWD)-compatible custom DLL with <code>OneCore.lib</code> . The recommended versions of Windows SDK are 10.0.17134.0 or higher. If 'uwd_compat'=yes, then <code>threading = sequential</code> and <code>crt = ucrt.lib</code> by default. The default value of is <code>uwd_compat</code> is no.
<code>buf_lib</code>	Manages resolution of the <code>__security_cookie</code> external references in the custom DLL on systems based on the Intel® 64 architecture. By default, the makefile uses the <code>bufferoverflowu.lib</code> library of Microsoft SDK builds 1289 or higher. This library resolves the <code>__security_cookie</code> external references. To avoid using this library, set the empty value of this parameter. Therefore, if you are using an older SDK, set <code>buf_lib=</code> . Caution Use the <code>buf_lib</code> parameter only with the empty value. Incorrect value of the parameter causes builder errors.
<code>crt = <c run-time library></code>	Specifies the name of the Microsoft C run-time library to be used to build the custom DLL. By default, the builder uses <code>msvcrt.lib</code> .
<code>manifest =</code> <code>{yes no embed}</code>	Manages the creation of a Microsoft manifest for the custom DLL: • If <code>manifest=yes</code>, the manifest file with the name defined by the <code>name</code> parameter above and the <code>manifest</code> extension is created. • If <code>manifest=no</code>, the manifest file is not be created. • If <code>manifest=embed</code>, the manifest is embedded into the DLL. By default, the builder does not use the <code>manifest</code> parameter.

All of the above parameters are optional. However, you must make the system and c-runtime (crt) libraries and link.exe available by setting the `PATH` and `LIB` environment variables appropriately. You can do this in the following ways:

- Manually
- If you are using Microsoft Visual Studio (VS), call the `vcvarsall.bat` script with the appropriate 32-bit (x86) or 64-bit (x64 or amd-64) architecture flag.
- If you are using the Intel compiler, use the `compilervars.bat` script with the appropriate 32-bit (x86) or 64-bit (x64 or amd-64) architecture flag.

In the simplest case, the command line is:

```
#source Visual Studio environment variables
call vcvarsall.bat x86
#run custom dll builder script
nmake ia32
```

and the missing options have default values. This command creates the `mkl_custom.dll` and `mkl_custom.lib` libraries with the `cdecl` interface for processors using the IA-32 architecture. The command takes the list of functions from the `functions_list` file and uses the native Intel® oneAPI Math Kernel Library error handler `xerbla`.

Here is an example of a more complex case:

```
#source Visual Studio environment variables
call vcvarsall.bat x86
#run custom dll builder script
nmake ia32 interface=cdecl export=my_func_list.txt name=mkl_small xerbla=my_xerbla.obj
```

In this case, the command creates the `mkl_small.dll` and `mkl_small.lib` libraries with the `cdecl` interface for processors using the IA-32 architecture. The command takes the list of functions from `my_func_list.txt` file and uses the error handler of the user `my_xerbla.obj`.

To build a UWD-compatible custom dll, use the `uwd_compat=yes` option. For this purpose, you must make a different set of universal system (`OneCore.lib`) and universal c-runtime (`ucrt.lib`) libraries available. You can get these libraries by downloading Windows 10 SDK 10.0.17134.0 (version 1803) or newer. Make sure to source the Visual Studio environment with the appropriate native architecture to add the libraries to your path.

This example shows how to create a 64-bit architecture library, `mkl_uwd_compat.dll`, that is UWD-compatible with the `lp64` interface using `my_function_list.txt` for specific functionality:

```
#source Visual Studio environment variables, LIB should have paths to desired OneCore.lib and
universal crt libraries
call vcvarsall.bat x64
#run custom dll builder script
nmake intel64 interface=lp64 export=my_func_list.txt uwd_compat=yes name=mkl_uwd_compat
```

See Also

[Linking with System Libraries](#)

Composing a List of Functions

To compose a list of functions for a minimal custom DLL needed for your application, you can use the following procedure:

1. Link your application with installed Intel® oneAPI Math Kernel Library libraries to make sure the application builds.
2. Remove all Intel® oneAPI Math Kernel Library libraries from the link line and start linking. Unresolved symbols indicate Intel® oneAPI Math Kernel Library functions that your application uses.
3. Include these functions in the list.

Important

Each time your application starts using more Intel® oneAPI Math Kernel Library functions, update the list to include the new functions.

See Also

Specifying Function Names

Specifying Function Names

In the file with the list of functions for your custom DLL, adjust function names to the required interface. For example, you can list the cdecl entry points as follows:

```
DGEMM
DTRSM
DDOT
DGETRF
DGETRS
cblas_dgemm
cblas_ddot
```

For more examples, see domain-specific lists of function names in the `<mkl_directory>\tools\builder` folder. This folder contains lists of function names for cdecl interfaces.

NOTE

The lists of function names are provided in the `<mkl_directory>\tools\builder` folder merely as examples. See [Composing a List of Functions](#) for how to compose lists of functions for your custom DLL.

Tip

Names of Fortran-style routines (BLAS, LAPACK, etc.) can be both upper-case or lower-case, with or without the trailing underscore. For example, these names are equivalent:

BLAS: `dgemm`, `DGEMM`, `dgemm_`, `DGEMM_`
LAPACK: `dgetrf`, `DGETRF`, `dgetrf_`, `DGETRF_`.

Properly capitalize names of C support functions in the function list. To do this, follow the guidelines below:

1. In the `mkl_service.h` include file, look up a `#define` directive for your function (`mkl_service.h` is included in the `mkl.h` header file).
2. Take the function name from the replacement part of that directive.

For example, the `#define` directive for the `mkl_disable_fast_mm` function is
`#define mkl_disable_fast_mm MKL_Disable_Fast_MM.`

Capitalize the name of this function in the list like this: `MKL_Disable_Fast_MM`.

For the names of the Fortran support functions, see the [tip](#).

Building a Custom Dynamic-link Library in the Visual Studio* Development System

You can build a custom dynamic-link library (DLL) in the Microsoft Visual Studio* Development System (VS*). To do this, use projects available in the `tools\builder\MSVS_Projects` subdirectory of the Intel® oneAPI Math Kernel Library directory. The directory contains subdirectories with projects for the respective versions of the Visual Studio Development System, for example, `VS2012`. For each version of VS two solutions are available:

- `libia32.sln` builds a custom DLL for the IA-32 architecture.
- `libintel64.sln` builds a custom DLL for the Intel® 64 architecture.

The builder uses the following default settings for the custom DLL:

Interface:	cdecl for the IA-32 architecture and LP64 for the Intel 64 architecture
Error handler:	Native Intel® oneAPI Math Kernel Library <code>xerbla</code>
Create Microsoft manifest:	yes
List of functions:	in the project's source file <code>examples.def</code>

To build a custom DLL:

1. Set the `MKLROOT` environment variable with the installation directory of the Intel® oneAPI Math Kernel Library version you are going to use.
2. Open the `libia32.sln` or `libintel64.sln` solution depending on the architecture of your system.

The solution includes the following projects:

- `i_malloc_dll`
 - `vml_dll_core`
 - `cdecl_parallel` (in `libia32.sln`) or `lp64_parallel` (in `libintel64.sln`)
 - `cdecl_sequential` (in `libia32.sln`) or `lp64_sequential` (in `libintel64.sln`)
3. [Optional] To change any of the default settings, select the project depending on whether the DLL will use Intel® oneAPI Math Kernel Library functions in the sequential or multi-threaded mode:
 - In the `libia32` solution, select the `cdecl_sequential` or `cdecl_parallel` project.
 - In the `libintel64` solution, select the `lp64_sequential` or `lp64_parallel` project.
 4. [Optional] To include your own error handler in the DLL:
 - a. Select **Project>Properties>Configuration Properties>Linker>Input**.
 - b. Add `<user_xerbla>.obj`
 5. [Optional] To turn off creation of the manifest:
 - a. Select **Project>Properties>Configuration Properties>Linker>Manifest File>Generate Manifest**.
 - b. Select: no.
 6. [Optional] To change the list of functions to be included in the DLL:
 - a. Select **Source Files**.
 - b. Edit the `examples.def` file. Refer to [Specifying Function Names](#) for how to specify entry points.
 7. To build the library, select **Build>Build Solution**.

See Also

[Using the Custom Dynamic-link Library Builder in the Command-line Mode](#)

Distributing Your Custom Dynamic-link Library

To enable use of your custom DLL in a threaded mode, distribute `libiomp5md.dll` along with the custom DLL.

Building a Universal Windows Driver

Intel® oneAPI Math Kernel Library is compatible with Universal Windows Drivers (UWDs) whose binaries are Win10 SDK version 1803 (also called Redstone 4, RS4 or Win 10 SDK build 10.0.17134.0) or higher. The Windows system calls used in the statically linked libraries are defined in:

- `OneCore.Lib` for headless applications like Nano Server.
- `OneCoreUAP.Lib` for a UWD.

A Windows application or driver that links to either of these libraries can use functionality from statically linked Intel® oneAPI Math Kernel Library to accelerate computations. For more information on design criteria for a UWD, see the Microsoft website.

This table summarizes important information about compatibility between Intel® oneAPI Math Kernel Library and UWDs.

Compatibility between Intel MKL and Universal Windows Driver (UWD)

Criteria	Compatibility
SDK	10.0.17134.0 (also called Redstone 4, RS4, version 1803) or higher
Architecture	- ia32 - intel64
Universal API sets	- OneCore.Lib for headless applications like Nano Server - OneCoreUAP.Lib for a UWD
Threading layers	- Sequential - Intel® Threading Building Blocks (Intel® TBB). See Using Intel®TBB in UWP Applications.
UWD Compliance	- Intel® oneAPI Math Kernel Library with static linking libraries - For dynamic linking libraries, use the custom DLL builder with <code>uwd_compat=yes</code> to construct a custom DLL that is compatible with UWD. See Using the Custom DLL Builder.

NOTE The UWD links to Intel® oneAPI Math Kernel Library using the same link line settings as drivers or executables that are not universal. See [Link Line Advisor](#) for recommended linking models.

UWD Compliance for 32-bit Architecture (ia32)

These tables describe compliance between Intel® oneAPI Math Kernel Library and UWD on 32-bit (ia32) architecture for static and dynamic linked libraries.

Static Linking

Layer	Library	Compatibility
Fortran interface layer	mkl_blas95.lib mkl_lapack95.lib	Not applicable ¹
Interface layer	mkl_intel_c.lib (cdecl)	✓
Threading layer	mkl_sequential.lib mkl_tbb_thread.lib mkl_intel_thread.lib (OpenMP)	✓ ✓ (if using UWD Intel® TBB ²) ✗
Computational layer	mkl_core.lib	✓

¹ These are Fortran 95 wrappers for BLAS (BLAS95) and LAPACK (LAPACK95). Windows drivers are typically written in C, C++ or ASM. Fortran wrappers are UWD-compliant but not useful for building UWDs.

² See [Using Intel®TBB in UWP Applications](#). Starting with the Intel® TBB 2018 release, you can find a prebuilt UWD-compliant Intel® TBB library in `<tbb_distribution_root>\lib\<target_architecture>\vc14_uwd`.

Dynamic Linking

Layer	Library	Compatibility
Fortran interface layer	mkl_blas95_dll.lib mkl_lapack95_dll.lib	Not applicable ¹
Interface layer	mkl_intel_c_dll.lib (cdecl)	See UWD Compliance in the compatibility table above.
Threading layer	mkl_sequential_dll.lib mkl_tbb_thread_dll.lib mkl_intel_thread_dll.lib (OpenMP)	See UWD Compliance in the compatibility table above.
Computational layer	mkl_core_dll.lib	See UWD Compliance in the compatibility table above.

¹ These are Fortran 95 wrappers for BLAS (BLAS95) and LAPACK (LAPACK95). Windows drivers are typically written in C, C++ or ASM. Fortran wrappers are UWD-compliant but not useful for building UWDs.

UWD Compliance for 64-bit Architecture (intel64)

These tables describe compliance between Intel® oneAPI Math Kernel Library and UWD on 64-bit (intel64) architecture for static and dynamic linked libraries.

Static Linking

Layer	Library	Compatibility
Fortran interface layer	mkl_blas95_lp64.lib mkl_blas95_ilp64.lib mkl_lapack95_lp64.lib mkl_lapack95_ilp64.lib	Not applicable ¹
Interface layer	mkl_intel_lp64.lib mkl_intel_ilp64.lib	✓ ✓
Threading layer	mkl_sequential.lib mkl_tbb_thread.lib mkl_intel_thread.lib (OpenMP) mkl_pgi_thread.lib (PGI OpenMP)	✓ ✓ (if using UWD Intel® TBB ²) ✗ ✗
Computational layer	mkl_core.lib	✓
MPI layer	mkl_blacs_lp64.lib mkl_blacs_ilp64.lib mkl_cdft_core.lib mkl_scalapack_lp64.lib	✗ ✗ ✗ ✗

Layer	Library	Compatibility
	mkl_scalapack_ilp64.lib	

¹ These are Fortran 95 wrappers for BLAS (BLAS95) and LAPACK (LAPACK95). Windows drivers are typically written in C, C++ or ASM. Fortran wrappers are UWD-compliant but not useful for building UWDs.

² See [Using Intel® TBB in UWP Applications](#). Starting with the Intel® TBB 2018 release, you can find a prebuilt UWD-compliant Intel® TBB library in <tbbs_distribution_root>\lib\<target_architecture>\vc14_uwd.

Dynamic Linking

Layer	Library	Compatibility
Fortran interface layer	mkl_blas95_lp64.lib mkl_blas95_ilp64.lib mkl_lapack95_lp64.lib mkl_lapack95_ilp64.lib	Not applicable ¹
Interface layer	mkl_intel_lp64_dll.lib mkl_intel_ilp64_dll.lib	See UWD Compliance in the compatibility table above.
Threading layer	mkl_sequential_dll.lib mkl_tbb_thread_dll.lib mkl_intel_thread_dll.lib (OpenMP) mkl_pgi_thread_dll.lib (PGI OpenMP)	See UWD Compliance in the compatibility table above.
Computational layer	mkl_core_dll.lib	See UWD Compliance in the compatibility table above.
MPI layer	mkl_blacs_lp64_dll.lib mkl_blacs_ilp64_dll.lib mkl_cdft_core_dll.lib mkl_scalapack_lp64_dll.lib b mkl_scalapack_ilp64_dll.lib ib	See UWD Compliance in the compatibility table above.

¹ These are Fortran 95 wrappers for BLAS (BLAS95) and LAPACK (LAPACK95). Windows drivers are typically written in C, C++ or ASM. Fortran wrappers are UWD-compliant but not useful for building UWDs.

Managing Performance and Memory

4

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Improving Performance with Threading

is extensively parallelized. See [OpenMP* Threaded Functions and Problems](#) and [Functions Threaded with Intel® Threading Building Blocks](#) for lists of threaded functions and problems that can be threaded.

Intel® oneAPI Math Kernel Library is *thread-safe*, which means that all Intel® oneAPI Math Kernel Library functions (except the LAPACK deprecated routine `?lacon`) work correctly during simultaneous execution by multiple threads. In particular, any chunk of threaded Intel® oneAPI Math Kernel Library code provides access for multiple threads to the same shared data, while permitting only one thread at any given time to access a shared piece of data. Therefore, you can call Intel® oneAPI Math Kernel Library from multiple threads and not worry about the function instances interfering with each other.

If you are using OpenMP* threading technology, you can use the environment variable `OMP_NUM_THREADS` to specify the number of threads or the equivalent OpenMP run-time function calls. Intel® oneAPI Math Kernel Library also offers variables that are independent of OpenMP, such as `asmkl_num_threads`, and equivalent Intel® oneAPI Math Kernel Library functions for thread management. The Intel® oneAPI Math Kernel Library variables are always inspected first, then the OpenMP variables are examined, and if neither is used, the OpenMP software chooses the default number of threads.

By default, Intel® oneAPI Math Kernel Library uses the number of OpenMP threads equal to the number of physical cores on the system.

If you are using the Intel TBB threading technology, the OpenMP threading controls, such as the `OMP_NUM_THREADS` environment variable or `mkf_num_threads` function, have no effect. Use the Intel TBB application programming interface to control the number of threads.

To achieve higher performance, set the number of threads to the number of processors or physical cores, as summarized in [Techniques to Set the Number of Threads](#).

See Also

[Managing Multi-core Performance](#)

OpenMP* Threaded Functions and Problems

The following Intel® oneAPI Math Kernel Library function domains are threaded with the OpenMP* technology:

- Direct sparse solver.
- LAPACK.

For a list of threaded routines, see [LAPACK Routines](#).

- Level1 and Level2 BLAS.

For a list of threaded routines, see [BLAS Level1 and Level2 Routines](#).

- All Level 3 BLAS and all Sparse BLAS routines except Level 2 Sparse Triangular solvers.
- All Vector Mathematics functions (except service functions).
- FFT.

For a list of FFT transforms that can be threaded, see [Threaded FFT Problems](#).

LAPACK Routines

In this section, ? stands for a precision prefix of *each* flavor of the respective routine and may have the value of s, d, c, or z.

The following LAPACK routines are threaded with OpenMP*:

- Linear equations, computational routines:
 - Factorization: ?getrf, ?getrfnp, ?gbtrf, ?potrf, ?pptrf, ?sytrf, ?hetrf, ?sptrf, ?hptrf
 - Solving: ?dtttrs, ?gbtrs, ?gttrs, ?pptrs, ?pbtrs, ?pttrs, ?sytrs, ?sptrs, ?hptrs, ?tptrs, ?tbtrs
- Orthogonal factorization, computational routines:
 - ?geqrf, ?ormqr, ?unmqr, ?ormlq, ?unmlq, ?ormql, ?unmql, ?ormrq, ?unmrq
- Singular Value Decomposition, computational routines:
 - ?gebrd, ?bdsqr
- Symmetric Eigenvalue Problems, computational routines:
 - ?sytrd, ?hetrd, ?sptrd, ?hptrd, ?steqr, ?stedc.
- Generalized Nonsymmetric Eigenvalue Problems, computational routines:
 - chgeqz/zhgeqz.

A number of other LAPACK routines, which are based on threaded LAPACK or BLAS routines, make effective use of OpenMP* parallelism:

?gesv, ?posv, ?gels, ?gesvd, ?syev, ?heev, cgegs/zgegs, cgegv/zgegv, cgges/zgges, cggesx/zggesx, cggev/zggev, cggevz/zggevz, and so on.

Threaded BLAS Level1 and Level2 Routines

In the following list, ? stands for a precision prefix of *each* flavor of the respective routine and may have the value of s, d, c, or z.

The following routines are threaded with OpenMP*:

- Level1 BLAS:
 - ?axpy, ?copy, ?swap, ddot/sdot, cdotc, drot/srot
- Level2 BLAS:
 - ?gemv, ?trsv, ?trmv, dsyr/ssyr, dsyr2/ssyr2, dsymv/ssymv

Threaded FFT Problems

The following characteristics of a specific problem determine whether your FFT computation may be threaded with OpenMP*:

- rank
- domain
- size/length
- precision (single or double)
- placement (in-place or out-of-place)
- strides
- number of transforms
- layout (for example, interleaved or split layout of complex data)

Most FFT problems are threaded. In particular, computation of multiple transforms in one call (number of transforms > 1) is threaded. Details of which transforms are threaded follow.

One-dimensional (1D) transforms

1D transforms are threaded in many cases.

1D complex-to-complex (c2c) transforms of size N using interleaved complex data layout are threaded under the following conditions depending on the architecture:

Architecture	Conditions
Intel® 64	N is a power of 2, $\log_2(N) > 9$, the transform is double-precision out-of-place, and input/output strides equal 1.
IA-32	N is a power of 2, $\log_2(N) > 13$, and the transform is single-precision. N is a power of 2, $\log_2(N) > 14$, and the transform is double-precision.
Any	N is composite, $\log_2(N) > 16$, and input/output strides equal 1.

1D complex-to-complex transforms using split-complex layout are not threaded.

Multidimensional transforms

All multidimensional transforms on large-volume data are threaded.

Functions Threaded with Intel® Threading Building Blocks

In this section, α stands for a precision prefix or suffix of the routine name and may have the value of s , d , c , or z .

The following Intel® oneAPI Math Kernel Library function domains are threaded with Intel® Threading Building Blocks (Intel® TBB):

- LAPACK.
For a list of threaded routines, see [LAPACK Routines](#).
 - Entire Level3 BLAS.
 - Level2 BLAS – α GEMV.
 - Fast Poisson, Laplace, and Helmholtz Solver (Poisson Library).
 - All Vector Mathematics functions (except service functions).
 - Intel® oneAPI Math Kernel Library PARDISO, a direct sparse solver based on Parallel Direct Sparse Solver (PARDISO*).
- For details, see [oneMKL PARDISO Steps](#).
- Sparse BLAS.
For a list of threaded routines, see [Sparse BLAS Routines](#).

LAPACK Routines

The following LAPACK routines are threaded with Intel TBB:

α geqrf, α gelqf, α getrf, α potrf, α unmqr*, α ormqr*, α unmrq*, α ormrq*, α unmlq*, α ormlq*, α unmql*, α ormql*, α sytrd, α hetrd, α syev, α heev, and α latrd.

A number of other LAPACK routines, which are based on threaded LAPACK or BLAS routines, make effective use of Intel TBB threading:

α getrs, α gesv, α potrs, α bdsqr, and α gels.

oneMKL PARDISO Steps

Intel® oneAPI Math Kernel Library PARDISO is threaded with Intel TBB in the reordering and factorization steps. However, routines performing the solving step are still called sequentially when using Intel TBB.

Sparse BLAS Routines

The Sparse BLAS inspector-executor application programming interface routines `mkl_sparse_?_mv` are threaded with Intel TBB for the general compressed sparse row (CSR) and block sparse row (BSR) formats.

The following Sparse BLAS inspector-executor application programming routines are threaded with Intel TBB:

- `mkl_sparse_?_mv` using the general compressed sparse row (CSR) and block sparse row (BSR) matrix formats.
- `mkl_sparse_?_mm` using the general CSR sparse matrix format and both row and column major storage formats for the dense matrix.

Avoiding Conflicts in the Execution Environment

Certain situations can cause conflicts in the execution environment that make the use of threads in Intel® oneAPI Math Kernel Library problematic. This section briefly discusses why these problems exist and how to avoid them.

If your program is parallelized by other means than Intel® OpenMP* run-time library (RTL) and Intel TBB RTL, several calls to Intel® oneAPI Math Kernel Library may operate in a multithreaded mode at the same time and result in slow performance due to overuse of machine resources.

The following table considers several cases where the conflicts may arise and provides recommendations depending on your threading model:

Threading model	Discussion
You parallelize the program using the technology other than Intel OpenMP and Intel TBB (for example: Win32* threads on Windows*).	If more than one thread calls Intel® oneAPI Math Kernel Library, and the function being called is threaded, it may be important that you turn off Intel® oneAPI Math Kernel Library threading. Set the number of threads to one by any of the available means (see Techniques to Set the Number of Threads).
You parallelize the program using OpenMP directives and/or pragmas and compile the program using a non-Intel compiler.	To avoid simultaneous activities of multiple threading RTLs, link the program against the Intel® oneAPI Math Kernel Library threading library that matches the compiler you use (see Linking Examples on how to do this). If this is not possible, use Intel® oneAPI Math Kernel Library in the sequential mode. To do this, you should link with the appropriate threading library: <code>mkl_sequential.lib</code> or <code>mkl_sequential.dll</code> (see Appendix C: Directory Structure in Detail).
You thread the program using Intel TBB threading technology and compile the program using a non-Intel compiler.	To avoid simultaneous activities of multiple threading RTLs, link the program against the Intel® oneAPI Math Kernel Library Intel TBB threading library and Intel TBB RTL if it matches the compiler you use. If this is not possible, use Intel® oneAPI Math Kernel Library in the sequential mode. To do this, link with the appropriate threading library: <code>mkl_sequential.lib</code> or <code>mkl_sequential_dll.lib</code> (see Appendix C: Directory Structure in Detail).
You run multiple programs calling Intel® oneAPI Math Kernel Library on a multiprocessor system, for example, a program parallelized using a message-passing interface (MPI).	The threading RTLs from different programs you run may place a large number of threads on the same processor on the system and therefore overuse the machine resources. In this case, one of the solutions is to set the number of threads to one by any of the available means (see Techniques to Set the Number of Threads). The Intel® Distribution for LINPACK* Benchmark section discusses another solution for a Hybrid (OpenMP* + MPI) mode.

Using the `mkl_set_num_threads` and `mkl_domain_set_num_threads` functions to control parallelism of Intel® oneAPI Math Kernel Library from parallel user threads may result in a race condition that impacts the performance of the application because these functions operate on internal control variables that are global, that is, apply to all threads. For example, if parallel user threads call these functions to set different numbers of threads for the same function domain, the number of threads actually set is unpredictable. To avoid this kind of data races, use the `mkl_set_num_threads_local` function (see the "Support Functions" chapter in the *Intel® oneAPI Math Kernel Library Developer Reference* for the function description).

See Also

[Using Additional Threading Control](#)

Linking with Compiler Support RTLs

Techniques to Set the Number of Threads

Use the following techniques to specify the number of OpenMP threads to use in Intel® oneAPI Math Kernel Library:

- Set one of the OpenMP or Intel® oneAPI Math Kernel Library environment variables:
 - `OMP_NUM_THREADS`
 - `MKL_NUM_THREADS`
 - `MKL_DOMAIN_NUM_THREADS`
- Call one of the OpenMP or Intel® oneAPI Math Kernel Library functions:
 - `omp_set_num_threads()`
 - `mkl_set_num_threads()`
 - `mkl_domain_set_num_threads()`
 - `mkl_set_num_threads_local()`

NOTE

A call to the `mkl_set_num_threads` or `mkl_domain_set_num_threads` function changes the number of OpenMP threads available to all in-progress calls (in concurrent threads) and future calls to Intel® oneAPI Math Kernel Library and may result in slow Intel® oneAPI Math Kernel Library performance and/or race conditions reported by run-time tools, such as Intel® Inspector.

To avoid such situations, use the `mkl_set_num_threads_local` function (see the "Support Functions" section in the *Intel® oneAPI Math Kernel Library Developer Reference* for the function description).

When choosing the appropriate technique, take into account the following rules:

- The Intel® oneAPI Math Kernel Library threading controls take precedence over the OpenMP controls because they are inspected first.
- A function call takes precedence over any environment settings. The exception, which is a consequence of the previous rule, is that a call to the OpenMP subroutine `omp_set_num_threads()` does not have precedence over the settings of Intel® oneAPI Math Kernel Library environment variables such as `MKL_NUM_THREADS`. See [Using Additional Threading Control](#) for more details.
- You cannot change run-time behavior in the course of the run using the environment variables because they are read only once at the first call to Intel® oneAPI Math Kernel Library.

If you use the Intel TBB threading technology, read the documentation for the `tbb::task_scheduler_init` class at <https://www.threadingbuildingblocks.org/documentation> to find out how to specify the number of threads.

Setting the Number of Threads Using an OpenMP* Environment Variable

You can set the number of threads using the environment variable `OMP_NUM_THREADS`. To change the number of OpenMP threads, in the command shell in which the program is going to run, enter:

```
set OMP_NUM_THREADS=<number of threads to use>.
```

Some shells require the variable and its value to be exported:

```
export OMP_NUM_THREADS=<number of threads to use>.
```

You can alternatively assign value to the environment variable using Microsoft Windows* OS Control Panel.

Note that you will not benefit from setting this variable on Microsoft Windows* 98 or Windows* ME because multiprocessing is not supported.

See Also

[Using Additional Threading Control](#)

Changing the Number of OpenMP* Threads at Run Time

You cannot change the number of OpenMP threads at run time using environment variables. However, you can call OpenMP routines to do this. Specifically, the following sample code shows how to change the number of threads during run time using the `omp_set_num_threads()` routine. For more options, see also [Techniques to Set the Number of Threads](#).

The example is provided for both C and Fortran languages. To run the example in C, use the `omp.h` header file from the Intel(R) compiler package. If you do not have the Intel compiler but wish to explore the functionality in the example, use Fortran API for `omp_set_num_threads()` rather than the C version. For example, `omp_set_num_threads_( &i_one );`

```
// ***** C language *****
#include "omp.h"
#include "mkl.h"
#include <stdio.h>
#define SIZE 1000
int main(int args, char *argv[]){
double *a, *b, *c;
a = (double*)malloc(sizeof(double)*SIZE*SIZE);
b = (double*)malloc(sizeof(double)*SIZE*SIZE);
c = (double*)malloc(sizeof(double)*SIZE*SIZE);
double alpha=1, beta=1;
int m=SIZE, n=SIZE, k=SIZE, lda=SIZE, ldb=SIZE, ldc=SIZE, i=0, j=0;
char transa='n', transb='n';
for( i=0; i<SIZE; i++)
{
    for( j=0; j<SIZE; j++)
    {
        a[i*SIZE+j]= (double) (i+j);
        b[i*SIZE+j]= (double) (i*j);
        c[i*SIZE+j]= (double)0;
    }
}
cblas_dgemm(CblasRowMajor, CblasNoTrans, CblasNoTrans,
m, n, k, alpha, a, lda, b, ldb, beta, c, ldc);
printf("row\ta\tc\n");
for ( i=0;i<10;i++)
{
printf("%d:\t%f\t%f\n", i, a[i*SIZE], c[i*SIZE]);
}
omp_set_num_threads(1);
for( i=0; i<SIZE; i++)
{
    for( j=0; j<SIZE; j++)
    {
        a[i*SIZE+j]= (double) (i+j);
        b[i*SIZE+j]= (double) (i*j);
        c[i*SIZE+j]= (double)0;
    }
}
cblas_dgemm(CblasRowMajor, CblasNoTrans, CblasNoTrans,
m, n, k, alpha, a, lda, b, ldb, beta, c, ldc);
printf("row\ta\tc\n");
for ( i=0;i<10;i++)
{
printf("%d:\t%f\t%f\n", i, a[i*SIZE], c[i*SIZE]);
}
```

```

}
omp_set_num_threads(2);
for( i=0; i<SIZE; i++)
{
    for( j=0; j<SIZE; j++)
    {
        a[i*SIZE+j]= (double) (i+j);
        b[i*SIZE+j]= (double) (i*j);
        c[i*SIZE+j]= (double) 0;
    }
}
cblas_dgemm(CblasRowMajor, CblasNoTrans, CblasNoTrans,
m, n, k, alpha, a, lda, b, ldb, beta, c, ldc);
printf("row\ta\tc\n");
for ( i=0;i<10;i++)
{
    printf("%d:\t%f\t%f\n", i, a[i*SIZE],
c[i*SIZE]);
}
free (a);
free (b);
free (c);
return 0;
}

```

```

// ***** Fortran language *****
PROGRAM DGEMM_DIFF_THREADS
INTEGER N, I, J
PARAMETER (N=100)
REAL*8 A(N,N),B(N,N),C(N,N)
REAL*8 ALPHA, BETA

ALPHA = 1.1
BETA = -1.2
DO I=1,N
    DO J=1,N
        A(I,J) = I+J
        B(I,J) = I*j
        C(I,J) = 0.0
    END DO
END DO
CALL DGEMM('N','N',N,N,N,ALPHA,A,N,B,N,BETA,C,N)
print *, 'Row A C'
DO i=1,10
write(*,'(I4,F20.8,F20.8)') I, A(1,I),C(1,I)
END DO
CALL OMP_SET_NUM_THREADS(1);
DO I=1,N
    DO J=1,N
        A(I,J) = I+J
        B(I,J) = I*j
        C(I,J) = 0.0
    END DO
END DO
CALL DGEMM('N','N',N,N,N,ALPHA,A,N,B,N,BETA,C,N)
print *, 'Row A C'

```

```

DO i=1,10
write(*,'(I4,F20.8,F20.8)') I, A(1,I),C(1,I)
END DO
CALL OMP_SET_NUM_THREADS(2);
DO I=1,N
    DO J=1,N
        A(I,J) = I+J
        B(I,J) = I*j
        C(I,J) = 0.0
    END DO
END DO
CALL DGEMM('N','N',N,N,N,ALPHA,A,N,B,N,BETA,C,N)
print *, 'Row A C'
DO i=1,10
write(*,'(I4,F20.8,F20.8)') I, A(1,I),C(1,I)
END DO
STOP
END

```

Using Additional Threading Control

oneMKL-specific Environment Variables for OpenMP Threading Control

Intel® oneAPI Math Kernel Library provides environment variables and support functions to control Intel® oneAPI Math Kernel Library threading independently of OpenMP. The Intel® oneAPI Math Kernel Library-specific threading controls take precedence over their OpenMP equivalents. Use the Intel® oneAPI Math Kernel Library-specific threading controls to distribute OpenMP threads between Intel® oneAPI Math Kernel Library and the rest of your program.

NOTE

Some Intel® oneAPI Math Kernel Library routines may use fewer OpenMP threads than suggested by the threading controls if either the underlying algorithms do not support the suggested number of OpenMP threads or the routines perform better with fewer OpenMP threads because of lower OpenMP overhead and/or better data locality. Set the `MKL_DYNAMIC` environment variable to `FALSE` or call `mkl_set_dynamic(0)` to use the suggested number of OpenMP threads whenever the algorithms permit and regardless of OpenMP overhead and data locality.

Section "Number of User Threads" in the "Fourier Transform Functions" chapter of the *Intel® oneAPI Math Kernel Library Developer References* shows how the Intel® oneAPI Math Kernel Library threading controls help to set the number of threads for the FFT computation.

The table below lists the Intel® oneAPI Math Kernel Library environment variables for threading control, their equivalent functions, and OMP counterparts:

Environment Variable	Support Function	Comment	Equivalent OpenMP* Environment Variable
MKL_NUM_THREADS	mkl_set_num_threads mkl_set_num_threads_local	Suggests the number of OpenMP threads to use.	OMP_NUM_THREADS

Environment Variable	Support Function	Comment	Equivalent OpenMP* Environment Variable
MKL_DOMAIN_NUM_THREADS	mkl_domain_set_num_threads	Suggests the number of OpenMP threads for a particular function domain.	
MKL_DYNAMIC	mkl_set_dynamic	Enables Intel® oneAPI Math Kernel Library to dynamically change the number of OpenMP threads.	OMP_DYNAMIC

NOTE

Call `mkl_set_num_threads()` to force Intel® oneAPI Math Kernel Library to use a given number of OpenMP threads and prevent it from reacting to the environment variables `MKL_NUM_THREADS`, `MKL_DOMAIN_NUM_THREADS`, and `OMP_NUM_THREADS`.

The example below shows how to force Intel® oneAPI Math Kernel Library to use one thread:

```
// ***** C language *****

#include <mkl.h>
...
mkl_set_num_threads ( 1 );
```

```
// ***** Fortran language *****

...
call mkl_set_num_threads( 1 )
```

See the *Intel® oneAPI Math Kernel Library Developer Reference* for the detailed description of the threading control functions, their parameters, calling syntax, and more code examples.

MKL_DYNAMIC

The `MKL_DYNAMIC` environment variable enables Intel® oneAPI Math Kernel Library to dynamically change the number of threads.

The default value of `MKL_DYNAMIC` is `TRUE`, regardless of `OMP_DYNAMIC`, whose default value may be `FALSE`.

When `MKL_DYNAMIC` is `TRUE`, Intel® oneAPI Math Kernel Library may use fewer OpenMP threads than the maximum number you specify.

For example, `MKL_DYNAMIC` set to `TRUE` enables optimal choice of the number of threads in the following cases:

- If the requested number of threads exceeds the number of physical cores (perhaps because of using the Intel® Hyper-Threading Technology), Intel® oneAPI Math Kernel Library scales down the number of OpenMP threads to the number of physical cores.
- If you are able to detect the presence of a message-passing interface (MPI), but cannot determine whether it has been called in a thread-safe mode, Intel® oneAPI Math Kernel Library runs one OpenMP thread.

When `MKL_DYNAMIC` is `FALSE`, Intel® oneAPI Math Kernel Library uses the suggested number of OpenMP threads whenever the underlying algorithms permit. For example, if you attempt to do a size one matrix-matrix multiply across eight threads, the library may instead choose to use only one thread because it is impractical to use eight threads in this event.

If Intel® oneAPI Math Kernel Library is called from an OpenMP parallel region in your program, Intel® oneAPI Math Kernel Library uses only one thread by default. If you want Intel® oneAPI Math Kernel Library to go parallel in such a call, link your program against an OpenMP threading RTL supported by Intel® oneAPI Math Kernel Library and set the environment variables:

- `OMP_NESTED` to `TRUE`
- `OMP_DYNAMIC` and `MKL_DYNAMIC` to `FALSE`
- `MKL_NUM_THREADS` to some reasonable value

With these settings, Intel® oneAPI Math Kernel Library uses `MKL_NUM_THREADS` threads when it is called from the OpenMP parallel region in your program.

In general, set `MKL_DYNAMIC` to `FALSE` only under circumstances that Intel® oneAPI Math Kernel Library is unable to detect, for example, to use nested parallelism where the library is already called from a parallel section.

MKL_DOMAIN_NUM_THREADS

The `MKL_DOMAIN_NUM_THREADS` environment variable suggests the number of OpenMP threads for a particular function domain.

`MKL_DOMAIN_NUM_THREADS` accepts a string value `<MKL-env-string>`, which must have the following format:

```
<MKL-env-string> ::= <MKL-domain-env-string> { <delimiter><MKL-domain-env-string> }
<delimiter> ::= [ <space-symbol>* ] ( <space-symbol> | <comma-symbol> | <semicolon-symbol> | <colon-symbol> ) [ <space-symbol>* ]
<MKL-domain-env-string> ::= <MKL-domain-env-name><uses><number-of-threads>
<MKL-domain-env-name> ::= MKL_DOMAIN_ALL | MKL_DOMAIN_BLAS | MKL_DOMAIN_FFT |
MKL_DOMAIN_VML | MKL_DOMAIN_PARDISO
<uses> ::= [ <space-symbol>* ] ( <space-symbol> | <equality-sign> | <comma-symbol> )
[ <space-symbol>* ]
<number-of-threads> ::= <positive-number>
<positive-number> ::= <decimal-positive-number> | <octal-number> | <hexadecimal-number>
```

In the syntax above, values of `<MKL-domain-env-name>` indicate function domains as follows:

<code>MKL_DOMAIN_ALL</code>	All function domains
<code>MKL_DOMAIN_BLAS</code>	BLAS Routines
<code>MKL_DOMAIN_FFT</code>	non-cluster Fourier Transform Functions
<code>MKL_DOMAIN_LAPACK</code>	LAPACK Routines
<code>MKL_DOMAIN_VML</code>	Vector Mathematics (VM)
<code>MKL_DOMAIN_PARDISO</code>	Intel® oneAPI Math Kernel Library PARDISO, a direct sparse solver based on Parallel Direct Sparse Solver (PARDISO*)

For example, you could set the `MKL_DOMAIN_NUM_THREADS` environment variable to any of the following string variants, in this case, defining three specific domain variables internal to Intel® oneAPI Math Kernel Library:

```
MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=2, MKL_DOMAIN_BLAS=1, MKL_DOMAIN_FFT=4"
```

```
MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL 2 : MKL_DOMAIN_BLAS 1 : MKL_DOMAIN_FFT 4"
```

```
MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=2 : MKL_DOMAIN_BLAS=1 : MKL_DOMAIN_FFT=4"
```

```

MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=2; MKL_DOMAIN_BLAS=1; MKL_DOMAIN_FFT=4"
MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=2 MKL_DOMAIN_BLAS 1, MKL_DOMAIN_FFT 4"
MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL,2: MKL_DOMAIN_BLAS 1, MKL_DOMAIN_FFT,4"

```

NOTE Prepend the appropriate `set/export/setenv` command for your command shell and operating system. Refer to [Setting the Environment Variables for Threading Control](#) for more details.

The global variables `MKL_DOMAIN_ALL`, `MKL_DOMAIN_BLAS`, `MKL_DOMAIN_FFT`, `MKL_DOMAIN_VML`, and `MKL_DOMAIN_PARDISO`, as well as the interface for the Intel® oneAPI Math Kernel Library threading control functions, can be found in the `themkl.h` header file.

NOTE You can retrieve the values of the specific domain variables that you have set in your code with a call to the `mkl_get_domain_max_threads(domain_name)` function per the [Fortran](#) and [C interface](#) with the desired domain variable name.

This table illustrates how values of `MKL_DOMAIN_NUM_THREADS` are interpreted.

Value of <code>MKL_DOMAIN_NUM_THREADS</code>	Interpretation
<code>MKL_DOMAIN_ALL=4</code>	All parts of Intel® oneAPI Math Kernel Library should try four OpenMP threads. The actual number of threads may be still different because of the <code>MKL_DYNAMIC</code> setting or system resource issues. The setting is equivalent to <code>MKL_NUM_THREADS = 4</code> .
<code>MKL_DOMAIN_ALL=1,</code> <code>MKL_DOMAIN_BLAS=4</code>	All parts of Intel® oneAPI Math Kernel Library should try one OpenMP thread, except for BLAS, which is suggested to try four threads.
<code>MKL_DOMAIN_VML=2</code>	VM should try two OpenMP threads. The setting affects no other part of Intel® oneAPI Math Kernel Library.

Be aware that the domain-specific settings take precedence over the overall ones. For example, the "`MKL_DOMAIN_BLAS=4`" value of `MKL_DOMAIN_NUM_THREADS` suggests trying four OpenMP threads for BLAS, regardless of later setting `MKL_NUM_THREADS`, and a function call "`mkl_domain_set_num_threads ( 4, MKL_DOMAIN_BLAS );`" suggests the same, regardless of later calls to `mkl_set_num_threads()`. However, a function call with input "`MKL_DOMAIN_ALL`", such as "`mkl_domain_set_num_threads (4, MKL_DOMAIN_ALL);`" is equivalent to "`mkl_set_num_threads(4)`", and thus it will be overwritten by later calls to `mkl_set_num_threads`. Similarly, the environment setting of `MKL_DOMAIN_NUM_THREADS` with "`MKL_DOMAIN_ALL=4`" will be overwritten with `MKL_NUM_THREADS = 2`.

Whereas the `MKL_DOMAIN_NUM_THREADS` environment variable enables you set several variables at once, for example, "`MKL_DOMAIN_BLAS=4,MKL_DOMAIN_FFT=2`", the corresponding function does not take string syntax. So, to do the same with the function calls, you may need to make several calls, which in this example are as follows:

```

mkl_domain_set_num_threads ( 4, MKL_DOMAIN_BLAS );
mkl_domain_set_num_threads ( 2, MKL_DOMAIN_FFT );

```

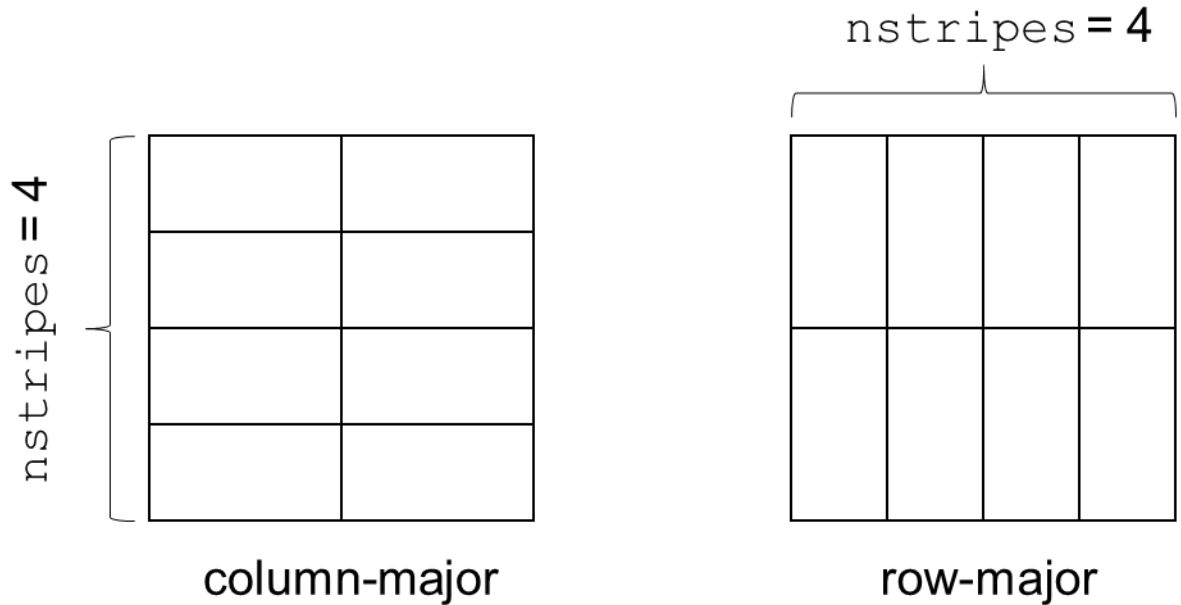
`MKL_NUM_STRIPES`

The `MKL_NUM_STRIPES` environment variable controls the Intel® oneAPI Math Kernel Library threading algorithm for `?gemm` functions. When `MKL_NUM_STRIPES` is set to a positive integer value *nstripes*, Intel® oneAPI Math Kernel Library tries to use a number of partitions equal to *nstripes* along the leading dimension of the output matrix.

The following table explains how the value *nstripes* of `MKL_NUM_STRIPES` defines the partitioning algorithm used by Intel® oneAPI Math Kernel Library for `gemm` output matrix; *max_threads_for_mkl* denotes the maximum number of OpenMP threads for Intel® oneAPI Math Kernel Library:

Value of <code>MKL_NUM_STRIPES</code>	Partitioning Algorithm
$1 < nstripes < (\text{max_threads_for_mkl} / 2)$	2D partitioning with the number of partitions equal to <i>nstripes</i> : - Horizontal, for column-major ordering. - Vertical, for row-major ordering.
$nstripes = 1$	1D partitioning algorithm along the opposite direction of the leading dimension.
$nstripes \geq (\text{max_threads_for_mkl} / 2)$	1D partitioning algorithm along the leading dimension.
$nstripes < 0$	The default Intel® oneAPI Math Kernel Library threading algorithm.

The following figure shows the partitioning of an output matrix for *nstripes* = 4 and a total number of 8 OpenMP threads for column-major and row-major orderings:



You can use support functions `mkl_set_num_stripes` and `mkl_get_num_stripes` to set and query the number of stripes, respectively.

Setting the Environment Variables for Threading Control

To set the environment variables used for threading control, in the command shell in which the program is going to run, enter:

```
set <VARIABLE NAME>=<value>
```

For example:

```
set MKL_NUM_THREADS=4
set MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=1, MKL_DOMAIN_BLAS=4"
set MKL_DYNAMIC=FALSE
set MKL_NUM_STRIPES=4
```

Some shells require the variable and its value to be exported:

```
export <VARIABLE NAME>=<value>
```

For example:

```
export MKL_NUM_THREADS=4
export MKL_DOMAIN_NUM_THREADS="MKL_DOMAIN_ALL=1, MKL_DOMAIN_BLAS=4"
export MKL_DYNAMIC=FALSE
export MKL_NUM_STRIPES=4
```

You can alternatively assign values to the environment variables using Microsoft Windows* OS Control Panel.

Calling oneMKL Functions from Multi-threaded Applications

This section summarizes typical usage models and available options for calling Intel® oneAPI Math Kernel Library functions from multi-threaded applications. These recommendations apply to any multi-threading environments: OpenMP*, Intel® Threading Building Blocks, Windows* threads, and others.

Usage model: disable oneMKL internal threading for the whole application

When used: Intel® oneAPI Math Kernel Library internal threading interferes with application's own threading or may slow down the application.

Example: the application is threaded at top level, or the application runs concurrently with other applications.

Options:

- Link statically or dynamically with the sequential library
- Link with the Single Dynamic Library `mk1_rt.lib` and select the sequential library using an environment variable or a function call:
 - Set `MKL_THREADING_LAYER=sequential`
 - Call `mk1_set_threading_layer(MKL_THREADING_SEQUENTIAL)`[‡]

Usage model: partition system resources among application threads

When used: application threads are specialized for a particular computation.

Example: one thread solves equations on all cores but one, while another thread running on a single core updates a database.

Linking Options:

- Link statically or dynamically with a threading library
- Link with the Single Dynamic Library `mk1_rt.lib` and select a threading library using an environment variable or a function call:
 - set `MKL_THREADING_LAYER=intel` or `MKL_THREADING_LAYER=tbb`
 - call `mk1_set_threading_layer(MKL_THREADING_INTEL)` or `mk1_set_threading_layer(MKL_THREADING_TBB)`

Other Options for OpenMP Threading:

- Set the `MKL_NUM_THREADS` environment variable to a desired number of OpenMP threads for Intel® oneAPI Math Kernel Library.
- Set the `MKL_DOMAIN_NUM_THREADS` environment variable to a desired number of OpenMP threads for Intel® oneAPI Math Kernel Library for a particular function domain.

Use if the application threads work with different Intel® oneAPI Math Kernel Library function domains.

- Call `mkl_set_num_threads()`
Use to globally set a desired number of OpenMP threads for Intel® oneAPI Math Kernel Library at run time.
- Call `mkl_domain_set_num_threads()`.
Use if at some point application threads start working with different Intel® oneAPI Math Kernel Library function domains.
- Call `mkl_set_num_threads_local()`.
Use to set the number of OpenMP threads for Intel® oneAPI Math Kernel Library called from a particular thread.

NOTE

If your application uses OpenMP* threading, you may need to provide additional settings:

- Set the environment variable `OMP_NESTED=TRUE`, or alternatively call `omp_set_nested(1)`, to enable OpenMP nested parallelism.
- Set the environment variable `MKL_DYNAMIC=FALSE`, or alternatively call `mkl_set_dynamic(0)`, to prevent Intel® oneAPI Math Kernel Library from dynamically reducing the number of OpenMP threads in nested parallel regions.

* For details of the mentioned functions, see the Support Functions section of the *Intel® oneAPI Math Kernel Library Developer Reference*, available in the Intel Software Documentation Library.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[Linking with Threading Libraries](#)

[Dynamically Selecting the Interface and Threading Layer](#)

[oneMKL-specific Environment Variables for OpenMP Threading Control](#)

[MKL_DOMAIN_NUM_THREADS](#)

[Avoiding Conflicts in the Execution Environment](#)

[Intel Software Documentation Library](#)

Using Intel® Hyper-Threading Technology

Intel® Hyper-Threading Technology (Intel® HT Technology) is especially effective when each thread performs different types of operations and when there are under-utilized resources on the processor. However, Intel® oneAPI Math Kernel Library fits neither of these criteria because the threaded portions of the library execute at high efficiencies using most of the available resources and perform identical operations on each thread. You may obtain higher performance by disabling Intel HT Technology.

If you run with Intel HT Technology enabled, performance may be especially impacted if you run on fewer threads than physical cores. Moreover, if, for example, there are two threads to every physical core, the thread scheduler may assign two threads to some cores and ignore the other cores altogether. If you are using the OpenMP* library of the Intel Compiler, read the respective User Guide on how to best set the thread affinity interface to avoid this situation. For Intel® oneAPI Math Kernel Library, apply the following setting:

```
export KMP_AFFINITY=granularity=fine,compact,1,0
```

If you are using the Intel TBB threading technology, read the documentation on the `tbb::affinity_partitioner` class at <https://www.threadingbuildingblocks.org/documentation> to find out how to affinity Intel TBB threads.

Managing Multi-core Performance

You can obtain best performance on systems with multi-core processors by requiring that threads do not migrate from core to core. To do this, bind threads to the CPU cores by setting an affinity mask to threads. Use one of the following options:

- OpenMP facilities (if available), for example, the `KMP_AFFINITY` environment variable using the Intel OpenMP library
- A system function, as explained below
- Intel TBB facilities (if available), for example, the `tbb::affinity_partitioner` class (for details, see <https://www.threadingbuildingblocks.org/documentation>)

Consider the following performance issue:

- The system has two sockets with two cores each, for a total of four cores (CPUs).
- The application sets the number of OpenMP threads to four and calls an Intel® oneAPI Math Kernel Library LAPACK routine. This call takes considerably different amounts of time from run to run.

To resolve this issue, before calling Intel® oneAPI Math Kernel Library, set an affinity mask for each OpenMP thread using the `KMP_AFFINITY` environment variable or the `SetThreadAffinityMask` system function. The following code example shows how to resolve the issue by setting an affinity mask by operating system means using the Intel compiler. The code calls the function `SetThreadAffinityMask` to bind the threads to appropriate cores, preventing migration of the threads. Then the Intel® oneAPI Math Kernel Library LAPACK routine is called:

```
// Set affinity mask
#include <windows.h>
#include <omp.h>
int main(void) {
    #pragma omp parallel default(shared)
    {
        int tid = omp_get_thread_num();
        // 2 packages x 2 cores/pkg x 1 threads/core (4 total cores)
        DWORD_PTR mask = (1 << (tid == 0 ? 0 : 2));
        SetThreadAffinityMask( GetCurrentThread(), mask );
    }
    // Call Intel MKL LAPACK routine
    return 0;
}
```

Compile the application with the Intel compiler using the following command:

```
icl /Qopenmp test_application.c
```

where `test_application.c` is the filename for the application.

Build the application. Run it in four threads, for example, by using the environment variable to set the number of threads:

```
set OMP_NUM_THREADS=4
test_application.exe
```

See Windows API documentation at msdn.microsoft.com/ for the restrictions on the usage of Windows API routines and particulars of the `SetThreadAffinityMask` function used in the above example.

See also a similar example at en.wikipedia.org/wiki/Affinity_mask.

Managing Performance with Heterogeneous Cores

A hybrid architecture offers heterogeneous CPU cores. For example, the 12th Gen Intel® Core™ processor (Alder Lake) contains two types of cores: Performance-cores (P-cores) and Efficient-cores (E-cores).

Achieving the best performance on a hybrid architecture is harder because load balancing with heterogeneous cores is more complicated. Therefore, for hybrid architectures like Alder Lake, we recommend running threads on the P-cores only. This approach might not yield the best performance, but it is simple and predictable.

To specify P-cores with OpenMP, users can use the environment variable KMP_HW_SUBSET. For a detailed description of this environment variable, refer to the [Intel® C++ Compiler Classic Developer Guide and Reference](#). In the case of an Alder Lake processor with eight P-cores, either of the following two commands can be used for restricting threads to run only on the P-cores:

```
set KMP_HW_SUBSET=8c:intel_core
```

—or—

```
set KMP_HW_SUBSET=8c:eff1
```

Note that for higher performance, Intel® Hyper-Threading Technology on P-cores must be disabled. You can achieve this either by changing the BIOS setting or by using KMP_HW_SUBSET to specify P-cores and one-thread-per-core with the following command:

```
set KMP_HW_SUBSET=8c:intel_core,1t
```

—or—

```
set KMP_HW_SUBSET=8c:eff1,1t
```

If the user decides to adopt the more difficult approach of running on both P-cores and E-cores to maximize performance, there are a few aspects to take into consideration:

- Static versus dynamic load balancing
- Problem size
- Number of P-cores and E-cores
- OpenMP versus oneTBB

If there are similar or equal numbers of P-cores and E-cores and if both core types are used, using static load balancing for splitting the work items is likely to result in lower performance because E-cores will take longer to complete the work items assigned to them. For large GEMMs and {S,D}GETRF routines, oneMKL has implemented dynamic load balancing with OpenMP and will automatically select the best load balancing scheme. For most cases with small or regular problem sizes, static load balancing on P-cores is likely to give better performance. If the problem size is very large, the overhead of dynamic scheduling is small compared to overall computation time and dynamic load balancing will make more efficient use of P-cores and E-cores.

If the number of P-cores is much smaller than the number of E-cores, running on all cores may outperform limiting computations to only P-cores. Additional performance measurements would be needed to determine the best strategy.

As an alternative to OpenMP, users can also try [oneTBB](#), which might give better results for a given set of supported operations.

Improving Performance for Small Size Problems

The overhead of calling an Intel® oneAPI Math Kernel Library function for small problem sizes can be significant when the function has a large number of parameters or internally checks parameter errors. To reduce the performance overhead for these small size problems, the Intel® oneAPI Math Kernel Library *direct call* feature works in conjunction with the compiler to preprocess the calling parameters to supported Intel® oneAPI Math Kernel Library functions and directly call or inline special optimized small-matrix kernels that bypass error checking. For a list of functions supporting direct call, see [Limitations of the Direct Call](#).

To activate the feature, do the following:

- Compile your C or Fortran code with the preprocessor macro depending on whether a threaded or sequential mode of Intel® oneAPI Math Kernel Library is required by supplying the compiler option as explained below:

Intel® oneAPI Math Kernel Library Mode	Macro	Compiler Option
Threaded	MKL_DIRECT_CALL	/DMKL_DIRECT_CALL
Sequential	MKL_DIRECT_CALL_SEQ	/DMKL_DIRECT_CALL_SEQ

- For Fortran applications:
 - Enable preprocessor by using the `/fpp` option for Intel® Fortran Compiler.
 - Include the Intel® oneAPI Math Kernel Library Fortran include file `mkl_direct_call.fi`.

Intel® oneAPI Math Kernel Library skips error checking and intermediate function calls if the problem size is small enough (for example: a call to a function that supports direct call, such as `asdgemm`, with matrix ranks smaller than 50).

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Using MKL_DIRECT_CALL in C Applications

The following examples of code and link lines show how to activate direct calls to Intel® oneAPI Math Kernel Library kernels in C applications:

- Include the `mkl.h` header file:

```
#include "mkl.h"
int main(void) {

// Call Intel MKL DGEMM

return 0;
}
```

- For multi-threaded Intel® oneAPI Math Kernel Library, compile with `MKL_DIRECT_CALL` preprocessor macro:

```
icl /DMKL_DIRECT_CALL /Qstd=c99 your_application.c mkl_intel_lp64.lib mkl_core.lib
mkl_intel_thread.lib /Qopenmp -I%MKLROOT%/include
```

- To use Intel® oneAPI Math Kernel Library in the sequential mode, compile with `MKL_DIRECT_CALL_SEQ` preprocessor macro:

```
icl /DMKL_DIRECT_CALL_SEQ /Qstd=c99 your_application.c mkl_intel_lp64.lib mkl_core.lib
mkl_sequential.lib -I%MKLROOT%/include
```

Using MKL_DIRECT_CALL in Fortran Applications

The following examples of code and link lines show how to activate direct calls to Intel® oneAPI Math Kernel Library kernels in Fortran applications:

- Include `mkl_direct_call.fi`, to be preprocessed by the Fortran compiler preprocessor

```
#    include "mkl_direct_call.fi"
program    DGEMM_MAIN
....
*    Call Intel MKL DGEMM
```

```

....
    call sub1()
    stop 1
end

*      A subroutine that calls DGEMM
subroutine sub1
*      Call Intel MKL DGEMM

end

```

- For multi-threaded Intel® oneAPI Math Kernel Library, compile with `/fpp` option for Intel Fortran compiler and with `MKL_DIRECT_CALL` preprocessor macro:

```

ifort /DMKL_DIRECT_CALL /fpp your_application.f mkl_intel_lp64.lib mkl_core.lib
mkl_intel_thread.lib
/Qopenmp -I%MKLROOT%/include

```

- To use Intel® oneAPI Math Kernel Library in the sequential mode, compile with `/fpp` option for Intel Fortran compiler (or with `-Mpreprocess` for PGI compilers) and with `MKL_DIRECT_CALL_SEQ` preprocessor macro:

```

ifort /DMKL_DIRECT_CALL_SEQ /fpp your_application.f mkl_intel_lp64.lib mkl_core.lib
mkl_sequential.lib
-I%MKLROOT%/include

```

Limitations of the Direct Call

Directly calling the Intel® oneAPI Math Kernel Library kernels has the following limitations:

- If the `MKL_DIRECT_CALL` or `MKL_DIRECT_CALL_SEQ` macro is used, Intel® oneAPI Math Kernel Library may skip error checking.

Important

With a limited error checking, you are responsible for checking the correctness of function parameters to avoid unsafe and incorrect code execution.

- The feature is only available for the following functions:
 - BLAS: `?gemm`, `?gemm3m`, `?syrk`, `?trsm`, `?axpy`, and `?dot`
 - LAPACK: `?getrf`, `?getrs`, `?getri`, `?potrf`, and `?geqrf`. (available for C applications only)
- Intel® oneAPI Math Kernel Library Verbose mode, Conditional Numerical Reproducibility, and BLAS95 interfaces are not supported.
- GNU* Fortran compilers are not supported.
- For C applications, you must enable mixing declarations and user code by providing the `/Qstd=c99` option for Intel® compilers.

Other Tips and Techniques to Improve Performance

See Also

[Managing Performance of the Cluster Fourier Transform Functions](#)

Coding Techniques

This section discusses coding techniques to improve performance on processors based on supported architectures.

To improve performance, properly align arrays in your code. Additional conditions can improve performance for specific function domains.

Data Alignment and Leading Dimensions

To improve performance of your application that calls Intel® oneAPI Math Kernel Library, align your arrays on 64-byte boundaries and ensure that the leading dimensions of the arrays are divisible by $64/element_size$, where *element_size* is the number of bytes for the matrix elements (4 for single-precision real, 8 for double-precision real and single-precision complex, and 16 for double-precision complex). For more details, see [Example of Data Alignment](#).

For Intel® Xeon Phi™ processor x200 product family, codenamed Knights Landing, align your matrices on 4096-byte boundaries and set the leading dimension to the following integer expression:
 $((n * element_size + 511) / 512) * 512 + 64) / element_size$,
 where *n* is the matrix dimension along the leading dimension.

LAPACK Packed Routines

The routines with the names that contain the letters HP, OP, PP, SP, TP, UP in the matrix type and storage position (the second and third letters respectively) operate on the matrices in the packed format (see LAPACK "Routine Naming Conventions" sections in the Intel® oneAPI Math Kernel Library Developer Reference). Their functionality is strictly equivalent to the functionality of the unpacked routines with the names containing the letters HE, OR, PO, SY, TR, UN in the same positions, but the performance is significantly lower.

If the memory restriction is not too tight, use an unpacked routine for better performance. In this case, you need to allocate $N^2/2$ more memory than the memory required by a respective packed routine, where *N* is the problem size (the number of equations).

For example, to speed up solving a symmetric eigenproblem with an expert driver, use the unpacked routine:

```
call dsyevx(jobz, range, uplo, n, a, lda, vl, vu, il, iu, abstol, m, w, z, ldz, work, lwork,
iwork, ifail, info)
```

where *a* is the dimension *lda-by-n*, which is at least N^2 elements, instead of the packed routine:

```
call dspevx(jobz, range, uplo, n, ap, vl, vu, il, iu, abstol, m, w, z, ldz, work, iwork, ifail,
info)
```

where *ap* is the dimension $N*(N+1)/2$.

See Also

[Managing Performance of the Cluster Fourier Transform Functions](#)

Improving oneMKL Performance on Specific Processors

Dual-Core Intel® Xeon® Processor 5100 Series

To get the best performance with Intel® oneAPI Math Kernel Library on Dual-Core Intel® Xeon® processor 5100 series systems, enable the Hardware DPL (streaming data) Prefetcher functionality of this processor. To configure this functionality, use the appropriate BIOS settings, as described in your BIOS documentation.

Operating on Denormals

The IEEE 754-2008 standard, "An IEEE Standard for Binary Floating-Point Arithmetic", defines *denormal* (or *subnormal*) numbers as non-zero numbers smaller than the smallest possible normalized numbers for a specific floating-point format. Floating-point operations on denormals are slower than on normalized

operands because denormal operands and results are usually handled through a software assist mechanism rather than directly in hardware. This software processing causes Intel® oneAPI Math Kernel Library functions that consume denormals to run slower than with normalized floating-point numbers.

You can mitigate this performance issue by setting the appropriate bit fields in the MXCSR floating-point control register to flush denormals to zero (FTZ) or to replace any denormals loaded from memory with zero (DAZ). Check your compiler documentation to determine whether it has options to control FTZ and DAZ. Note that these compiler options may slightly affect accuracy.

Using Memory Functions

Avoiding Memory Leaks in oneMKL

When running, Intel® oneAPI Math Kernel Library (oneMKL) may allocate and deallocate internal buffers to facilitate better performance. Memory leaks can occur if the Intel® oneAPI Math Kernel Library is unloaded before freeing the internal buffers.

You can free the internal buffers by calling the `mkl_free_buffers()` function or, for more granular control, the `mkl_thread_free_buffers()` function.

Alternatively, setting the `MKL_DISABLE_FAST_MM` environment variable to 1 or calling the `mkl_disable_fast_mm()` function disables the internal memory manager. Be aware that this change may negatively impact the performance of some oneMKL functions, especially for small problem sizes.

See Also

[Intel Software Documentation Library](#)

Redefining Memory Functions

In C/C++ programs, you can replace Intel® oneAPI Math Kernel Library memory functions that the library uses by default with your own functions. To do this, use the *memory renaming* feature.

Memory Renaming

Intel® oneAPI Math Kernel Library memory management by default uses standard C run-time memory functions to allocate or free memory. These functions can be replaced using memory renaming.

Intel® oneAPI Math Kernel Library accesses the memory functions by pointers `i_malloc`, `i_free`, `i_calloc`, and `i_realloc`, which are visible at the application level. These pointers initially hold addresses of the standard C run-time memory functions `malloc`, `free`, `calloc`, and `realloc`, respectively. You can programmatically redefine values of these pointers to the addresses of your application's memory management functions.

Redirecting the pointers is the only correct way to use your own set of memory management functions. If you call your own memory functions without redirecting the pointers, the memory will get managed by two independent memory management packages, which may cause unexpected memory issues.

How to Redefine Memory Functions

To redefine memory functions, use the following procedure:

If you are using the statically linked Intel® oneAPI Math Kernel Library,

1. Include the `i_malloc.h` header file in your code.
This header file contains all declarations required for replacing the memory allocation functions. The header file also describes how memory allocation can be replaced in those Intel libraries that support this feature.

2. Redefine values of pointers `i_malloc`, `i_free`, `i_calloc`, and `i_realloc` prior to the first call to Intel® oneAPI Math Kernel Library functions, as shown in the following example:

```
#include "i_malloc.h"
. . .
i_malloc = my_malloc;
i_calloc = my_calloc;
i_realloc = my_realloc;
i_free   = my_free;
. . .
// Now you may call Intel MKL functions
```

If you are using the dynamically linked Intel® oneAPI Math Kernel Library,

1. Include the `i_malloc.h` header file in your code.
2. Redefine values of pointers `i_malloc_dll`, `i_free_dll`, `i_calloc_dll`, and `i_realloc_dll` prior to the first call to Intel® oneAPI Math Kernel Library functions, as shown in the following example:

```
#include "i_malloc.h"
. . .
i_malloc_dll = my_malloc;
i_calloc_dll = my_calloc;
i_realloc_dll = my_realloc;
i_free_dll   = my_free;
. . .
// Now you may call Intel MKL functions
```

Language-specific Usage Options

5

The provides broad support for Fortran and C/C++ programming. However, not all functions support both Fortran and C interfaces. For example, some LAPACK functions have no C interface. You can call such functions from C using mixed-language programming.

If you want to use LAPACK or BLAS functions that support Fortran 77 in the Fortran 95 environment, additional effort may be initially required to build compiler-specific interface libraries and modules from the source code provided with Intel® oneAPI Math Kernel Library.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[Language Interfaces Support, by Function Domain](#)

Using Language-Specific Interfaces with Intel® oneAPI Math Kernel Library

This section discusses mixed-language programming and the use of language-specific interfaces with Intel® oneAPI Math Kernel Library.

See also the "FFTW Interface to Intel® oneAPI Math Kernel Library" Appendix in the Intel® oneAPI Math Kernel Library Developer Reference for details of the FFTW interfaces to Intel® oneAPI Math Kernel Library.

Interface Libraries and Modules

You can create the following interface libraries and modules using the respective makefiles located in the interfaces directory.

File name	Contains
Libraries, in Intel® oneAPI Math Kernel Library architecture-specific directories	
<code>mk1_blas95.lib¹</code>	Fortran 95 wrappers for BLAS (BLAS95) for IA-32 architecture.
<code>mk1_blas95_ilp64.lib¹</code>	Fortran 95 wrappers for BLAS (BLAS95) supporting LP64 interface.
<code>mk1_blas95_lp64.lib¹</code>	Fortran 95 wrappers for BLAS (BLAS95) supporting ILP64 interface.
<code>mk1_lapack95.lib¹</code>	Fortran 95 wrappers for LAPACK (LAPACK95) for IA-32 architecture.
<code>mk1_lapack95_lp64.lib¹</code>	Fortran 95 wrappers for LAPACK (LAPACK95) supporting LP64 interface.
<code>mk1_lapack95_ilp64.lib¹</code>	Fortran 95 wrappers for LAPACK (LAPACK95) supporting ILP64 interface.

File name	Contains
<code>fftw2xc_intel.lib</code> ¹	Interfaces for FFTW version 2.x (C interface for Intel compilers) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw2xc_ms.lib</code>	Contains interfaces for FFTW version 2.x (C interface for Microsoft compilers) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw2xf_intel.lib</code>	Interfaces for FFTW version 2.x (Fortran interface for Intel compilers) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw3xc_intel.lib</code> ²	Interfaces for FFTW version 3.x (C interface for Intel compiler) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw3xc_ms.lib</code>	Interfaces for FFTW version 3.x (C interface for Microsoft compilers) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw3xf_intel.lib</code> ²	Interfaces for FFTW version 3.x (Fortran interface for Intel compilers) to call Intel® oneAPI Math Kernel Library FFT.
<code>fftw2x_cdft_SINGLE.lib</code>	Single-precision interfaces for MPI FFTW version 2.x (C interface) to call Intel® oneAPI Math Kernel Library cluster FFT.
<code>fftw2x_cdft_DOUBLE.lib</code>	Double-precision interfaces for MPI FFTW version 2.x (C interface) to call Intel® oneAPI Math Kernel Library cluster FFT.
<code>fftw3x_cdft.lib</code>	Interfaces for MPI FFTW version 3.x (C interface) to call Intel® oneAPI Math Kernel Library cluster FFT.
<code>fftw3x_cdft_ilp64.lib</code>	Interfaces for MPI FFTW version 3.x (C interface) to call Intel® oneAPI Math Kernel Library cluster FFT supporting the ILP64 interface.
Modules, in architecture- and interface-specific subdirectories of the Intel® oneAPI Math Kernel Library include directory	
<code>blas95.mod</code> ¹	Fortran 95 interface module for BLAS (BLAS95).
<code>lapack95.mod</code> ¹	Fortran 95 interface module for LAPACK (LAPACK95).
<code>f95_precision.mod</code> ¹	Fortran 95 definition of precision parameters for BLAS95 and LAPACK95.
<code>mkl_service.mod</code> ¹	Fortran 95 interface module for Intel® oneAPI Math Kernel Library support functions.

¹ Prebuilt for the Intel® Fortran compiler

² FFTW3 interfaces are integrated with Intel® oneAPI Math Kernel Library. Look into `<mkl directory>\interfaces\fftw3x*\makefile` for options defining how to build and where to place the standalone library with the wrappers.

See Also

[Fortran 95 Interfaces to LAPACK and BLAS](#)

Fortran 95 Interfaces to LAPACK and BLAS

Fortran 95 interfaces are compiler-dependent. Intel® oneAPI Math Kernel Library provides the interface libraries and modules precompiled with the Intel® Fortran compiler. Additionally, the Fortran 95 interfaces and wrappers are delivered as sources. (For more information, see [Compiler-dependent Functions and Fortran 90 Modules](#)). If you are using a different compiler, build the appropriate library and modules with your compiler and link the library as a user's library:

1. Go to the respective directory `<mkl_directory>\interfaces\blas95` or `<mkl_directory>\interfaces\lapack95`
2. Type:
 - For the IA 32 architecture, make `libia32 install_dir=<user_dir>`
 - `nmake libintel64 [interface=lp64|ilp64] install_dir=<user_dir>`

Important

The parameter `install_dir` is required.

As a result, the required library is built and installed in the `<user_dir>\lib` directory, and the `.mod` files are built and installed in the `<user_dir>\include\<arch>[\{lp64|ilp64\}]` directory, where `<arch>` is {ia32, intel64}.

By default, the ifort compiler is assumed. You may change the compiler with an additional parameter of `nmake`:

`FC=<compiler>.`

For example, the command

```
nmake libintel64 FC=f95 install_dir=<userf95_dir> interface=lp64
```

builds the required library and `.mod` files and installs them in subdirectories of `<userf95_dir>`.

To delete the library from the building directory, type:

- For the IA-32 architecture, make `cleania32 INSTALL_DIR=<user_dir>`
- `nmake cleanintel64 [interface=lp64|ilp64] install_dir=<user_dir>`
- `make clean INSTALL_DIR=<user_dir>`

Caution

Even if you have administrative rights, avoid setting `install_dir=..\..` or `install_dir=<mkl_directory>` in a build or clean command above because these settings replace or delete the Intel® oneAPI Math Kernel Library prebuilt Fortran 95 library and modules.

Compiler-dependent Functions and Fortran 90 Modules

Compiler-dependent functions occur whenever the compiler inserts into the object code function calls that are resolved in its run-time library (RTL). Linking of such code without the appropriate RTL will result in undefined symbols. Intel® oneAPI Math Kernel Library has been designed to minimize RTL dependencies.

In cases where RTL dependencies might arise, the functions are delivered as source code and you need to compile the code with whatever compiler you are using for your application.

In particular, Fortran 90 modules result in the compiler-specific code generation requiring RTL support. Therefore, Intel® oneAPI Math Kernel Library delivers these modules compiled with the Intel compiler, along with source code, to be used with different compilers.

Mixed-language Programming with the Intel Math Kernel Library

[Appendix A Intel® oneAPI Math Kernel Library Language Interfaces Support](#) lists the programming languages supported for each Intel® oneAPI Math Kernel Library function domain. However, you can call Intel® oneAPI Math Kernel Library routines from different language environments.

See also these Knowledge Base articles:

- <https://software.intel.com/content/www/us/en/develop/articles/performance-tools-for-software-developers-how-do-i-use-intel-mkl-with-java.html> for how to call Intel® oneAPI Math Kernel Library from Java* applications.
- <https://software.intel.com/content/www/us/en/develop/articles/how-to-use-boost-ublas-with-intel-mkl.html> for how to perform BLAS matrix-matrix multiplication in C++ using Intel® oneAPI Math Kernel Library substitution of Boost* uBLAS functions.
- <https://software.intel.com/content/www/us/en/develop/articles/intel-mkl-and-third-party-applications-how-to-use-them-together.html> for a list of articles describing how to use Intel® oneAPI Math Kernel Library with third-party libraries and applications.

Calling LAPACK, BLAS, and CBLAS Routines from C/C++ Language Environments

Not all Intel® oneAPI Math Kernel Library function domains support both C and Fortran environments. To use Intel® oneAPI Math Kernel Library Fortran-style functions in C/C++ environments, you should observe certain conventions, which are discussed for LAPACK and BLAS in the subsections below.

Caution

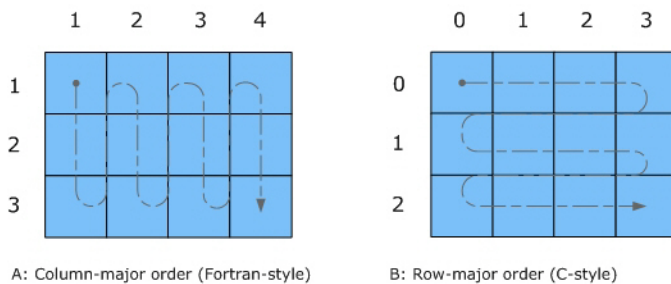
Avoid calling BLAS 95/LAPACK 95 from C/C++. Such calls require skills in manipulating the descriptor of a deferred-shape array, which is the Fortran 90 type. Moreover, BLAS95/LAPACK95 routines contain links to a Fortran RTL.

LAPACK and BLAS

Because LAPACK and BLAS routines are Fortran-style, when calling them from C-language programs, follow the Fortran-style calling conventions:

- Pass variables by *address*, not by *value*.
Function calls in [Example "Calling a Complex BLAS Level 1 Function from C++"](#) and [Example "Using CBLAS Interface Instead of Calling BLAS Directly from C"](#) illustrate this.
- Store your data in Fortran style, that is, column-major rather than row-major order.

With row-major order, adopted in C, the last array index changes most quickly and the first one changes most slowly when traversing the memory segment where the array is stored. With Fortran-style column-major order, the last index changes most slowly whereas the first index changes most quickly (as illustrated by the figure below for a two-dimensional array).



For example, if a two-dimensional matrix A of size $m \times n$ is stored densely in a one-dimensional array B, you can access a matrix element like this:

$A[i][j] = B[i*n+j]$ in C ($i=0, \dots, m-1, j=0, \dots, n-1$)

$A(i,j) = B((j-1)*m+i)$ in Fortran ($i=1, \dots, m, j=1, \dots, n$).

When calling LAPACK or BLAS routines from C, be aware that because the Fortran language is case-insensitive, the routine names can be both upper-case or lower-case, with or without the trailing underscore. For example, the following names are equivalent:

- LAPACK: `dgetrf`, `DGETRF`, `dgetrf_`, and `DGETRF_`
- BLAS: `dgemm`, `DGEMM`, `dgemm_`, and `DGEMM_`

See [Example "Calling a Complex BLAS Level 1 Function from C++"](#) on how to call BLAS routines from C.

See also the Intel® oneAPI Math Kernel Library Developer Reference for a description of the C interface to LAPACK functions.

CBLAS

Instead of calling BLAS routines from a C-language program, you can use the CBLAS interface.

CBLAS is a C-style interface to the BLAS routines. You can call CBLAS routines using regular C-style calls. Use the `mk1.h` header file with the CBLAS interface. `mk1.h` includes the `mk1_cblas.h` header file, which specifies enumerated values and prototypes of all the functions. It also determines whether the program is being compiled with a C++ compiler, and if it is, the included file will be correct for use with C++ compilation.

[Example "Using CBLAS Interface Instead of Calling BLAS Directly from C"](#) illustrates the use of the CBLAS interface.

C Interface to LAPACK

Instead of calling LAPACK routines from a C-language program, you can use the C interface to LAPACK provided by Intel® oneAPI Math Kernel Library.

The C interface to LAPACK is a C-style interface to the LAPACK routines. This interface supports matrices in row-major and column-major order, which you can define in the first function argument `matrix_order`. Use the `mk1.h` header file with the C interface to LAPACK. `mk1.h` includes the `mk1_lapacke.h` header file, which specifies constants and prototypes of all the functions. It also determines whether the program is being compiled with a C++ compiler, and if it is, the included file will be correct for use with C++ compilation. You can find examples of the C interface to LAPACK in the `examples\lapackesub` directory in the Intel® oneAPI Math Kernel Library installation directory.

Using Complex Types in C/C++

As described in the documentation for the Intel® Visual Fortran Compiler, C/C++ does not directly implement the Fortran types `COMPLEX(4)` and `COMPLEX(8)`. However, you can write equivalent structures. The type `COMPLEX(4)` consists of two 4-byte floating-point numbers. The first of them is the real-number component, and the second one is the imaginary-number component. The type `COMPLEX(8)` is similar to `COMPLEX(4)` except that it contains two 8-byte floating-point numbers.

Intel® oneAPI Math Kernel Library provides complex types `MKL_Complex8` and `MKL_Complex16`, which are structures equivalent to the Fortran complex types `COMPLEX(4)` and `COMPLEX(8)`, respectively. The `MKL_Complex8` and `MKL_Complex16` types are defined in the `mk1_types.h` header file. You can use these types to define complex data. You can also redefine the types with your own types before including the `mk1_types.h` header file. The only requirement is that the types must be compatible with the Fortran complex layout, that is, the complex type must be a pair of real numbers for the values of real and imaginary parts.

For example, you can use the following definitions in your C++ code:

```
#define MKL_Complex8 std::complex<float>
```

and

```
#define MKL_Complex16 std::complex<double>
```

See [Example "Calling a Complex BLAS Level 1 Function from C++"](#) for details. You can also define these types in the command line:

```
-DMKL_Complex8="std::complex<float>"
-MDKL_Complex16="std::complex<double>"
```

See Also

[Intel® Software Documentation Library](#) for the Intel®Visual Fortran Compiler documentation for the Intel®Visual Fortran Compiler documentation

Calling BLAS Functions that Return the Complex Values in C/C++ Code

Complex values that functions return are handled differently in C and Fortran. Because BLAS is Fortran-style, you need to be careful when handling a call from C to a BLAS function that returns complex values. However, in addition to normal function calls, Fortran enables calling functions as though they were subroutines, which provides a mechanism for returning the complex value correctly when the function is called from a C program. When a Fortran function is called as a subroutine, the return value is the first parameter in the calling sequence. You can use this feature to call a BLAS function from C.

The following example shows how a call to a Fortran function as a subroutine converts to a call from C and the hidden parameter result gets exposed:

Normal Fortran function call: `result = cdotc( n, x, 1, y, 1 )`

A call to the function as a subroutine: `call cdotc( result, n, x, 1, y, 1)`

A call to the function from C: `cdotc( &result, &n, x, &one, y, &one )`

NOTE

Intel® oneAPI Math Kernel Library has both upper-case and lower-case entry points in the Fortran-style (case-insensitive) BLAS, with or without the trailing underscore. So, all these names are equivalent and acceptable: `cdotc`, `CDOTC`, `cdotc_`, and `CDOTC_`.

The above example shows one of the ways to call several level 1 BLAS functions that return complex values from your C and C++ applications. An easier way is to use the CBLAS interface. For instance, you can call the same function using the CBLAS interface as follows:

```
cblas_cdotc( n, x, 1, y, 1, &result )
```

NOTE

The complex value comes last on the argument list in this case.

The following examples show use of the Fortran-style BLAS interface from C and C++, as well as the CBLAS (C language) interface:

- [Example "Calling a Complex BLAS Level 1 Function from C"](#)
- [Example "Calling a Complex BLAS Level 1 Function from C++"](#)
- [Example "Using CBLAS Interface Instead of Calling BLAS Directly from C"](#)

Example "Calling a Complex BLAS Level 1 Function from C"

The example below illustrates a call from a C program to the complex BLAS Level 1 function `zdotc()`. This function computes the dot product of two double-precision complex vectors.

In this example, the complex dot product is returned in the structure `c`.

```
#include "mkl.h"
#define N 5
int main()
{
    int n = N, inca = 1, incb = 1, i;
    MKL_Complex16 a[N], b[N], c;
    for( i = 0; i < n; i++ )
    {
        a[i].real = (double)i; a[i].imag = (double)i * 2.0;
        b[i].real = (double)(n - i); b[i].imag = (double)i * 2.0;
    }
    zdotc( &c, &n, a, &inca, b, &incb );
    printf( "The complex dot product is: ( %6.2f, %6.2f)\n", c.real, c.imag );
    return 0;
}
```

In this example, the complex dot product for large data size is returned in the structure `c`.

```
#include "mkl.h"
#define N 5
int main()
{
    MKL_INT64 n = N, inca = 1, incb = 1, i;
    MKL_Complex16 a[N], b[N], c;
    for( i = 0; i < n; i++ )
    {
        a[i].real = (double)i; a[i].imag = (double)i * 2.0;
        b[i].real = (double)(n - i); b[i].imag = (double)i * 2.0;
    }
    zdotc_64( &c, &n, a, &inca, b, &incb );
    printf( "The complex dot product is: ( %6.2f, %6.2f)\n", c.real, c.imag );
    return 0;
}
```

Example "Calling a Complex BLAS Level 1 Function from C++"

Below is the C++ implementation:

```
#include <complex>
#include <iostream>
#define MKL_Complex16 std::complex<double>
#include "mkl.h"

#define N 5

int main()
{
    int n, inca = 1, incb = 1, i;
    std::complex<double> a[N], b[N], c;
    n = N;

    for( i = 0; i < n; i++ )
    {
        a[i] = std::complex<double>(i, i*2.0);
        b[i] = std::complex<double>(n-i, i*2.0);
    }
}
```

```
    zdotc(&c, &n, a, &inca, b, &incb );
    std::cout << "The complex dot product is: " << c << std::endl;
    return 0;
}
```

Example "Using CBLAS Interface Instead of Calling BLAS Directly from C"

This example uses CBLAS:

```
#include <stdio.h>
#include "mkl.h"
typedef struct{ double re; double im; } complex16;
#define N 5
int main()
{
    int n, inca = 1, incb = 1, i;
    complex16 a[N], b[N], c;
    n = N;
    for( i = 0; i < n; i++ )
    {
        a[i].re = (double)i; a[i].im = (double)i * 2.0;
        b[i].re = (double)(n - i); b[i].im = (double)i * 2.0;
    }
    cblas_zdotc_sub(n, a, inca, b, incb, &c );
    printf( "The complex dot product is: ( %.2f, %.2f)\n", c.re, c.im );
    return 0;
}
```

Obtaining Numerically Reproducible Results

offers functions and environment variables that help you obtain Conditional Numerical Reproducibility (CNR) of floating-point results when calling the library functions from your application. These new controls enable Intel® oneAPI Math Kernel Library to run in a special mode, when functions return bitwise reproducible floating-point results from run to run under the following conditions:

- Calls to Intel® oneAPI Math Kernel Library occur in a single executable
- The number of computational threads used by the library does not change in the run

For a limited set of routines, you can eliminate the second condition by using Intel® oneAPI Math Kernel Library in [strict CNR mode](#).

It is well known that for general single and double precision IEEE floating-point numbers, the associative property does not always hold, meaning $(a+b)+c$ may not equal $a+(b+c)$. Let's consider a specific example. In infinite precision arithmetic $2^{-63} + 1 + -1 = 2^{-63}$. If this same computation is done on a computer using double precision floating-point numbers, a rounding error is introduced, and the order of operations becomes important:

$$(2^{-63} + 1) + (-1) \simeq 1 + (-1) = 0$$

versus

$$2^{-63} + (1 + (-1)) \simeq 2^{-63} + 0 = 2^{-63}$$

This inconsistency in results due to order of operations is precisely what the new functionality addresses.

The application related factors that affect the order of floating-point operations within a single executable program include selection of a code path based on run-time processor dispatching, alignment of data arrays, variation in number of threads, threaded algorithms and internal floating-point control settings. You can control most of these factors by controlling the number of threads and floating-point settings and by taking steps to align memory when it is allocated (see the [Getting Reproducible Results with Intel® MKL](#) knowledge base article for details). However, run-time dispatching and certain threaded algorithms do not allow users to make changes that can ensure the same order of operations from run to run.

Intel® oneAPI Math Kernel Library does run-time processor dispatching in order to identify the appropriate internal code paths to traverse for the Intel® oneAPI Math Kernel Library functions called by the application. The code paths chosen may differ across a wide range of Intel processors and Intel architecture compatible processors and may provide differing levels of performance. For example, an Intel® oneAPI Math Kernel Library function running on an Intel® Pentium® 4 processor may run one code path, while on the latest Intel® Xeon® processor it will run another code path. This happens because each unique code path has been optimized to match the features available on the underlying processor. One key way that the new features of a processor are exposed to the programmer is through the instruction set architecture (ISA). Because of this, code branches in Intel® oneAPI Math Kernel Library are designated by the latest ISA they use for optimizations: from the Intel® Streaming SIMD Extensions 2 (Intel® SSE2) to the Intel® Advanced Vector Extensions2 (Intel® AVX2). The feature-based approach introduces a challenge: if any of the internal floating-point operations are done in a different order or are re-associated, the computed results may differ.

Dispatching optimized code paths based on the capabilities of the processor on which the code is running is central to the optimization approach used by Intel® oneAPI Math Kernel Library. So it is natural that consistent results require some performance trade-offs. If limited to a particular code path, performance of Intel® oneAPI Math Kernel Library can in some circumstances degrade by more than a half. To understand this, note that matrix-multiply performance nearly doubled with the introduction of new processors supporting Intel AVX2 instructions. Even if the code branch is not restricted, performance can degrade by 10-20% because the new functionality restricts algorithms to maintain the order of operations.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Product and Performance Information

Notice revision #20201201

Getting Started with Conditional Numerical Reproducibility

Intel® oneAPI Math Kernel Library offers functions and environment variables to help you get reproducible results. You can configure Intel® oneAPI Math Kernel Library using functions or environment variables, but the functions provide more flexibility.

The following specific examples introduce you to the conditional numerical reproducibility.

While these examples recommend aligning input and output data, you can supply unaligned data to Intel® oneAPI Math Kernel Library functions running in the CNR mode, but refer to [Reproducibility Conditions](#) for details related to data alignment.

Intel CPUs supporting Intel AVX2

To ensure Intel® oneAPI Math Kernel Library calls return the same results on every Intel CPU supporting Intel AVX2 instructions:

1. Make sure that your application uses a fixed number of threads
2. (Recommended) Properly align input and output arrays in Intel® oneAPI Math Kernel Library function calls
3. Do either of the following:

- Call

```
mkl_cbwr_set(MKL_CBWR_AVX2)
```

- Set the environment variable:

```
set MKL_CBWR = AVX2
```

NOTE

On non-Intel CPUs and on Intel CPUs that do not support Intel AVX2, this environment setting may cause results to differ because the `AUTO` branch is used instead, while the above function call returns an error and does not enable the CNR mode.

Intel CPUs supporting Intel SSE2

To ensure Intel® oneAPI Math Kernel Library calls return the same results on every Intel CPU supporting Intel SSE2 instructions:

1. Make sure that your application uses a fixed number of threads
2. (Recommended) Properly align input and output arrays in Intel® oneAPI Math Kernel Library function calls
3. Do either of the following:

- Call

```
mkl_cbwr_set(MKL_CBWR_SSE2)
```

- Set the environment variable:

```
set MKL_CBWR = SSE2
```

NOTE

On non-Intel CPUs, this environment setting may cause results to differ because the `AUTO` branch is used instead, while the above function call returns an error and does not enable the CNR mode.

Intel or Intel compatible CPUs supporting Intel SSE2

On non-Intel CPUs, only the `MKL_CBWR_AUTO` and `MKL_CBWR_COMPATIBLE` options are supported for function calls and only `AUTO` and `COMPATIBLE` options for environment settings.

To ensure Intel® oneAPI Math Kernel Library calls return the same results on all Intel or Intel compatible CPUs supporting Intel SSE2 instructions:

1. Make sure that your application uses a fixed number of threads
2. (Recommended) Properly align input and output arrays in Intel® oneAPI Math Kernel Library function calls
3. Do either of the following:

- Call

```
mkl_cbwr_set(MKL_CBWR_COMPATIBLE)
```

- Set the environment variable:

```
set MKL_CBWR = COMPATIBLE
```

NOTE

The special `MKL_CBWR_COMPATIBLE/COMPATIBLE` option is provided because Intel and Intel compatible CPUs have a few instructions, such as approximation instructions `rcpps/rsqrtps`, that may return different results. This option ensures that Intel® oneAPI Math Kernel Library does not use these instructions and forces a single Intel SSE2 only code path to be executed.

Next steps

See [Specifying the Code Branches](#)

for details of specifying the branch using environment variables.

See the following sections in the *Intel® oneAPI Math Kernel Library Developer Reference*:

Support Functions for Conditional Numerical Reproducibility

for how to configure the CNR mode of Intel® oneAPI Math Kernel Library using functions.

Intel® oneAPI Math Kernel Library PARDISO - Parallel Direct Sparse Solver Interface

for how to configure the CNR mode for PARDISO.

See Also

[Code Examples](#)

Specifying Code Branches

Intel® oneAPI Math Kernel Library provides a conditional numerical reproducibility (CNR) functionality that enables you to obtain reproducible results from oneMKL routines. When enabling CNR, you choose a specific code branch of Intel® oneAPI Math Kernel Library that corresponds to the instruction set architecture (ISA) that you target. You can specify the code branch and other CNR options using the `MKL_CBWR` environment variable.

- `MKL_CBWR=<branch>[,STRICT]"` or
- `MKL_CBWR="BRANCH=<branch>[,STRICT]"`

Use the `STRICT` flag to enable strict CNR mode. For more information, see [Reproducibility Conditions](#).

The `<branch>` placeholder specifies the CNR branch with one of the following values:

Value	Description
AUTO	CNR mode uses the standard ISA-based dispatching model while ensuring fixed cache sizes, deterministic reductions, and static scheduling
COMPATIBLE	Intel® Streaming SIMD Extensions 2 (Intel® SSE2) without <code>rcpps/</code> <code>rsqrtps</code> instructions
SSE2	Intel SSE2
SSE3	DEPRECATED. Intel® Streaming SIMD Extensions 3 (Intel® SSE3). This setting is kept for backward compatibility and is equivalent to <code>SSE2</code> .
SSSE3	Supplemental Streaming SIMD Extensions 3 (SSSE3)
SSE4_2	Intel® Streaming SIMD Extensions 4.2 (Intel® SSE4.2)
AVX	Intel® Advanced Vector Extensions (Intel® AVX)
AVX2	Intel® Advanced Vector Extensions 2 (Intel® AVX2)
AVX512	Intel AVX-512 on Intel® Xeon® processors
AVX512_E1	Intel® Advanced Vector Extensions 512 (Intel® AVX-512) with support for Vector Neural Network Instructions
AVX512_MIC	DEPRECATED. Intel® Advanced Vector Extensions 512 (Intel® AVX-512) on Intel® Xeon Phi™ processors. This setting is kept for backward compatibility and is equivalent to <code>AVX2</code> .
AVX512_MIC_E1	DEPRECATED. Intel® Advanced Vector Extensions 512 (Intel® AVX-512) with support for Vector Neural Network Instructions on Intel® Xeon Phi™ processors. This setting is kept for backward compatibility and is equivalent to <code>AVX2</code> .

When specifying the CNR branch, be aware of the following:

- Reproducible results are provided under [Reproducibility Conditions](#).
- Settings other than `AUTO` or `COMPATIBLE` are available only for Intel processors.
- To get the CNR branch optimized for the processor where your program is currently running, choose the value of `AUTO` or call the `mkl_cbwr_get_auto_branch` function.
- Strict CNR mode is supported only for `AVX2`, `AVX512`, `AVX512_E1`, `AVX512_MIC`, and `AVX512_MIC_E1` branches. You can also use strict CNR mode with the `AUTO` branch when running on Intel processors that support one of these instruction set architectures (ISAs).

Setting the `MKL_CBWR` environment variable or a call to an equivalent `mkl_cbwr_set` function fixes the code branch and sets the reproducibility mode.

NOTE

- If the value of the branch is incorrect or your processor or operating system does not support the specified ISA, CNR ignores this value and uses the `AUTO` branch without providing any warning messages.
 - Calls to functions that define the behavior of CNR must precede any of the math library functions that they control.
 - Settings specified by the functions take precedence over the settings specified by the environment variable.
-

See the *Intel® oneAPI Math Kernel Library Developer Reference* for how to specify the branches using functions.

See Also

[Getting Started with Conditional Numerical Reproducibility](#)

Reproducibility Conditions

To get reproducible results from run to run, ensure that the number of threads is fixed and constant. Specifically:

- If you are running your program with OpenMP* parallelization on different processors, explicitly specify the number of threads.
- To ensure that your application has deterministic behavior with OpenMP* parallelization and does not adjust the number of threads dynamically at run time, set `MKL_DYNAMIC` and `OMP_DYNAMIC` to `FALSE`. This is especially needed if you are running your program on different systems.
- If you are running your program with the Intel® Threading Building Blocks parallelization, numerical reproducibility is not guaranteed.

Strict CNR Mode

In strict CNR mode, oneAPI Math Kernel Library provides bitwise reproducible results for a limited set of functions and code branches even when the number of threads changes. These routines and branches support strict CNR mode (64-bit libraries only):

- `?gemm`, `?symm`, `?hemm`, `?trsm` and their CBLAS equivalents (`cblas_?gemm`, `cblas_?symm`, `cblas_?hemm`, and `cblas_?trsm`).
- Intel® Advanced Vector Extensions 2 (Intel® AVX2) or Intel® Advanced Vector Extensions 512 (Intel® AVX-512).

When using other routines or CNR branches, oneAPI Math Kernel Library operates in standard (non-strict) CNR mode, subject to the restrictions described above. Enabling strict CNR mode can reduce performance.

NOTE

- As usual, you should align your data, even in CNR mode, to obtain the best possible performance. While CNR mode also fully supports unaligned input and output data, the use of it might reduce the performance of some oneAPI Math Kernel Library functions on earlier Intel processors. Refer to coding techniques that improve performance for more details.
 - Conditional Numerical Reproducibility does not ensure that bitwise-identical NaN values are generated when the input data contains NaN values.
 - If dynamic memory allocation fails on one run but succeeds on another run, you may fail to get reproducible results between these two runs.
-

See Also

[MKL_DYNAMIC](#)

Coding Techniques

Setting the Environment Variable for Conditional Numerical Reproducibility

The following examples illustrate the use of the `MKL_CBWR` environment variable. The first command sets Intel® oneAPI Math Kernel Library to run in the CNR mode based on the default dispatching for your platform. The other two commands are equivalent and set the CNR branch to Intel AVX:

- `set MKL_CBWR=AUTO`
- `set MKL_CBWR=AVX`
- `set MKL_CBWR=BRANCH=AVX`

See Also

Specifying Code Branches

Code Examples

The following simple programs show how to obtain reproducible results from run to run of Intel® oneAPI Math Kernel Library functions. See the *Intel® oneAPI Math Kernel Library Developer Reference* for more examples.

C Example of CNR

```
#include <mkl.h>
int main(void) {
    int my_cbwr_branch;
    /* Align all input/output data on 64-byte boundaries */
    /* "for best performance of Intel® oneAPI Math Kernel Library */
    void *darray;
    int darray_size=1000;
    /* Set alignment value in bytes */
    int alignment=64;
    /* Allocate aligned array */
    darray = mkl_malloc (sizeof(double)*darray_size, alignment);
    /* Find the available MKL_CBWR_BRANCH automatically */
    my_cbwr_branch = mkl_cbwr_get_auto_branch();
    /* User code without oneMKL calls */
    /* Piece of the code where CNR of oneMKL is needed */
    /* The performance of oneMKL functions might be reduced for CNR mode */
    /* If the "IF" statement below is commented out, Intel® oneAPI Math Kernel Library will run in a
    regular mode, */
    /* and data alignment will allow you to get best performance */
    if (mkl_cbwr_set(my_cbwr_branch)) {
        printf("Error in setting MKL_CBWR_BRANCH! Aborting...\n");
        return;
    }
    /* CNR calls to oneMKL + any other code */
    /* Free the allocated aligned array */
    mkl_free(darray);
}
```

Fortran Example of CNR

```
PROGRAM MAIN
  INCLUDE 'mkl.fi'
  INTEGER*4 MY_CBWR_BRANCH
  ! Align all input/output data on 64-byte boundaries
```

```

! "for best performance of Intel® oneAPI Math Kernel Library
! Declare oneMKL memory allocation routine
#ifdef _IA32
    INTEGER MKL_MALLOC
#else
    INTEGER*8 MKL_MALLOC
#endif
EXTERNAL MKL_MALLOC, MKL_FREE
DOUBLE PRECISION DARRAY
POINTER (P_DARRAY,DARRAY(1))
INTEGER DARRAY_SIZE
PARAMETER (DARRAY_SIZE=1000)
! Set alignment value in bytes
INTEGER ALIGNMENT
PARAMETER (ALIGNMENT=64)
! Allocate aligned array
P_DARRAY = MKL_MALLOC (%VAL(8*DARRAY_SIZE), %VAL(ALIGNMENT));
! Find the available MKL_CBWR_BRANCH automatically
MY_CBWR_BRANCH = MKL_CBWR_GET_AUTO_BRANCH()
! User code without oneMKL calls
! Piece of the code where CNR of oneMKL is needed
! The performance of oneMKL functions may be reduced for CNR mode
! If the "IF" statement below is commented out, Intel® oneAPI Math Kernel Library will run in a
regular mode,
! and data alignment will allow you to get best performance
IF (MKL_CBWR_SET (MY_CBWR_BRANCH) .NE. MKL_CBWR_SUCCESS) THEN
    PRINT *, 'Error in setting MKL_CBWR_BRANCH! Aborting...'
    RETURN
ENDIF
! CNR calls to oneMKL + any other code
! Free the allocated aligned array
CALL MKL_FREE(P_DARRAY)

END

```

Use of CNR with Unaligned Data in C

```

#include <mkl.h>
int main(void) {
    int my_cbwr_branch;
    /* If it is not possible to align all input/output data on 64-byte boundaries */
    /* to achieve performance, use unaligned IO data with possible performance */
    /* penalty */
    /* Using unaligned IO data */
    double *darray;
    int darray_size=1000;
    /* Allocate array, malloc aligns data on 8/16-byte boundary only */
    darray = (double *)malloc (sizeof(double)*darray_size);
    /* Find the available MKL_CBWR_BRANCH automatically */
    my_cbwr_branch = mkl_cbwr_get_auto_branch();
    /* User code without oneMKL calls */
    /* Piece of the code where CNR of oneMKL is needed */
    /* The performance of oneMKL functions might be reduced for CNR mode */
    /* If the "IF" statement below is commented out, oneMKL will run in a regular mode, */
    /* and you will NOT get best performance without data alignment */
    if (mkl_cbwr_set(my_cbwr_branch)) {
        printf("Error in setting MKL_CBWR_BRANCH! Aborting...\n");
        return;
    }
}

```

```

}
/* CNR calls to oneMKL + any other code */
/* Free the allocated array */
free(darray);

```

Use of CNR with Unaligned Data in Fortran

```

PROGRAM MAIN
INCLUDE 'mkl.fi'
INTEGER*4 MY_CBWR_BRANCH
! If it is not possible to align all input/output data on 64-byte boundaries
! to achieve performance, use unaligned IO data with possible performance
! penalty
DOUBLE PRECISION, DIMENSION(:), ALLOCATABLE :: DARRAY
INTEGER DARRAY_SIZE, STATUS
PARAMETER (DARRAY_SIZE=1000)
! Allocate array with undefined alignment
ALLOCATE(DARRAY(DARRAY_SIZE));
! Find the available MKL_CBWR_BRANCH automatically
MY_CBWR_BRANCH = MKL_CBWR_GET_AUTO_BRANCH()
! User code without oneMKL calls
! Piece of the code where CNR of oneMKL is needed
! The performance of oneMKL functions might be reduced for CNR mode
! If the "IF" statement below is commented out, oneMKL will run in a regular mode,
! and you will NOT get best performance without data alignment
IF (MKL_CBWR_SET(MY_CBWR_BRANCH) .NE. MKL_CBWR_SUCCESS) THEN
    PRINT *, 'Error in setting MKL_CBWR_BRANCH! Aborting...'
    RETURN
ENDIF
! CNR calls to oneMKL + any other code
! Free the allocated array
DEALLOCATE(DARRAY)
END

```

Coding Tips

This section provides coding tips for managing data alignment and version-specific compilation.

See Also

[Mixed-language Programming with the Intel® oneAPI Math Kernel Library](#) Tips on language-specific programming

[Managing Performance and Memory](#) Coding tips related to performance improvement and use of memory functions

[Obtaining Numerically Reproducible Results](#) Tips for obtaining numerically reproducible results of computations

Example of Data Alignment

Needs for best performance with Intel® oneAPI Math Kernel Library or for reproducible results from run to run of Intel® oneAPI Math Kernel Library functions require alignment of data arrays. The following example shows how to align an array on 64-byte boundaries. To do this, use `mkl_malloc()` in place of system provided memory allocators, as shown in the code example below.

Aligning Addresses on 64-byte Boundaries

```
// ***** C language *****
...
#include <stdlib.h>
#include <mkl.h>
...
void *darray;
int workspace;
// Set value of alignment
int alignment=64;
...
// Allocate aligned workspace
darray = mkl_malloc( sizeof(double)*workspace, alignment );
...
// call the program using oneMKL
mkl_app( darray );
...
// Free workspace
mkl_free( darray );
```

```
! ***** Fortran language *****
...
! Set value of alignment
integer alignment
parameter (alignment=64)
...
! Declare oneMKL routines
#ifdef _IA32
integer mkl_malloc
#else
```

```

integer*8 mkl_malloc
#endif
external mkl_malloc, mkl_free, mkl_app
...
double precision darray
pointer (p_wrk,darray(1))
integer workspace
...
! Allocate aligned workspace
p_wrk = mkl_malloc( %val(8*workspace), %val(alignment) )
...
! call the program using oneMKL
call mkl_app( darray )
...
! Free workspace
call mkl_free(p_wrk)

```

Using Predefined Preprocessor Symbols for Intel® MKL Version-Dependent Compilation

Preprocessor symbols (macros) substitute values in a program before it is compiled. The substitution is performed in the preprocessing phase.

The following preprocessor symbols are available:

Predefined Preprocessor Symbol	Description
<code>__INTEL_MKL__</code>	Intel® oneAPI Math Kernel Library major version
<code>__INTEL_MKL_MINOR__</code>	Intel® oneAPI Math Kernel Library minor version
<code>__INTEL_MKL_UPDATE__</code>	Intel® oneAPI Math Kernel Library update number
<code>INTEL_MKL_VERSION</code>	Intel® oneAPI Math Kernel Library full version in the following format: $INTEL_MKL_VERSION = (_INTEL_MKL_ _ * 100 + _INTEL_MKL_MINOR_ _) * 100 + _INTEL_MKL_UPDATE_ _$

These symbols enable conditional compilation of code that uses new features introduced in a particular version of the library.

To perform conditional compilation:

1. Depending on your compiler, include in your code the file where the macros are defined:

C/C++ compiler: `mkl_version.h`,
or `mkl.h`, which includes `mkl_version.h`

Intel®Fortran compiler: `mkl.fi`

Any Fortran compiler with enabled preprocessing: `mkl_version.h`

Read the documentation for your compiler for the option that enables preprocessing.

2. [Optionally] Use the following preprocessor directives to check whether the macro is defined:

- `#ifdef`, `#endif` for C/C++
- `!DEC$IF DEFINED`, `!DEC$ENDIF` for Fortran

3. Use preprocessor directives for conditional inclusion of code:

- #if, #endif for C/C++
- !DEC\$IF, !DEC\$ENDIF for Fortran

Example

This example shows how to compile a code segment conditionally for a specific version of Intel® oneAPI Math Kernel Library. In this case, the version is 11.2 Update 4:

Intel®Fortran Compiler:

```
include "mkl.fi"
!DEC$IF DEFINED INTEL_MKL_VERSION
!DEC$IF INTEL_MKL_VERSION .EQ. 110204
*   Code to be conditionally compiled
!DEC$ENDIF
!DEC$ENDIF
```

C/C++ Compiler. Fortran Compiler with Enabled Preprocessing:

```
#include "mkl.h"
#ifdef INTEL_MKL_VERSION
#if INTEL_MKL_VERSION == 110204
...   Code to be conditionally compiled
#endif
#endif
```

8

Managing Output

Using oneMKL Verbose Mode

When building applications that call Intel® oneAPI Math Kernel Library functions, it may be useful to determine:

- which computational functions are called
- what parameters are passed to them
- how much time is spent to execute the functions
- (for GPU applications) which GPU device the kernel is executed on

You can get an application to print this information to a standard output device by enabling Intel® oneAPI Math Kernel Library Verbose. Functions that can print this information are referred to as *verbose-enabled functions*.

When Verbose mode is active in an Intel® oneAPI Math Kernel Library domain, every call of a verbose-enabled function finishes with printing a human-readable line describing the call. However, if your application gets terminated for some reason during the function call, no information for that function will be printed. The first call to a verbose-enabled function also prints a version information line.

For GPU applications, additional information (one or more GPU information lines) will also be printed by the first call to a verbose-enabled function, following the version information line which will be printed for the host CPU. If there is more than one GPU detected, each GPU device will be printed in a separate line.

We have different implementations for verbose with CPU applications and verbose with GPU applications. The Intel® MKL Verbose mode has 2 modes when used with CPU applications: disabled (default) and enabled. The Intel® MKL Verbose mode has three modes when used with GPU applications: disabled (default), enabled without timing, and enabled with synchronous timing.

To change the verbose mode, either set the environment variable `MKL_VERBOSE`:

	CPU application	GPU application
Set <code>MKL_VERBOSE</code> to 0	to disable Verbose	to disable Verbose
Set <code>MKL_VERBOSE</code> to 1	to enable Verbose	to enable Verbose without timing
Set <code>MKL_VERBOSE</code> to 2	to enable Verbose	to enable Verbose with synchronous timing

or call the support function `mkl_verbose(int mode)`:

	CPU application	GPU application
Call <code>mkl_verbose(0)</code>	to disable Verbose	to disable Verbose
Call <code>mkl_verbose(1)</code>	to enable Verbose	to enable Verbose without timing
Call <code>mkl_verbose(2)</code>	to enable Verbose	to enable Verbose with synchronous timing

Verbose with CPU Applications

Verbose output will be consisted of version information line and call description lines for CPU.

For CPU applications, you can enable Intel® oneAPI Math Kernel Library Verbose mode in these domains:

- BLAS (and BLAS-like extensions)
- LAPACK
- ScaLAPACK (selected functionality)
- FFT

Verbose with GPU Applications

The verbose feature is enabled for GPU applications that uses DPC++ API or C/Fortran API with OpenMP offload. When used with GPU applications, verbose allows the measurement of execution time to be enabled or disabled with verbose mode. Timing is taken synchronously, so if verbose is enabled with timing, kernel executions will become synchronous (previous kernel will block later kernels)

Verbose output will be consisted of version information line, GPU information lines, and call description lines for GPU.

NOTE Timing for GPU applications is reported for overall execution. For selected functionality device execution time can be also reported if the input queue was created with [profiling information](#).

For GPU applications, you can enable Intel® oneAPI Math Kernel Library Verbose mode in these domains:

- BLAS (and BLAS-like extensions)
- LAPACK
- FFT

For Both CPU and GPU Verbose

Both enabling and disabling of the Verbose mode using the function call takes precedence over the environment setting. For a full description of the `mkl_verbose` function, see either the Intel® oneAPI Math Kernel Library Developer Reference for C or the Intel® oneAPI Math Kernel Library Developer Reference for Fortran. Both references are available in the Intel® Software Documentation Library.

Intel® oneAPI Math Kernel Library Verbose mode is not a thread-local but a global state. In other words, if an application changes the mode from multiple threads, the result is undefined.

WARNING

The performance of an application may degrade with the Verbose mode enabled, especially when the number of calls to verbose-enabled functions is large, because every call to a verbose-enabled function requires an output operation.

See Also

[Intel Software Documentation Library](#)

Version Information Line

In the Intel® oneAPI Math Kernel Library Verbose mode, the first call to a verbose-enabled function prints a version information line. The line begins with the `thmKL_VERBOSE` character string and uses spaces as delimiters. The format of the rest of the line may change in a future release.

The following table lists information contained in a version information line and provides available links for more information:

Information	Description	Related Links
Intel® oneAPI Math Kernel Library version.	This information is separated by a comma from the rest of the line.	
Operating system.	Possible values: • <code>LnX</code> for Linux* OS • <code>Win</code> for Windows* OS • <code>OSX</code> for macOS*	

Information	Description	Related Links
The host CPU frequency.		
Intel® oneAPI Math Kernel Library interface layer used by the application.	Possible values: - cdecl - lp64 or ilp64 on systems based on the Intel® 64 architecture.	Using the ILP64 Interface vs. LP64 Interface
Intel® oneAPI Math Kernel Library threading layer used by the application.	Possible values: intel_thread, tbb_thread, or sequential.	Linking with Threading Libraries

The following is an example of a version information line:

```
MKL_VERBOSE Intel(R) MKL 11.2 Beta build 20131126 for Intel(R) 64 architecture Intel(R)
Advanced Vector Extensions (Intel(R) AVX) Enabled Processor, Win 3.10GHz lp64
intel_thread
```

GPU Information Line

In Intel® oneAPI Math Kernel Library Verbose mode for GPU applications, the first call to a verbose-enabled function prints out the GPU information line or lines for all detected GPU devices, each in a separate line. The line begins with the `MKL_VERBOSE Detected` character string and uses spaces as delimiters. The format of the rest of the line may change in a future release.

Only Intel® GPU is supported.

The following table lists information contained in a GPU information line.

See Also [Call Description Line for GPU](#)

Information	Description
GPU index	The index of the GPU device will be printed after the character string "GPU" (e.g. GPU0, GPU1, GPU2, etc). This GPU index will be used as a nickname of the device in call description lines to refer to the device.
Intel® GPU architecture	The value can be one of the following: - Intel(R) Gen9 - Intel(R) Xe_LP - Intel(R) Xe_HP - Intel(R) Xe_HPG - Intel(R) Xe_HPC - Unknown GPU
Runtime backend	The value printed is prefixed with Backend:
Vector Engine number	The value printed is prefixed with VE:
Stack number	The value printed is prefixed with Stack:
Maximum workgroup size	The value printed is prefixed with maxWGsize:

The following is an example of a GPU information line:

```
MKL_VERBOSE Detected GPU0 Intel(R)_Gen9 Backend:OpenCL VE:72 Stack:1 maxWGsize:256
```

Call Description Line

Call Description Line for CPU

In Intel® oneAPI Math Kernel Library Verbose mode, each verbose-enabled function called from your application prints a call description line. The line begins with the `MKL_VERBOSE` character string and uses spaces as delimiters. The format of the rest of the line is subject to change in a future release.

The following table lists information contained in a call description line for Verbose with CPU applications and provides available links for more information:

Information	Description	Related Links
The name of the function.	Although the name printed may differ from the name used in the source code of the application (for example, the <code>cblas_</code> prefix of CBLAS functions is not printed), you can easily recognize the function by the printed name.	
Values of the arguments.	- The values are listed in the order of the formal argument list. The list directly follows the function name, it is parenthesized and comma-separated. - Arrays are printed as addresses (to see the alignment of the data). - Integer scalar parameters passed by reference are printed by value. Zero values are printed for <code>NULL</code> references. - Character values are printed without quotes. - For all parameters passed by reference, the values printed are the values <i>returned by the function</i>. For example, the printed value of the <code>info</code> parameter of a LAPACK function is its value after the function execution. - For verbose-enabled functions in the ScaLAPACK domain, in addition to the standard input parameters, information about blocking factors, MPI rank, and process grid is also printed.	
Time taken by the function.	- The time is printed in convenient units (seconds, milliseconds, and so on), which are explicitly indicated. - The time may fluctuate from run to run. - The time printed may occasionally be larger than the time actually taken by the function call, especially for small problem sizes and multi-socket machines. To reduce this effect, bind threads that call Intel® oneAPI Math Kernel Library to CPU cores by setting an affinity mask.	Managing Multi-core Performance for options to set an affinity mask.
Value of the <code>MKL_CBWR</code> environment variable.	The value printed is prefixed with <code>CNR</code> :	Getting Started with Conditional Numerical Reproducibility
Value of the <code>MKL_DYNAMIC</code> environment variable.	The value printed is prefixed with <code>Dyn</code> :	MKL_DYNAMIC

Information	Description	Related Links
Status of the Intel® oneAPI Math Kernel Library memory manager.	The value printed is prefixed with <code>FastMM:</code>	Avoiding Memory Leaks in oneMKL for a description of the Intel® oneAPI Math Kernel Library memory manager
OpenMP* thread number of the calling thread.	The value printed is prefixed with <code>TID:</code>	
Values of Intel® oneAPI Math Kernel Library environment variables defining the general and domain-specific numbers of threads, separated by a comma.	The first value printed is prefixed with <code>NThr:</code>	oneMKL-specific Environment Variables for Threading Control

The following is an example of a call description line (with OpenMP threading):

```
MKL_VERBOSE
DGEMM(n,n,1000,1000,240,0x7ffff708bb30,0x7ff2aea4c000,1000,0x7ff28e92b000,240,0x7ffff708bb38,0x7ff28e08d000,1000) 1.66ms CNR:OFF Dyn:1 FastMM:1 TID:0 NThr:16
```

The following is an example of a call description line (with TBB threading):

```
MKL_VERBOSE
DGEMM(n,n,1000,1000,240,0x7ffff708bb30,0x7ff2aea4c000,1000,0x7ff28e92b000,240,0x7ffff708bb38,0x7ff28e08d000,1000) 1.66ms CNR:OFF Dyn:1 FastMM:1
```

NOTE For more information about selected threading, refer to [Version Information Line](#).

The following information is not printed because of limitations of Intel® oneAPI Math Kernel Library Verbose mode:

- Input values of parameters passed by reference if the values were changed by the function.
For example, if a LAPACK function is called with a workspace query, that is, the value of the `lwork` parameter equals -1 on input, the call description line prints the result of the query and not -1.
- Return values of functions.
For example, the value returned by the function `ilaenv` is not printed.
- Floating-point scalars passed by reference.

Call Description Line for GPU

In Intel® oneAPI Math Kernel Library Verbose mode, each verbose-enabled function called from your application prints a call description line. The line begins with the `MKL_VERBOSE` character string and uses spaces as delimiters. The format of the rest of the line may change in a future release.

The following table lists information contained in a call description line for verbose with GPU applications.

See Also [GPU Information Line](#)

Information	Description
The name of the function	Although the name printed may differ from the name used in the source code of the application, you can easily recognize the function by the printed name.
The values of the arguments	- The values are listed in the order of the formal argument list. The list directly follows the function name, and it is parenthesized and comma-separated. - Arrays are printed as addresses (to show the alignment of the data). - Integer scalar parameters passed by reference are printed by value. Zero values are printed for NULL references. - Character values are printed without quotation marks. - For all parameters passed by reference, the values printed are the values returned by the function.
Time taken by the function	- If verbose is enabled with timing for GPU applications, kernel executions will become synchronous (previous kernel will block later kernels) and the measured time may include potential data transfers and/or data copies in host and devices. - If Verbose is enabled without timing for GPU applications, time will be printed out as 0. - The time is printed in convenient units (seconds, milliseconds, and so on), which are explicitly indicated. - The time may fluctuate from run to run. - The time printed may occasionally be larger than the time actually taken by the function call, especially for small problem sizes.
Device index	The index of the GPU device on which the kernel is being executed will be printed after the character string "GPU" (e.g. GPU0, GPU1, GPU2, etc). Use the index and refer to the GPU information lines for more information about the specific device. If the kernel is executed on the host CPU, this field will be empty.

The following is an example of a call description line:

```
MKL_VERBOSE FFT(dcfi64) 224.30us GPU0
```

Some Limitations

For GPU applications, the call description lines may be printed out-of-order (the order of the call description lines printed in the verbose output may not be the order in which the kernels are submitted in the functions) for the following two cases:

- Verbose is enabled without timing and the kernel executions stay asynchronous.
 - The kernel is not executed on one of the GPU devices, but on the host CPU (the device index will not be printed in this case).
-

Working with the Intel® oneAPI Math Kernel Library Cluster Software

9

includes distributed memory function domains for use on clusters:

- ScaLAPACK
- Cluster Fourier Transform Functions (Cluster FFT)
- Parallel Direct Sparse Solvers for Clusters (Cluster Sparse Solver)

ScaLAPACK, Cluster FFT, and Cluster Sparse Solver are only provided for the Intel® 64 and Intel® Many Integrated Core architectures.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[Intel® oneAPI Math Kernel Library Structure](#)

[Managing Performance of the Cluster Fourier Transform Functions](#)

[Intel® Distribution for LINPACK* Benchmark](#)

Message-Passing Interface Support

Intel® oneAPI Math Kernel Library ScaLAPACK, Cluster FFT, and Cluster Sparse Solver support implementations of the message-passing interface (MPI) identified in the *Release Notes*.

To link applications with ScaLAPACK, Cluster FFT, or Cluster Sparse Solver, you need to configure your system depending on your message-passing interface (MPI) implementation as explained below.

If you are using the Microsoft MPI, do the following:

1. Add Microsoft Compute Cluster Pack\include to the include path (assuming the default installation of the Microsoft MPI).
2. Add Microsoft Compute Cluster Pack\Lib\AMD64 to the library path.
3. Add msmpl.lib to your link command.

If you are using the Intel® MPI, do the following:

1. Add the following string to the include path: %ProgramFiles%\Intel\MPI\<ver>\intel64\include, where <ver> is the directory for a particular MPI version, for example, %ProgramFiles%\Intel\MPI\5.1\intel64\include.
2. Add the following string to the library path: %ProgramFiles%\Intel\MPI\<ver>\intel64\lib, for example, %ProgramFiles%\Intel\MPI\5.1\intel64\lib.
3. Add impi.lib and impicxx.lib to your link command.

Check the documentation that comes with your MPI implementation for implementation-specific details of linking.

Linking with oneMKL Cluster Software

The Intel® oneAPI Math Kernel Library ScaLAPACK, Cluster FFT, and Cluster Sparse Solver support MPI implementations identified in the *Intel® oneAPI Math Kernel Library Release Notes*.

To link with ScaLAPACK, Cluster FFT, and/or Cluster Sparse Solver, use the following commands:

```
set lib =<path to MKL libraries>;<path to MPI libraries>;%lib%
<linker> <files to link> [<MKL cluster library>] <BLACS><MKL core libraries><MPI
libraries>
```

where the placeholders stand for paths and libraries as explained in the following table:

<path to MKL libraries>	<mkl directory>\lib\intel64. If you performed the Scripts to Set Environment Variables Setting Environment Variables step of the Getting Started process, you do not need to add this directory to the lib environment variable.
<path to MPI libraries>	Typically the lib subdirectory in the MPI installation directory.
<linker>	One of icl, ifort, xilink.
<MKL cluster library>	One of libraries for ScaLAPACK or Cluster FFT listed in Appendix C: Directory Structure in Detail . For example, for the LP64 interface, it is mkl_scalapack_lp64.lib or mkl_cdft_core.lib. Cluster Sparse Solver does not require an additional computation library.
<BLACS>	The BLACS library corresponding to your , programming interface (LP64 or ILP64), and MPI version. These libraries are listed in Appendix C: Directory Structure in Detail . For example, for the LP64 interface, choose one of mkl_blacs_intelmpi_lp64.lib, or mkl_blacs_msmapi_lp64.lib in the case of static linking and mkl_blacs_lp64_dll.lib in the case of dynamic linking.
<MKL core libraries>	Intel® oneAPI Math Kernel Library libraries other than libraries with ScaLAPACK, Cluster FFT, or Cluster Sparse Solver.

Tip

Use the [Using the Link-line Advisor](#) to quickly choose the appropriate set of <MKL cluster Library>, <BLACS>, and <MKL core libraries>.

Intel MPI provides prepackaged scripts for its linkers to help you link using the respective linker. Therefore, if you are using Intel MPI, the best way to link is to use the following commands:

```
<path to Intel MPI binaries>\vars.bat
set lib = <path to MKL libraries>;%lib%
<mpilinker><files to link> [<MKL cluster Library>] <BLACS><MKL core libraries>
```

where the placeholders that are not yet defined are explained in the following table:

<code><path to MPI binaries></code>	By default, the <code>bin</code> subdirectory in the MPI installation directory.
<code><MPI linker></code>	<code>mpicl</code> or <code>mpiifort</code>

See Also

[Linking Your Application with the Intel® oneAPI Math Kernel Library](#)

[Examples of Linking for Clusters](#)

Determining the Number of OpenMP* Threads

The OpenMP* run-time library responds to the environment variable `OMP_NUM_THREADS`. Intel® oneAPI Math Kernel Library also has other mechanisms to set the number of OpenMP threads, such as the `MKL_NUM_THREADS` or `MKL_DOMAIN_NUM_THREADS` environment variables (see [Using Additional Threading Control](#)).

Make sure that the relevant environment variables have the same and correct values on all the nodes. Intel® oneAPI Math Kernel Library does not set the default number of OpenMP threads to one, but depends on the OpenMP libraries used with the compiler to set the default number. For the threading layer based on the Intel compiler (`mkl_intel_thread.lib`), this value is the number of CPUs according to the OS.

Caution

Avoid over-prescribing the number of OpenMP threads, which may occur, for instance, when the number of MPI ranks per node and the number of OpenMP threads per node are both greater than one. The number of MPI ranks per node multiplied by the number of OpenMP threads per node should not exceed the number of hardware threads per node.

The `OMP_NUM_THREADS` environment variable is assumed in the discussion below.

Set `OMP_NUM_THREADS` so that the product of its value and the number of MPI ranks per node equals the number of real processors or cores of a node. If the Intel® Hyper-Threading Technology is enabled on the node, use only half number of the processors that are visible on Windows OS.

Important

For Cluster Sparse Solver, set the number of OpenMP threads to a number greater than one because the implementation of the solver only supports a multithreaded algorithm.

See Also

[Setting Environment Variables on a Cluster](#)

Using DLLs

All the needed DLLs must be visible on all the nodes at run time, and you should install on each node of the cluster. You can use Remote Installation Services (RIS) provided by Microsoft to remotely install the library on each of the nodes that are part of your cluster. The best way to make the DLLs visible is to point to these libraries in the `PATH` environment variable. See [Setting Environment Variables on a Cluster](#) on how to set the value of the `PATH` environment variable.

The ScaLAPACK DLLs in the `<mkl_directory>\redist\intel64` directory use the MPI dispatching mechanism. MPI dispatching is based on the `MKL_BLACS_MPI` environment variable. The BLACS DLL uses `MKL_BLACS_MPI` for choosing the needed MPI libraries. The table below lists possible values of the variable.

Value	Comment
INTELMPI	Default value. Intel MPI is used for message passing
MSMPI	Microsoft MPI is used for message passing
CUSTOM	Intel® oneAPI Math Kernel Library MPI wrappers built with a custom MPI are used for message passing

If you are using a non-default MPI, assign the same appropriate value to `MKL_BLACS_MPI` on all nodes.

See Also

[Setting Environment Variables on a Cluster](#)
[Notational Conventions](#)

Setting Environment Variables on a Cluster

By default, when you call the MPI launch command `mpiexec`, the entire launching node environment is passed to the MPI processes. However, if there are undefined variables or variables that are different from what is stored in your environment, you can use `-env` or `-genv` options with `mpiexec`. Each of these options take two arguments- the name and the value of the environment variable to be passed.

```
-genv NAME1 VALUE1 -genv NAME2 VALUE2
```

```
-env NAME VALUE -genv
```

See these Intel MPI examples on how to set the value of `OMP_NUM_THREADS` explicitly:

```
mpiexec -genv OMP_NUM_THREADS 2 ....
```

```
mpiexec -n 1 -host first -env OMP_NUM_THREADS 2 test.exe : -n 2 -host second -env  
OMP_NUM_THREADS 3 test.exe ....
```

See these Intel MPI examples on how to set the value of `MKL_BLACS_MPI` explicitly:

```
mpiexec -genv MKL_BLACS_MPI INTELMPI ....
```

```
mpiexec -n 1 -host first -env MKL_BLACS_MPI INTELMPI test.exe : -n 1 -host second -env  
MKL_BLACS_MPI INTELMPI test.exe.
```

If you want to pass all environment variables by default and avoid passing these values explicitly, modify the user or system environment variables on each Windows node. From the **Start** menu, select **Settings > Control Panel > System > Advanced > Environment Variables**.

If you are using Microsoft MPI, the ways of setting environment variables described above are also applicable if the Microsoft Single Program Multiple Data (SPMD) process managers are running in a debug mode on all nodes of the cluster. However, the best way to set environment variables is by using the Job Scheduler with the Microsoft Management Console (MMC) and/or the Command Line Interface (CLI) to submit a job and pass environment variables. For more information about MMC and CLI, see the Microsoft Help and Support page at the Microsoft Web site (<http://www.microsoft.com/>).

Interaction with the Message-passing Interface

To improve performance of cluster applications, it is critical for Intel® oneAPI Math Kernel Library to use the optimal number of threads, as well as the correct thread affinity. Usually, the optimal number is the number of available cores per node divided by the number of MPI processes per node. You can set the number of threads using one of the available methods, described in [Techniques to Set the Number of Threads](#).

If the number of threads is not set, Intel® oneAPI Math Kernel Library checks whether it runs under MPI provided by the Intel® MPI Library. If this is true, the following environment variables define Intel® oneAPI Math Kernel Library threading behavior:

- `I_MPI_THREAD_LEVEL`
- `MKL_MPI_PPN`
- `I_MPI_NUMBER_OF_MPI_PROCESSES_PER_NODE`
- `I_MPI_PIN_MAPPING`
- `OMPI_COMM_WORLD_LOCAL_SIZE`
- `MPI_LOCALNRANKS`

The threading behavior depends on the value of `I_MPI_THREAD_LEVEL` as follows:

- 0 or undefined.

Intel® oneAPI Math Kernel Library considers that thread support level of Intel MPI Library is `isMPI_THREAD_SINGLE` and defaults to sequential execution.

- 1, 2, or 3.

This value determines Intel® oneAPI Math Kernel Library conclusion of the thread support level:

- 1 - `MPI_THREAD_FUNNELED`
- 2 - `MPI_THREAD_SERIALIZED`
- 3 - `MPI_THREAD_MULTIPLE`

In all these cases, Intel® oneAPI Math Kernel Library determines the number of MPI processes per node using the other environment variables listed and defaults to the number of threads equal to the number of available cores per node divided by the number of MPI processes per node.

Important

Instead of relying on the discussed implicit settings, explicitly set the number of threads for Intel® oneAPI Math Kernel Library.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[Managing Multi-core Performance](#)

[Intel® Software Documentation Library](#) for more information on Intel MPI Library
for more information on Intel MPI Library

Using a Custom Message-Passing Interface

While different message-passing interface (MPI) libraries are compatible at the application programming interface (API) level, they are often incompatible at the application binary interface (ABI) level. Therefore, Intel® oneAPI Math Kernel Library provides a set of prebuilt BLACS libraries that support certain MPI libraries, but this, however, does not enable use of Intel® oneAPI Math Kernel Library with other MPI libraries. To fill this gap, Intel® oneAPI Math Kernel Library also includes the MKL MPI wrapper, which provides an MPI-independent ABI to Intel® oneAPI Math Kernel Library. The adaptor is provided as source code. To use Intel® oneAPI Math Kernel Library with an MPI library that is not supported by default, you can use the adaptor to build custom static or dynamic BLACS libraries and use them similarly to the prebuilt libraries.

Building a Custom BLACS Library

The MKL MPI wrapper is located in the `<mkl_directory>\interfaces\mklmpi` directory.

To build a custom BLACS library, from the above directory run the `nmake` command.

For example: the command

```
nmake libintel64
```

builds a static custom BLACS library `mkl_blacs_custom_lp64.lib` using the default MPI compiler on your system. Look into the `<mkl_directory>\interfaces\mklnmpi\makefile` for targets and variables that define how to build the custom library. In particular, you can specify the compiler through the `MPICC` variable.

For more control over the building process, refer to the documentation available through the command

```
nmake help
```

Using a Custom BLACS Library

In the case of static linking, use custom BLACS libraries exactly the same way as you use the prebuilt BLACS libraries, but pass the custom library to the linker. For example, instead of passing the `mkl_blacs_intelmpi_lp64.lib` static library, pass `mkl_blacs_custom_lp64.lib`.

To use a dynamic custom BLACS library:

1. Link your application the same way as when you use the prebuilt BLACS library.
2. Call the `mkl_set_mpi` support function or set the `MKL_BLACS_MPI` environment variable to one of the following values:
 - `CUSTOM`
to load a custom library with the default name `mkl_blacs_custom_lp64.dll` or `mkl_blacs_custom_ilp64.dll`, depending on whether the BLACS interface linked against your application is LP64 or ILP64.
 - `<dll_name>`
to load the specified BLACS DLL.

NOTE

Intel® oneAPI Math Kernel Library looks for the specified DLL either in the directory with Intel® oneAPI Math Kernel Library dynamic libraries or in the directory with the application executable.

For a description of the `mkl_set_mpi` function, see the *Intel® oneAPI Math Kernel Library Developer Reference*.

See Also

[Linking with Intel® oneAPI Math Kernel Library Cluster Software](#)

Examples of Linking for Clusters

This section provides examples of linking with ScaLAPACK, Cluster FFT, and Cluster Sparse Solver.

Note that a binary linked with the Intel® oneAPI Math Kernel Library cluster function domains runs the same way as any other MPI application (refer to the documentation that comes with your MPI implementation).

For further linking examples, see the support website for Intel products at <https://software.intel.com/content/www/us/en/develop/support.html>.

See Also

[Directory Structure in Detail](#)

Examples for Linking a C Application

These examples illustrate linking of an application under the following conditions:

- Main module is in C.
- Intel MPI is installed in C:\Program Files(x86)\Intel\oneAPI\mpi. Instructions on how to install Intel MPI can be found on the [Get Started with Intel® MPI Library for Windows* OS](#) and [Intel® oneAPI Toolkits Installation Guide for Windows*](#) pages.
- You are using the Intel® C++ Compiler.
- Intel® oneAPI Math Kernel Library functions use LP64 interfaces.

To set the Intel MPI environment, you must run the `<path to Intel MPI binaries>\vars.bat` script.

To link with ScaLAPACK for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%

icl <user files to link> mkl_scalapack_lp64.lib blacs_intelmpi_lp64.lib
mkl_intel_lp64.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib
bufferoverflowu.lib
```

To link with Cluster FFT for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%

icl <user files to link> mkl_cdft_core.lib mkl_blacs_intelmpi_lp64.lib
mkl_intel_lp64.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib
bufferoverflowu.lib
```

To link with Cluster Sparse Solver for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%

icl <user files to link> mkl_blacs_intelmpi_lp64.lib mkl_intel_lp64.lib
mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib bufferoverflowu.lib
```

See Also

[Linking with oneMKL Cluster Software](#)

[Using the Link-line Advisor](#)

[Linking with System Libraries](#)

Examples for Linking a Fortran Application

These examples illustrate linking of an application under the following conditions:

- Main module is in Fortran.
- Intel MPI is installed in C:\Program Files(x86)\Intel\oneAPI\mpi. Instructions on how to install Intel MPI can be found on the [Get Started with Intel® MPI Library for Windows* OS](#) and [Intel® oneAPI Toolkits Installation Guide for Windows*](#) pages.
- You are using the Intel® Fortran Compiler.
- Intel® oneAPI Math Kernel Library functions use LP64 interfaces.

To set the Intel MPI environment, you must run the `<path to Intel MPI binaries>\vars.bat` script.

To link with ScaLAPACK for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%

ifort <user files to link> mkl_scalapack_lp64.lib mkl_blacs_intelmpi_lp64.lib
mkl_intel_lp64.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib
bufferoverflowu.lib
```

To link with Cluster FFT for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%  
  
ifort <user files to link> mkl_cdft_core.lib mkl_blacs_intelmpi_lp64.lib  
mkl_intel_lp64.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib  
bufferoverflowu.lib
```

To link with Cluster Sparse Solver for a cluster of Intel® 64 architecture based systems, set the environment variable and use the link line as follows:

```
set lib=<mkl directory>\lib\intel64_win;%lib%  
  
ifort <user files to link> mkl_blacs_intelmpi_lp64.lib mkl_intel_lp64.lib  
mkl_intel_thread.lib mkl_core.lib libiomp5md.lib impi.lib bufferoverflowu.lib
```

See Also

[Linking with oneMKL Cluster Software](#)

[Using the Link-line Advisor](#)

[Linking with System Libraries](#)

Managing Behavior of the Intel® oneAPI Math Kernel Library with Environment Variables

10

See Also

[Intel® oneAPI Math Kernel Library-specific Environment Variables for Threading Control](#)

Specifying the Code Branches

for how to use an environment variable to specify the code branch for Conditional Numerical Reproducibility

Using Intel® oneAPI Math Kernel Library Verbose Mode

for how to use an environment variable to set the verbose mode

Managing Behavior of Function Domains with Environment Variables

Setting the Default Mode of Vector Math with an Environment Variable

enables overriding the default setting of the Vector Mathematics (VM) global mode using the `MKL_VML_MODE` environment variable.

Because the mode is set or can be changed in different ways, their precedence determines the actual mode used. The settings and function calls that set or change the VM mode are listed below, with the precedence growing from lowest to highest:

1. The default setting
2. The `MKL_VML_MODE` environment variable
3. A call `vmlSetMode` function
4. A call to any VM function other than a service function

For more details, see the Vector Mathematical Functions section in the *Intel® oneAPI Math Kernel Library Developer Reference* and the description of the `vmlSetMode` function in particular.

To set the `MKL_VML_MODE` environment variable, use the following command:

```
set MKL_VML_MODE=<mode-string>
```

In this command, `<mode-string>` controls error handling behavior and computation accuracy, consists of one or several comma-separated values of the `mode` parameter listed in the table below, and meets these requirements:

- Not more than one accuracy control value is permitted
- Any combination of error control values except `VML_ERRMODE_DEFAULT` is permitted
- No denormalized numbers control values are permitted

Values of the `mode` Parameter

Value of <code>mode</code>	Description
Accuracy Control	
<code>VML_HA</code>	high accuracy versions of VM functions
<code>VML_LA</code>	low accuracy versions of VM functions
<code>VML_EP</code>	enhanced performance accuracy versions of VM functions
Denormalized Numbers Handling Control	

Value of <i>mode</i>	Description
VML_FTZDAZ_ON	Faster processing of denormalized inputs is enabled.
VML_FTZDAZ_OFF	Faster processing of denormalized inputs is disabled.
VML_FTZDAZ_CURRENT	Keep the current CPU settings for denormalized inputs.
Error Mode Control	
VML_ERRMODE_IGNORE	On computation error, VM Error status is updated, but otherwise no action is set. Cannot be combined with other VML_ERRMODE settings.
VML_ERRMODE_NOERR	On computation error, VM Error status is not updated and no action is set. Cannot be combined with other VML_ERRMODE settings.
VML_ERRMODE_STDERR	On error, the error text information is written to <i>stderr</i> .
VML_ERRMODE_EXCEPT	On error, an exception is raised.
VML_ERRMODE_CALLBACK	On error, an additional error handler function is called.
VML_ERRMODE_DEFAULT	On error, an exception is raised and an additional error handler function is called.

This command provides an example of valid settings for the `MKL_VML_MODE` environment variable:

```
set MKL_VML_MODE=VML_LA,VML_ERRMODE_ERRNO,VML_ERRMODE_STDERR
```

NOTE

VM ignores the `MKL_VML_MODE` environment variable in the case of incorrect or misspelled settings of *mode*.

Managing Performance of the Cluster Fourier Transform Functions

Performance of Intel® oneAPI Math Kernel Library Cluster FFT (CFFT) in different applications mainly depends on the cluster configuration, performance of message-passing interface (MPI) communications, and configuration of the run. Note that MPI communications usually take approximately 70% of the overall CFFT compute time. For more flexibility of control over time-consuming aspects of CFFT algorithms, Intel® oneAPI Math Kernel Library provides the `MKL_CDFT` environment variable to set special values that affect CFFT performance. To improve performance of your application that intensively calls CFFT, you can use the environment variable to set optimal values for your cluster, application, MPI, and so on.

The `MKL_CDFT` environment variable has the following syntax, explained in the table below:

```
MKL_CDFT=option1[=value1],option2[=value2],...,optionN[=valueN]
```

Important

While this table explains the settings that usually improve performance under certain conditions, the actual performance highly depends on the configuration of your cluster. Therefore, experiment with the listed values to speed up your computations.

Option	Possible Values	Description
<code>alltoallv</code>	0 (default)	Configures CFFT to use the standard <code>MPI_Alltoallv</code> function to perform global transpositions.
	1	Configures CFFT to use a series of calls to <code>MPI_Isend</code> and <code>MPI_Irecv</code> instead of the <code>MPI_Alltoallv</code> function.
	4	Configures CFFT to merge global transposition with data movements in the local memory. CFFT performs global transpositions by calling <code>MPI_Isend</code> and <code>MPI_Irecv</code> in this case.

Option	Possible Values	Description
wo_omatcopy		Use this value in a hybrid case (MPI + OpenMP), especially when the number of processes per node equals one.
	0	Configures CFFT to perform local FFT and local transpositions separately. CFFT usually performs faster with this value than with wo_omatcopy = 1 if the configuration parameter DFTI_TRANSPOSE has the value of DFTI_ALLOW. See the <i>Intel® oneAPI Math Kernel Library Developer Reference</i> for details.
	1	Configures CFFT to merge local FFT calls with local transpositions. CFFT usually performs faster with this value than with wo_omatcopy = 0 if DFTI_TRANSPOSE has the value of DFTI_NONE.
	-1 (default)	Enables CFFT to decide which of the two above values to use depending on the value of DFTI_TRANSPOSE.
enable_soi	Not applicable	A flag that enables low-communication Segment Of Interest FFT (SOI FFT) algorithm for one-dimensional complex-to-complex CFFT, which requires fewer MPI communications than the standard nine-step (or six-step) algorithm.
Caution While using fewer MPI communications, the SOI FFT algorithm incurs a minor loss of precision (about one decimal digit).		

The following example illustrates usage of the environment variable:

```
set MKL_CDFT=wo_omatcopy=1,alltoallv=4,enable_soi
mpirun -ppn 2 -n 16 mkl_cdft_app.exe
```

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Managing Invalid Input Checking in LAPACKE Functions

The high-level interface includes an optional, on by default, NaN check on all matrix inputs before calling any LAPACK routine. This option affects all routines. If an input matrix contains any NaNs, the input parameter corresponding to this matrix is flagged with a return value error. For example, if the fifth parameter is found to contain a NaN, the routine returns the value, -5. The middle-level interface does not contain the NaN check.

NaN checking on matrix input can be expensive. By default, NaN checking is turned **on**. LAPACKE provides a way to set it through the environment variable:

- Setting environment variable LAPACKE_NANCHECK to 0 turns OFF NaN-checking
- Setting environment variable LAPACKE_NANCHECK to 1 turns ON NaN-checking

The other way is to call the LAPACKE_set_nancheck function; see the Developer Reference for C's [LAPACK Auxiliary Routines](#) section for more information.

Note that the NaN-checking flag value set by the call to LAPACKE_set_nancheck always has higher priority than the environment variable, LAPACKE_NANCHECK.

Instruction Set Specific Dispatching on Intel® Architectures

Intel® oneAPI Math Kernel Library automatically queries and then dispatches the code path supported on your Intel® processor to the optimal instruction set architecture (ISA) by default. The `MKL_ENABLE_INSTRUCTIONS` environment variable or the `mkl_enable_instructions` support function enables you to dispatch to an ISA-specific code path of your choice. For example, you can run the Intel® Advanced Vector Extensions (Intel® AVX) code path on an Intel processor based on Intel® Advanced Vector Extensions 2 (Intel® AVX2). This feature is not available on non-Intel processors.

In some cases Intel® oneAPI Math Kernel Library also provides support for upcoming architectures ahead of hardware availability, but the library does not automatically dispatch the code path specific to an upcoming ISA by default. If for your exploratory work you need to enable an ISA for an Intel processor that is not yet released or if you are working in a simulated environment, you can use the `MKL_ENABLE_INSTRUCTIONS` environment variable or `mkl_enable_instructions` support function.

The following table lists possible values of `MKL_ENABLE_INSTRUCTIONS` alongside the corresponding ISA supported by a given processor. `MKL_ENABLE_INSTRUCTIONS` dispatches to the default ISA if the ISA requested is not supported on the particular Intel processor. For example, if you request to run the Intel AVX512 code path on a processor based on Intel AVX2, Intel® oneAPI Math Kernel Library runs the Intel AVX2 code path. The table also explains whether the ISA is dispatched by default on the processor that supports this ISA.

Value of <code>MKL_ENABLE_INSTRUCTIONS</code>	ISA	Dispatched by Default
AVX512	Intel® Advanced Vector Extensions (Intel® AVX-512) for systems based on Intel® Xeon® processors	Yes
AVX512_E1	Intel® Advanced Vector Extensions (Intel® AVX-512) with support for Vector Neural Network Instructions.	Yes
AVX512_E2	ICX: Intel® Advanced Vector Extensions (Intel® AVX-512) enabled processors.	Yes
AVX512_E3	Intel® Advanced Vector Extensions 512 (Intel® AVX-512) with support of Vector Neural Network Instructions supporting BF16 enabled processors.	Yes
AVX512_E4	Intel® Advanced Vector Extensions 512 (Intel® AVX-512) with Intel® Deep Learning Boost (Intel® DL Boost) and bfloat16 support and Intel® Advanced Matrix Extensions (Intel® AMX) with bfloat16 and 8-bit integer support.	Yes
AVX2	Intel® AVX2	Yes
AVX2_E1	Intel® Advanced Vector Extensions 2 (Intel® AVX2) with support for Intel® Deep Learning Boost (Intel® DL Boost).	Yes
AVX	Intel® AVX	Yes

For more details about the `mkl_enable_instructions` function, including the argument values, see the *Intel® oneAPI Math Kernel Library Developer Reference*.

For example:

- To turn on automatic CPU-based dispatching of Intel AVX-512 with support of Intel DL Boost, bfloat16, Intel AMX with bfloat16 and 8-bit integer, and FP16 instruction, do one of the following:
 - Call
`mkl_enable_instructions(AVX512_E4)`
 - Set the environment variable:
`set MKL_ENABLE_INSTRUCTIONS=AVX512_E4`
- To configure the library not to dispatch more recent architectures than Intel AVX2, do one of the following:
 - Call
`mkl_enable_instructions(MKL_ENABLE_AVX2)`
 - Set the environment variable:
`set MKL_ENABLE_INSTRUCTIONS=AVX2`

NOTE

Settings specified by the `mkl_enable_instructions` function take precedence over the settings specified by the `MKL_ENABLE_INSTRUCTIONS` environment variable.

Product and Performance Information
Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex . Notice revision #20201201

Programming with Intel® Math Kernel Library in Integrated Development Environments (IDE)

11

Configuring Your Integrated Development Environment to Link with Intel® oneAPI Math Kernel Library

See these Knowledge Base articles for how to configure your Integrated Development Environment for linking with Intel® oneAPI Math Kernel Library:

- Compiling and Linking Intel® oneAPI Math Kernel Library with Microsoft* Visual C/C++* (<http://software.intel.com/en-us/articles/intel-math-kernel-library-intel-mkl-compiling-and-linking-with-microsoft-visual-cc>)
- How to Build an Intel® oneAPI Math Kernel Library Application with Intel® Visual Fortran Compiler (<http://software.intel.com/en-us/articles/how-to-build-mkl-application-in-intel-visual-fortran-msvc2005>)

Configuring the Microsoft Visual C/C++* Development System to Link with Intel® MKL

Steps for configuring Microsoft Visual C/C++* development system for linking with depend on whether you installed the C++ Integration(s) in Microsoft Visual Studio* component of the Intel® Parallel Studio XE Composer Edition:

- If you installed the integration component, see [Automatically Linking Your Microsoft Visual C/C++* Project with Intel® MKL](#).
- If you did not install the integration component or need more control over Intel® oneAPI Math Kernel Library libraries to link, you can configure the Microsoft Visual C++* development system by performing the following steps. Though some versions of the Visual C++* development system may vary slightly in the menu items mentioned below, the fundamental configuring steps are applicable to all these versions.
 1. In Solution Explorer, right-click your project and click **Properties**
 2. Select **Configuration Properties > VC++ Directories**
 3. Select **Include Directories**. Add the directory for the Intel® oneAPI Math Kernel Library include files, that is, `<mkl_directory>\include`
 4. Select **Library Directories**. Add architecture-specific directories for Intel® oneAPI Math Kernel Library and OpenMP* libraries, for example: `<mkl_directory>\lib\ia32` and `<compiler_directory>\windows\compiler\lib\ia32_win`
 5. Select **Executable Directories**. Add architecture-specific directories with dynamic-link libraries:
 - For OpenMP* support, for example: `<compiler_directory>\windows\redist\ia32_win\compiler`
 - For Intel® oneAPI Math Kernel Library (only if you link dynamically), for example: `<mkl_directory>\redist\ia32`
 6. Select **Configuration Properties > Custom Build Setup > Additional Dependencies**. Add the libraries required, for example, `mkl_intel_c.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib`

See Also

[Intel® Software Documentation Library](#) for the documentation for Intel Parallel Studio XE Composer Edition

for the documentation for Intel Parallel Studio XE Composer Edition

[Linking in Detail](#)

[Notational Conventions](#)

Configuring Intel® Visual Fortran to Link with Intel MKL

Steps for configuring Intel® Visual Fortran for linking with depend on whether you installed the Visual Fortran Integration(s) in Microsoft Visual Studio* component of the Intel® Parallel Studio XE Composer Edition:

- If you installed the integration component, see [Automatically Linking Your Intel® Visual Fortran Project with Intel® MKL](#).
- If you did not install the integration component or need more control over Intel® oneAPI Math Kernel Library libraries to link, you can configure your project as follows:
 1. Select **Project > Properties > Linker > General > Additional Library Directories**. Add architecture-specific directories for Intel® oneAPI Math Kernel Library and OpenMP* libraries, for example: `<mkl_directory>\lib\ia32` and `<parent_directory>\compiler\lib\ia32`
 2. Select **Project > Properties > Linker > Input > Additional Dependencies**. Insert names of the required libraries, for example: `mkl_intel_c.lib mkl_intel_thread.lib mkl_core.lib libiomp5md.lib`
 3. Select **Project > Properties > Debugging > Environment**. Add architecture-specific paths to dynamic-link libraries:
 - For OpenMP* support; for example: enter `PATH=%PATH%;<compiler_directory>\windows\redist\ia32_win\compiler`
 - For Intel® oneAPI Math Kernel Library (only if you link dynamically); for example: enter `PATH=%PATH%;<mkl_directory>\redist\ia32`

See Also

[Intel® Software Documentation Library](#) for the documentation for Intel Parallel Studio XE Composer Edition

for the documentation for Intel Parallel Studio XE Composer Edition

[Notational Conventions](#)

Getting Assistance for Programming in the Microsoft Visual Studio* IDE

Using Context-Sensitive Help

You can get context-sensitive help when typing your code in the Visual Studio* IDE Code Editor. To open the help topic describing an Intel® oneAPI Math Kernel Library function called in your code, select the function name and press F1. The topic with the function description opens in the Microsoft Help Viewer or your Web browser depending on the Visual Studio IDE Help settings.

Using the IntelliSense* Capability

IntelliSense is a set of native Visual Studio*(VS) IDE features that make language references easily accessible.

The user programming with Intel® oneAPI Math Kernel Library in the VS Code Editor can employ two IntelliSense features: Parameter Info and Complete Word.

Both features use header files. Therefore, to benefit from IntelliSense, make sure the path to the include files is specified in the VS or solution settings. For example, see [Configuring the Microsoft Visual C/C++* Development System to Link with Intel® MKL](#) on how to do this.

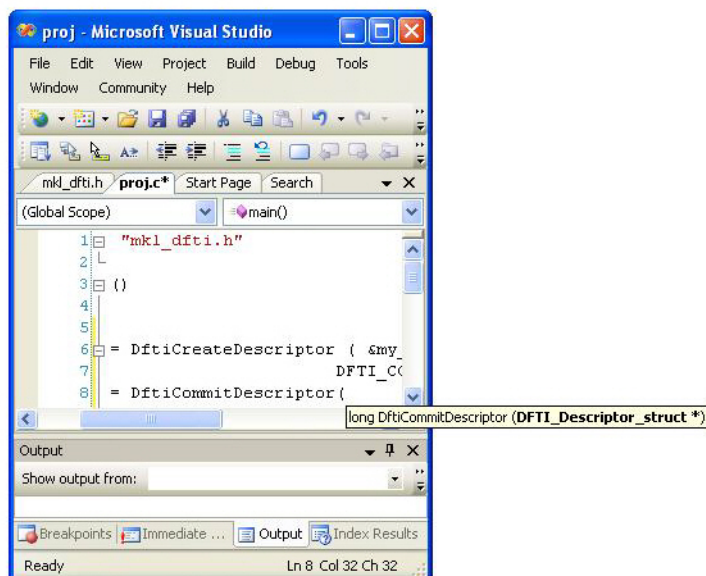
Parameter Info

The Parameter Info feature displays the parameter list for a function to give information on the number and types of parameters. This feature requires adding the `include` statement with the appropriate Intel® oneAPI Math Kernel Library header file to your code.

To get the list of parameters of a function specified in the header file,

1. Type the function name.
2. Type the opening parenthesis.

This brings up the tooltip with the list of the function parameters:

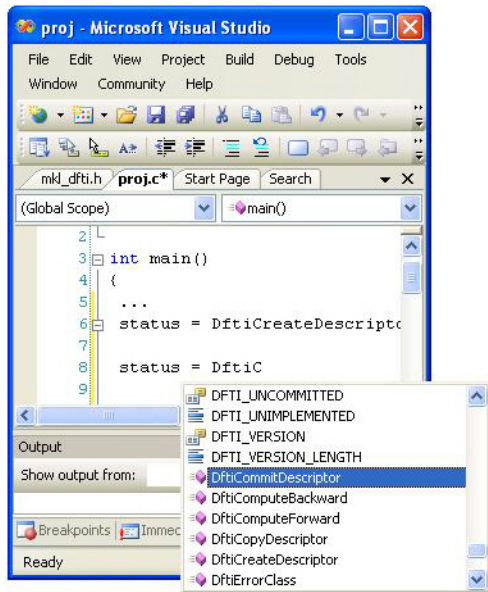


Complete Word

For a software library, the Complete Word feature types or prompts for the rest of the name defined in the header file once you type the first few characters of the name in your code. This feature requires adding the `include` statement with the appropriate Intel® oneAPI Math Kernel Library header file to your code.

To complete the name of the function or named constant specified in the header file,

1. Type the first few characters of the name.
2. Press `Alt+RIGHT ARROW` or `Ctrl+SPACEBAR`.
If you have typed enough characters to disambiguate the name, the rest of the name is typed automatically. Otherwise, a pop-up list appears with the names specified in the header file.
3. Select the name from the list, if needed.



Intel® oneAPI Math Kernel Library Benchmarks

12

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

Intel Optimized LINPACK Benchmark for Windows*

Intel Optimized LINPACK Benchmark for Windows* is a generalization of the LINPACK 1000 benchmark. It solves a dense ($\text{real} \times 8$) system of linear equations ($Ax=b$), measures the amount of time it takes to factor and solve the system, converts that time into a performance rate, and tests the results for accuracy. The generalization is in the number of equations (N) it can solve, which is not limited to 1000. It uses partial pivoting to assure the accuracy of the results.

Do not use this benchmark to report LINPACK 100 performance because that is a compiled-code only benchmark. This is a shared-memory (SMP) implementation which runs on a single platform. Do not confuse this benchmark with:

- Intel® Distribution for LINPACK* Benchmark, which is a distributed memory version of the same benchmark.
- LINPACK, the library, which has been expanded upon by the LAPACK library.

Intel provides optimized versions of the LINPACK benchmarks to help you obtain high LINPACK benchmark results on your genuine Intel processor systems more easily than with the High Performance Linpack (HPL) benchmark.

Additional information on this software, as well as on other Intel® software performance products, is available at <https://software.intel.com/content/www/us/en/develop/tools.html>.

Acknowledgement

This product includes software developed at the University of Tennessee, Knoxville, Innovative Computing Laboratories.

Contents of the Intel® Optimized LINPACK Benchmark

The Intel Optimized LINPACK Benchmark for Windows* contains the following files, located in the `benchmarks\linpack\` subdirectory of the directory:

File in <code>benchmarks\linpack\</code>	Description
<code>linpack_xeon32.exe</code>	The 32-bit program executable for a system based on Intel® Xeon® processor or Intel® Xeon® processor MP with or without Intel® Streaming SIMD Extensions 3 (SSE3).
<code>linpack_xeon64.exe</code>	The 64-bit program executable for a system with Intel Xeon processor using Intel® 64 architecture.
<code>runme_xeon32.bat</code>	A sample shell script for executing a pre-determined problem set for <code>linpack_xeon32.exe</code> .

File in benchmarks\linpack\	Description
runme_xeon64.bat	A sample shell script for executing a pre-determined problem set for linpack_xeon64.exe.
lininput_xeon32	Input file for a pre-determined problem for the runme_xeon32 script.
lininput_xeon64	Input file for a pre-determined problem for the runme_xeon64 script.
help.lpk	Simple help file.
xhelp.lpk	Extended help file.
These files are not available immediately after installation and appear as a result of execution of an appropriate runme script.	
win_xeon32.txt	Result of the runme_xeon32 script execution.
win_xeon64.txt	Result of the runme_xeon64 script execution.

See Also

High-level Directory Structure

Running the Software

To obtain results for the pre-determined sample problem sizes on a given system, type:

```
runme_xeon32.bat
```

```
runme_xeon64.bat
```

To run the software for other problem sizes, see the extended help included with the program. You can view extended help by running the program executable with the `-e` option:

```
linpack_xeon32.exe -e
```

```
linpack_xeon64.exe -e
```

The pre-defined data input files `lininput_xeon32` and `lininput_xeon64` are examples. Different systems have different numbers of processors or amounts of memory and therefore require new input files. The extended help can give insight into proper ways to change the sample input files.

Each input file requires the following minimum amount of memory:

```
lininput_xeon32          2 GB
```

```
lininput_xeon64         16 GB
```

If the system has less memory than the above sample data input requires, you may need to edit or create your own data input files, as explained in the extended help.

The Intel Optimized LINPACK Benchmark determines the optimal number of OpenMP threads to use. To run a different number, you can set the `OMP_NUM_THREADS` or `MKL_NUM_THREADS` environment variable inside a sample script. If you run the Intel Optimized LINPACK Benchmark without setting the number of threads, it defaults to the number of physical cores.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Product and Performance Information

Notice revision #20201201

Known Limitations of the Intel® Optimized LINPACK Benchmark

The following limitations are known for the Intel Optimized LINPACK Benchmark for Windows*:

- Intel Optimized LINPACK Benchmark supports only OpenMP threading
- Intel Optimized LINPACK Benchmark is threaded to effectively use multiple processors. So, in multi-processor systems, best performance will be obtained with the Intel® Hyper-Threading Technology turned off, which ensures that the operating system assigns threads to physical processors only.
- If an incomplete data input file is given, the binaries may either hang or fault. See the sample data input files and/or the extended help for insight into creating a correct data input file.

Intel® Distribution for LINPACK* Benchmark**Overview of the Intel® Distribution for LINPACK* Benchmark**

The Intel® Distribution for LINPACK* Benchmark is based on modifications and additions to High-Performance LINPACK (HPL) (<http://www.netlib.org/benchmark/hpl/>) from Innovative Computing Laboratories (ICL) at the University of Tennessee, Knoxville. The Intel® Distribution for LINPACK Benchmark can be used for TOP500 runs (see <http://www.top500.org>) and for benchmarking your cluster. To use the benchmark you need to be familiar with HPL usage. The Intel® Distribution for LINPACK Benchmark provides some enhancements designed to make the HPL usage more convenient and to use Intel® Message-Passing Interface (MPI) settings to improve performance.

The Intel® Distribution for LINPACK Benchmark measures the amount of time it takes to factor and solve a random dense system of linear equations ($Ax=b$) in `real*8` precision, converts that time into a performance rate, and tests the results for accuracy. The benchmark uses random number generation and full row pivoting to ensure the accuracy of the results.

Intel provides optimized versions of the LINPACK benchmarks to help you obtain high LINPACK benchmark results on your systems based on genuine Intel processors more easily than with the standard HPL benchmark. The prebuilt binaries require Intel® MPI library be installed on the cluster. The run-time version of Intel MPI library is free and can be downloaded from <https://www.software.intel.com/content/www/us/en/develop/tools.html>.

The Intel package includes software developed at the University of Tennessee, Knoxville, ICL, and neither the University nor ICL endorse or promote this product. Although HPL is redistributable under certain conditions, this particular package is subject to the license.

Intel® oneAPI Math Kernel Library provides prebuilt binaries that are linked against Intel MPI libraries either statically or dynamically. In addition, binaries linked with a customized MPI implementation can be created using the Intel® oneAPI Math Kernel Library MPI wrappers.

NOTE

Performance of statically and dynamically linked prebuilt binaries may be different. The performance of both depends on the version of Intel MPI you are using. You can build binaries statically or dynamically linked against a particular version of Intel MPI by yourself.

HPL code is homogeneous by nature: it requires that each MPI process runs in an environment with similar CPU and memory constraints. The Intel® Distribution for LINPACK Benchmark supports heterogeneity, meaning that the data distribution can be balanced to the performance requirements of each node, provided that there is enough memory on that node to support additional work. For information on how to configure Intel® oneAPI Math Kernel Library to use the internode heterogeneity, see [Heterogeneous Support in the Intel® Distribution for LINPACK Benchmark](#).

Contents of the Intel® Distribution for LINPACK* Benchmark

The Intel® Distribution for LINPACK Benchmark includes prebuilt binaries linked with Intel® MPI library. For a customized MPI implementation, tools are also included to build a binary using Intel® oneAPI Math Kernel Library MPI wrappers. All the files are located in the `.\benchmarks\mp_linpack\` subdirectory of the Intel® oneAPI Math Kernel Library directory.

File in <code><mkl_directory>\benchmarks\mp_linpack\</code>	Contents
<code>COPYRIGHT</code>	Original Netlib HPL copyright document.
<code>readme.txt</code>	Information about the files provided.
Prebuilt executables for performance testing	
<code>xhpl_intel64_dynamic.exe</code>	Prebuilt binary for the Intel® 64 architecture dynamically linked against Intel MPI library [†] .
Run scripts and an input file example	
<code>runme_intel64_dynamic.bat</code>	Sample run script for the Intel® 64 architecture and binary dynamically linked against Intel MPI library.
<code>runme_intel64_prv.bat</code>	Script that sets HPL environment variables. It is called by <code>runme_intel64_dynamic.bat</code> .
<code>HPL.dat</code>	Example of an HPL configuration file.
Prebuilt libraries and utilities for building with a customized MPI implementation	
<code>libhpl_intel64.lib</code>	Library file required to build Intel® Distribution for LINPACK Benchmark for the Intel® 64 architecture with a customized MPI implementation.
<code>HPL_main.c</code>	Source code required to build Intel® Distribution for LINPACK Benchmark for the Intel® 64 architecture with a customized MPI implementation.
<code>build.bat</code>	Build script for creating Intel® Distribution for LINPACK Benchmark for the Intel® 64 architecture with a customized MPI implementation.
Utilities for building Netlib HPL from source code	
<code>Make.Windows_Intel64</code>	Makefile for building Netlib HPL from source code.

[†]For a list of supported versions of the Intel MPI Library, see system requirements in the Intel® oneAPI Math Kernel Library Release Notes.

See Also

[High-level Directory Structure](#)

Building the Intel® Distribution for LINPACK* Benchmark for a Customized MPI Implementation

The Intel® Distribution for LINPACK Benchmark contains a sample build script `build.bat`. If you are using a customized MPI implementation, this script builds a binary using Intel® oneAPI Math Kernel Library MPI wrappers. To build the binary, follow these steps:

1. Specify the location of Intel® oneAPI Math Kernel Library to be used (`MKLROOT`)

2. Set up your MPI environment
3. Run the script `build.bat`

See Also

[Contents of the Intel® Distribution for LINPACK Benchmark](#)

Building the Netlib HPL from Source Code

The source code for Intel® Distribution for LINPACK* Benchmark is not provided. However, you can download reference Netlib HPL source code from <http://www.netlib.org/benchmark/hpl/> . To build the HPL:

1. Download and extract the source code.
2. Copy the makefile `<mk1 directory>\benchmarks\mp_linpack\Make.Windows_Intel64` to your HPL directory
3. Edit `Make.Windows_Intel64` as appropriate
4. Build the HPL binary:


```
$> nmake -f Make.Windows_Intel64
```
5. Check that the built binary is available in the current directory.

NOTE

The Intel® Distribution for LINPACK Benchmark may contain additional optimizations compared to the reference Netlib HPL implementation.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[Contents of the Intel® Distribution for LINPACK Benchmark](#)

Configuring Parameters

The most significant parameters in `HPL.dat` are P , Q , NB , and N . Specify them as follows:

- P and Q - the number of rows and columns in the process grid, respectively.
 $P*Q$ must be the number of MPI processes that HPL is using.
 Choose $P \leq Q$.
- NB - the block size of the data distribution.

The table below shows recommended values of NB for different Intel® processors:

Processor	NB
Intel® Xeon® Processor X56*/E56*/E7-*/E7*/X7* (codenamed Nehalem or Westmere)	256
Intel Xeon Processor E26*/E26* v2 (codenamed Sandy Bridge or Ivy Bridge)	256
Intel Xeon Processor E26* v3/E26* v4 (codenamed Haswell or Broadwell)	192
Intel® Core™ i3/i5/i7-6* Processor (codenamed Skylake Client)	192
Intel® Xeon Phi™ Processor 72* (codenamed Knights Landing)	336

Processor	NB
Intel Xeon Processor supporting Intel® Advanced Vector Extensions 512 (Intel® AVX-512) instructions (codenamed Skylake Server)	384

- N - the problem size:
 - For homogeneous runs, choose N divisible by $NB * LCM(P, Q)$, where LCM is the least common multiple of the two numbers.
 - For heterogeneous runs, see [Heterogeneous Support in the Intel® Distribution for LINPACK* Benchmark](#) for how to choose N .

NOTE
 Increasing N usually increases performance, but the size of N is bounded by memory. In general, you can compute the memory required to store the matrix (which does not count internal buffers) as $8 * N * N / (P * Q)$ bytes, where N is the problem size and P and Q are the process grids in `HPL.dat`. A general rule of thumb is to choose a problem size that fills 80% of memory.

Ease-of-use Command-line Parameters

The Intel® Distribution for LINPACK* Benchmark supports command-line parameters for HPL that help you to avoid making small changes in the `HPL.dat` input file every time you do a new run.

Placeholders in this command line illustrate these parameters:

```
xhpl.exe -n <problem size> -m <memory size in Mbytes> -b <block size> -p <grid row dimn> -q <grid column dimn>
```

You can also use command-line parameters with the sample `runme` script. For example:

```
runme_intel64_dynamic.bat -m <memory size in Mbytes> -b <block size> -p <grid row dimn> -q <grid column dimn>
```

For more command-line parameters, see [Heterogeneous Support in the Intel® Distribution for LINPACK Benchmark](#).

If you want to run for $N=10000$ on a 1x3 grid, execute this command, provided that the other parameters in `HPL.dat` and the script are correct:

```
runme_intel64_dynamic.bat -n 10000 -p 1 -q 3
```

By using the `m` parameter you can scale by the memory size instead of the problem size. The `m` parameter only refers to the size of the matrix storage. Therefore, to use matrices that fit in 50000 Mbytes with $NB=256$ on 16 nodes, adjust the script to set the total number of MPI processes to 16 and execute this command:

```
runme_intel64_dynamic.bat -m 50000 -b 256 -p 4 -q 4
```

Running the Intel® Distribution for LINPACK* Benchmark

To run the Intel® Distribution for LINPACK Benchmark on multiple nodes or on one node with multiple MPI processes, you need to use MPI and either modify `HPL.dat` or use [Ease-of-use Command-line Parameters](#). The following example describes how to run the dynamically-linked prebuilt Intel® Distribution for LINPACK Benchmark binary using the script provided. To run other binaries, adjust the steps accordingly.

1. Load the necessary environment variables for the Intel MPI Library and Intel® compiler:


```
<compiler_directory>\env\vars.bat
<mpi_directory>\env\vars.bat
```
2. In `HPL.dat`, set the problem size N to 10000. Because this setting is for a test run, the problem size should be small.

- For better performance, enable non-uniform memory access (NUMA) on your system and configure to run an MPI process for each NUMA socket as explained below.

NOTE

High-bandwidth Multi-Channel Dynamic Random Access Memory (MCDRAM) on the second-generation Intel® Xeon® Phi processors may appear to be a NUMA node. However, because there are no CPUs on this node, do not run an MPI process for it.

- Refer to your BIOS settings to enable NUMA on your system.
- Set the following variables at the top of the `runme_intel64_dynamic.bat` script according to your cluster configuration:

`MPI_PROC_NUM` The total number of MPI processes.

`MPI_PER_NODE` The number of MPI processes per each cluster node.

- In the `HPL.dat` file, set the parameters `Ps` and `Qs` so that $Ps * Qs$ equals the number of MPI processes. For example, for 2 processes, set `Ps` to 1 and `Qs` to 2. Alternatively, leave the `HPL.dat` file as is and launch with `-p` and `-q` command-line parameters.

- Execute `runme_intel64_dynamic.bat` script:

```
runme_intel64_dynamic.bat
```

- Rerun the test increasing the size of the problem until the matrix size uses about 80% of the available memory. To do this, either modify `Ns` in line 6 of `HPL.dat` or use the `-n` command-line parameter:

- For 16 GB: 40000 `Ns`
- For 32 GB: 56000 `Ns`
- For 64 GB: 83000 `Ns`

See Also

[Notational Conventions](#)

[Building the Intel® Distribution for LINPACK Benchmark for a Customized MPI Implementation](#)

[Building the Netlib HPL from Source Code](#)

Heterogeneous Support in the Intel® Distribution for LINPACK* Benchmark

Intel® Distribution for LINPACK Benchmark achieves heterogeneous support by distributing the matrix data unequally between the nodes. The heterogeneous factor command-line parameter `f` controls the amount of work to be assigned to the more powerful nodes, while the command-line parameter `c` controls the number of process columns for the faster nodes:

```
xhpl.exe -n <problem size> -b <block size> -p <grid row dimn> -q <grid column dimn> -f  
<heterogeneous factor> -c <number of faster processor columns>
```

If the heterogeneous factor is 2.5, roughly 2.5 times the work will be put on the more powerful nodes. The more work you put on the more powerful nodes, the more memory you might be wasting on the other nodes if all nodes have equal amount of memory. If your cluster includes many different types of nodes, you may need multiple heterogeneous factors.

Let P be the number of rows and Q the number of columns in your processor grid ($P \times Q$). The work must be *homogeneous* within each processor column because vertical operations, such as pivoting or panel factorization, are synchronizing operations. When there are two different types of nodes, use MPI to process all the faster nodes first and make sure the "PMAP process mapping" (line 9) of `HPL.dat` is set to 1 for Column-major mapping. Because all the nodes must be the same within a process column, the number of faster nodes must always be a multiple of P , and you can specify the faster nodes by setting the number of process columns C for the faster nodes with the `c` command-line parameter. The `-f 1.0 -c 0` setting corresponds to the default homogeneous behavior.

To understand how to choose the problem size N for a heterogeneous run, first consider a homogeneous system, where you might choose N as follows:

$$N \sim \sqrt{\text{Memory Utilization} * P * Q * \text{Memory Size in Bytes} / 8}$$

Memory Utilization is usually around 0.8 for homogeneous Intel® Xeon® processor systems. On a heterogeneous system, you may apply a different formula for N for each set of nodes that are the same and then choose the minimum N over all sets. Suppose you have a cluster with only one heterogeneous factor F and the number of processor columns (out of the total Q) in the group with that heterogeneous factor equal to C . That group contains $P \cdot C$ nodes. First compute the sum of the parts: $S = F \cdot P \cdot C + P \cdot (Q - C)$. Note that on a homogeneous system $S = P \cdot Q$, $F = 1$, and $C = Q$. Take N as

$$N \sim \sqrt{\text{Memory Utilization} * P * Q * ((F * P * C) / S) * \text{Memory Size in Bytes} / 8}$$

or simply scale down the value of N for the homogeneous system by $\sqrt{(F * P * C) / S}$.

Example

Suppose the cluster has 100 nodes each having 64 GB of memory, and 20 of the nodes are 2.7 times as powerful as the other 80. Run one MPI process per node for a total of 100 MPI processes. Assume a square processor grid $P = Q = 10$, which conveniently divides up the faster nodes evenly. Normally, the HPL documentation recommends choosing a matrix size that consumes 80 percent of available memory. If N is the size of the matrix, the matrix consumes $8N^2 / (P \cdot Q)$ bytes. So a homogeneous run might look like:

```
xhpl.exe -n 820000 -b 256 -p 10 -q 10
```

If you redistribute the matrix and run the heterogeneous Intel® Distribution for LINPACK Benchmark, you can take advantage of the faster nodes. But because some of the nodes will contain 2.7 times as much data as the other nodes, you must shrink the problem size (unless the faster nodes also happen to have 2.7 times as much memory). Instead of $0.8 \cdot 64\text{GB} \cdot 100$ total memory size, we have only $0.8 \cdot 64\text{GB} \cdot 20 + 0.8 \cdot 64\text{GB} / 2.7 \cdot 80$ total memory size, which is less than half the original space. So the problem size in this case would be 526000. Because $P = 10$ and there are 20 faster nodes, two processor columns are faster. If you arrange MPI to send these nodes first to the application, the command line looks like:

```
xhpl.exe -n 526000 -b 1024 -p 10 -q 10 -f 2.7 -c 2
```

The `m` parameter may be misleading for heterogeneous calculations because it calculates the problem size assuming all the nodes have the same amount of data.

Warning

The number of faster nodes must be $C \cdot P$. If the number of faster nodes is not divisible by P , you might not be able to take advantage of the extra performance potential by giving the faster nodes extra work.

While it suffices to simply provide `f` and `c` command-line parameters if you need only one heterogeneous factor, you must add lines to the `HPL.dat` input to support multiple heterogeneous factors. For the above example (two processor columns have nodes that are 2.7 times faster), instead of passing `f` and `c` command-line parameters you can modify the `HPL.dat` input file by adding these two lines to the end:

```
1      number of heterogeneous factors
0 1 2.7  [start_column, stop_column, heterogeneous factor for that range]
```

NOTE

Numbering of processor columns starts at 0. The start and stopping numbers must be between 0 and $Q - 1$ (inclusive).

If instead there are three different types of nodes in a cluster and you need at least two heterogeneous factors, change the number in the first row above from 1 to 2 and follow that line with two lines specifying the start column, stopping column, and heterogeneous factor.

When choosing parameters for heterogeneous support in `HPL.dat`, primarily focus on the most powerful nodes. The larger the heterogeneous factor, the more balanced the cluster may be from a performance viewpoint, but the more imbalanced from a memory viewpoint. At some point, further performance balancing

might affect the memory too much. If this is the case, try to reduce any changes done for the faster nodes (such as in block sizes). Experiment with values in `HPL.dat` carefully because wrong values may greatly hinder performance.

When tuning on a heterogeneous cluster, do not immediately attempt a heterogeneous run, but do the following:

1. Break the cluster down into multiple homogeneous clusters.
2. Make heterogeneous adjustments for performance balancing. For instance, if you have two different sets of nodes where one is three times as powerful as the other, it must do three times the work.
3. Figure out the approximate size of the problem (per node) that you can run on each piece.
4. Do some homogeneous runs with those problem sizes per node and the final block size needed for the heterogeneous run and find the best parameters.
5. Use these parameters for an initial heterogeneous run.

Environment Variables

The table below lists Intel® oneAPI Math Kernel Library environment variables to control runs of the Intel Distribution for LINPACK Benchmark.

Environment Variable	Description	Value
<code>HPL_LARGEPAGE</code>	Defines the memory mapping to be used for the Intel® Xeon® processor.	0 or 1: • 0 - normal memory mapping, default. • 1 - memory mapping with large pages (2 MB per page mapping). It may increase performance.
<code>HPL_LOG</code>	Controls the level of detail for the HPL output.	An integer ranging from 0 to 2: • 0 - no log is displayed. • 1 - only one root node displays a log, exactly the same as the <code>ASYOUGO</code> option provides. • 2 - the most detailed log is displayed. All <i>P</i> root nodes in the processor column that owns the current column block display a log.
<code>HPL_HOST_CORE</code> , <code>HPL_HOST_NODE</code>	Specifies cores or Non-Uniform Memory Access (NUMA) nodes to be used. <code>HPL_HOST_NODE</code> requires NUMA mode to be enabled. You can check whether it is enabled by the <code>numactl --hardware</code> command. The default behavior is auto-detection of the core or NUMA node.	A list of integers ranging from 0 to the largest number of a core or NUMA node in the cluster and separated as explained in example 3 .
<code>HPL_SWAPWIDTH</code>	Specifies width for each swap operation.	16 or 24. The default is 24.

You can set Intel Distribution for LINPACK Benchmark environment variables using the `PMI_RANK` and `PMI_SIZE` environment variables of the Intel MPI library, and you can create a shell script to automate the process.

Examples of Environment Settings

#	Settings	Behavior of the Intel Distribution for LINPACK Benchmark
1	Nothing specified	All Intel® Xeon® processors in the cluster are used.
2	HPL_MIC_DEVICE=0,2 HPL_HOST_CORE=1-3,8-10	Intel® Xeon® processor cores 1,2,3,8,9, and 10 are used.
3	HPL_HOST_NODE=1	Only Intel® Xeon® processor cores on NUMA node 1 are used.

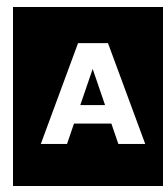
Improving Performance of Your Cluster

To improve cluster performance, follow these steps, provided all required software is installed on each node:

1. Reboot all nodes.
2. Ensure all nodes are in identical conditions and no zombie processes are left running from prior HPL runs. To do this, run single-node Stream and Intel® Distribution for LINPACK Benchmark on every node. Ensure results are within 10% of each other (problem size must be large enough depending on memory size and CPU speed). Investigate nodes with low performance for hardware/software problems.
3. Check that your cluster interconnects are working. Run a test over the complete cluster using an MPI test for bandwidth and latency, such as one found in the Intel® MPI Benchmarks package.
4. Run an Intel® Distribution for LINPACK Benchmark on pairs of two or four nodes and ensure results are within 10% of each other. The problem size must be large enough depending on the memory size and CPU speed.
5. Run a small problem size over the complete cluster to ensure correctness.
6. Increase the problem size and run the real test load.
7. In case of problems go back to step 2.

Before making a heterogeneous run, always run its homogeneous equivalent first.

Intel® oneAPI Math Kernel Library Language Interfaces Support



See Also

[Mixed-language Programming with Intel® oneAPI Math Kernel Library](#)

Language Interfaces Support, by Function Domain

The following table shows language interfaces that provides for each function domain. However, Intel® oneAPI Math Kernel Library routines can be called from other languages using mixed-language programming. See [Mixed-language Programming with the Intel Math Kernel Library](#) for an example of how to call Fortran routines from C/C++.

Function Domain	Fortran interface	C/C++ interface
Basic Linear Algebra Subprograms (BLAS)	Yes	through CBLAS
BLAS-like extension transposition routines	Yes	Yes
Sparse BLAS Level 1	Yes	through CBLAS
Sparse BLAS Level 2 and 3	Yes	Yes
LAPACK routines for solving systems of linear equations	Yes	Yes
LAPACK routines for solving least-squares problems, eigenvalue and singular value problems, and Sylvester's equations	Yes	Yes
Auxiliary and utility LAPACK routines	Yes	Yes
Parallel Basic Linear Algebra Subprograms (PBLAS)	Yes	
ScaLAPACK	Yes	†
Direct Sparse Solvers/ Intel® oneAPI Math Kernel Library PARDISO, a direct sparse solver based on Parallel Direct Sparse Solver (PARDISO*)	Yes	Yes
Parallel Direct Sparse Solvers for Clusters	Yes	Yes
Other Direct and Iterative Sparse Solver routines	Yes	Yes
Vector Mathematics (VM)	Yes	Yes
Vector Statistics (VS)	Yes	Yes
Fast Fourier Transforms (FFT)	Yes	Yes
Cluster FFT	Yes	Yes
Trigonometric Transforms	Yes	Yes
Fast Poisson, Laplace, and Helmholtz Solver (Poisson Library)	Yes	Yes
Optimization (Trust-Region) Solver	Yes	Yes

Function Domain	Fortran interface	C/C++ interface
Data Fitting	Yes	Yes
Extended Eigensolver	Yes	Yes
Support functions (including memory allocation)	Yes	Yes

[†] Supported using a mixed language programming call. See [Include Files](#) for the respective header file.

Include Files

The table below lists Intel® oneAPI Math Kernel Library include files.

Function Domain/ Purpose	Fortran Include Files	C/C++ Include Files
All function domains	<code>mkl.fi</code>	<code>mkl.h</code>
BLACS		<code>mkl_blacs.h^{††}</code>
BLAS	<code>blas.f90</code> <code>mkl_blas.fi[†]</code>	<code>mkl_blas.h[†]</code>
BLAS-like Extension Transposition Routines	<code>mkl_trans.fi[†]</code>	<code>mkl_trans.h[†]</code>
CBLAS Interface to BLAS		<code>mkl_cblas.h[†]</code>
Sparse BLAS	<code>mkl_spgblas.fi[†]</code>	<code>mkl_spgblas.h[†]</code>
LAPACK	<code>lapack.f90</code> <code>mkl_lapack.fi[†]</code>	<code>mkl_lapack.h[†]</code>
C Interface to LAPACK		<code>mkl_lapacke.h[†]</code>
PBLAS		<code>mkl_pblas.h^{††}</code>
ScaLAPACK		<code>mkl_scalapack.h^{††}</code>
Intel® oneAPI Math Kernel Library PARDISO	<code>mkl_pardiso.f90</code> <code>mkl_pardiso.fi[†]</code>	<code>mkl_pardiso.h[†]</code>
Parallel Direct Sparse Solvers for Clusters	<code>mkl_cluster_sparse_solver.f90</code>	<code>mkl_cluster_sparse_solver.h[†]</code>
Direct Sparse Solver (DSS)	<code>mkl_dss.f90</code> <code>mkl_dss.fi[†]</code>	<code>mkl_dss.h[†]</code>
RCI Iterative Solvers		
ILU Factorization	<code>mkl_rci.f90</code> <code>mkl_rci.fi[†]</code>	<code>mkl_rci.h[†]</code>
Optimization Solver	<code>mkl_rci.f90</code> <code>mkl_rci.fi[†]</code>	<code>mkl_rci.h[†]</code>
Vector Mathematics	<code>mkl_vml.90</code> <code>mkl_vml.fi[†]</code>	<code>mkl_vml.h[†]</code>
Vector Statistics	<code>mkl_vsl.f90</code> <code>mkl_vsl.fi[†]</code>	<code>mkl_vsl.h[†]</code>

Function Domain/ Purpose	Fortran Include Files	C/C++ Include Files
Fast Fourier Transforms	<code>mkl_dfti.f90</code>	<code>mkl_dfti.h</code> [†]
Cluster Fast Fourier Transforms	<code>mkl_cdft.f90</code>	<code>mkl_cdft.h</code> ^{††}
Partial Differential Equations Support		
Trigonometric Transforms	<code>mkl_trig_transforms.f90</code>	<code>mkl_trig_transform.h</code> [†]
Poisson Solvers	<code>mkl_poisson.f90</code>	<code>mkl_poisson.h</code> [†]
Data Fitting	<code>mkl_df.f90</code>	<code>mkl_df.h</code> [†]
Extended Eigensolver	<code>mkl_solvers_ee.fi</code> [†]	<code>mkl_solvers_ee.h</code> [†]
Support functions	<code>mkl_service.f90</code> <code>mkl_service.fi</code> [†]	<code>mkl_service.h</code> [†]
Declarations for replacing memory allocation functions. See Redefining Memory Functions for details.		<code>i_malloc.h</code>
Auxiliary macros to determine the version of Intel® oneAPI Math Kernel Library at compile time.	<code>mkl_version</code>	<code>mkl_version</code> [†]

[†] You can use the `mkl.fi` include file in your code instead.

[†] You can include the `mkl.h` header file in your code instead.

^{††} Also include the `mkl.h` header file in your code.

See Also

[Language Interfaces Support, by Function Domain](#)



Support for Third-Party Interfaces

FFTW Interface Support

offers two collections of wrappers for the FFTW interface (www.fftw.org). The wrappers are the superstructure of FFTW to be used for calling the Intel® oneAPI Math Kernel Library Fourier transform functions. These collections correspond to the FFTW versions 2.x and 3.x and the Intel® oneAPI Math Kernel Library versions 7.0 and later.

These wrappers enable using Intel® oneAPI Math Kernel Library Fourier transforms to improve the performance of programs that use FFTW without changing the program source code. See the "*FFTW Interface to Intel® oneAPI Math Kernel Library*" appendix in the *Intel® oneAPI Math Kernel Library Developer Reference* for details on the use of the wrappers.

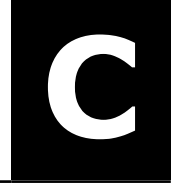
Important

For ease of use, the FFTW3 interface is also integrated in Intel® oneAPI Math Kernel Library.

Caution

The FFTW2 and FFTW3 interfaces are not compatible with each other. Avoid linking to both of them. If you must do so, first modify the wrapper source code for FFTW2:

1. Change every instance of `fftw_destroy_plan` in the `fftw2xc` interface to `fftw2_destroy_plan`.
 2. Change all the corresponding file names accordingly.
 3. Rebuild the pertinent libraries.
-



Directory Structure in Detail

Tables in this section show contents of the architecture-specific directories.

Product and Performance Information

Performance varies by use, configuration and other factors. Learn more at www.Intel.com/PerformanceIndex.

Notice revision #20201201

See Also

[High-level Directory Structure](#)

[Using Language-Specific Interfaces with Intel® oneAPI Math Kernel Library](#)

[Intel® oneAPI Math Kernel Library Benchmarks](#)

Detailed Structure of the IA-32 Architecture Directories

Static Libraries in the `lib\ia32` Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
Interface Layer			
mkl_intel_c.lib	cdecl interface library		
mkl_blas95.lib	Fortran 95 interface library for BLAS. Supports the Intel® Fortran compiler	Fortran 95 interfaces for BLAS and LAPACK	Yes
mkl_lapack95.lib	Fortran 95 interface library for LAPACK. Supports the Intel® Fortran compiler	Fortran 95 interfaces for BLAS and LAPACK	Yes
Threading Layer			
mkl_intel_thread.lib	OpenMP threading library for the Intel compilers		

File	Contents	Optional Component	
		Name	Installed by Default
mkl_tbb_thread.lib	Intel® Threading Building Blocks (Intel® TBB) threading library for the Intel compilers	Intel TBB threading support	Yes
mkl_sequential.lib	Sequential library		
Computational Layer			
mkl_core.lib	Kernel library for IA-32 architecture		

Dynamic Libraries in the `lib\ia32` Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
mkl_rt.lib	Single Dynamic Library to be used for linking		
Interface Layer			
mkl_intel_c_dll.lib	cdecl interface library for dynamic linking		
Threading Layer			
mkl_intel_thread_dll.lib	OpenMP threading library for dynamic linking with the Intel compilers		
mkl_tbb_thread_dll.lib	Intel TBB threading library for the Intel compilers	Intel TBB threading support	Yes
mkl_sequential_dll.lib	Sequential library for dynamic linking		
Computational Layer			
mkl_core_dll.lib	Core library for dynamic linking		

Contents of the `redist\ia32` Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
mk1_rt.dll	Single Dynamic Library		
Threading Layer			
mk1_intel_thread.dll	Dynamic OpenMP threading library for the Intel compilers		
mk1_tbb_thread.dll	Dynamic Intel TBB threading library for the Intel compilers	Intel TBB threading support	Yes
mk1_sequential.dll	Dynamic sequential library		
Computational Layer			
mk1_core.dll	Core library containing processor-independent code and a dispatcher for dynamic loading of processor-specific code		
mk1_p4.dll	Pentium® 4 processor kernel		
mk1_p4m.dll	Kernel library for Intel® Supplemental Streaming SIMD Extensions 3 (Intel® SSE3) enabled processors		
mk1_p4m3.dll	Kernel library for Intel® Streaming SIMD Extensions 4.2 (Intel® SSE4.2) enabled processors		
mk1_avx.dll	Kernel library for Intel® Advanced Vector Extensions (Intel® AVX) enabled processors		
mk1_avx2.dll	Kernel library for Intel® Advanced Vector Extensions 2 (Intel® AVX2) enabled processors		
mk1_avx512.dll	Kernel library for Intel® Advanced Vector Extensions 512 (Intel® AVX-512) enabled processors		

File	Contents	Optional Component	
		Name	Installed by Default
mk1_vml_p4.dll	Vector Mathematics (VM)/Vector Statistics (VS)/Data Fitting (DF) part of Pentium® 4 processor kernel		
mk1_vml_p4m.dll	VM/VS/DF for Intel® SSSE3 enabled processors		
mk1_vml_p4m2.dll	VM/VS/DF for 45nm Hi-k Intel® Core™2 and Intel Xeon® processor families		
mk1_vml_p4m3.dll	VM/VS/DF for Intel® SSE4.2 enabled processors		
mk1_vml_avx.dll	VM/VS/DF optimized for Intel® AVX enabled processors		
mk1_vml_avx2.dll	VM/VS/DF optimized for Intel® AVX2 enabled processors		
mk1_vml_avx512.dll	VM/VS/DF optimized for Intel® AVX-512 enabled processors		
mk1_vml_ia.dll	VM/VS/DF default kernel for newer Intel® architecture processors		
libmk1_vml_cmpt.dll	VM/VS/DF library for conditional numerical reproducibility		
libimalloc.dll	Dynamic library to support renaming of memory functions		
Message Catalogs			
1033\mk1_msg.dll	Catalog of messages in English		
1041\mk1_msg.dll	Catalog of Intel® oneAPI Math Kernel Library messages in Japanese. Available only if Intel® oneAPI Math Kernel Library provides Japanese localization. Please see the Release Notes for this information.		

Detailed Structure of the Intel® 64 Architecture Directories

Static Libraries in the `lib\intel64` Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
<code>mkl_sycl.lib</code>	DPC++ library for the Intel DPC++ compilers		
<code>mkl_sycl_d.lib</code>	DPC++ library for the Intel DPC++ compiler with debug runtime (/MDd)		
Interface Layer			
<code>mkl_intel_lp64.lib</code>	LP64 interface library for the Intel compilers		
<code>mkl_intel_ilp64.lib</code>	ILP64 interface library for the Intel compilers		
<code>mkl_blas95_lp64.lib</code>	Fortran 95 interface library for BLAS. Supports the Intel® Fortran compiler and LP64 interface	Fortran 95 interfaces for BLAS and LAPACK	Yes
<code>mkl_blas95_ilp64.lib</code>	Fortran 95 interface library for BLAS. Supports the Intel® Fortran compiler and ILP64 interface	Fortran 95 interfaces for BLAS and LAPACK	Yes
<code>mkl_lapack95_lp64.lib</code>	Fortran 95 interface library for LAPACK. Supports the Intel® Fortran compiler and LP64 interface	Fortran 95 interfaces for BLAS and LAPACK	Yes
<code>mkl_lapack95_ilp64.lib</code>	Fortran 95 interface library for LAPACK. Supports the Intel® Fortran compiler and ILP64 interface	Fortran 95 interfaces for BLAS and LAPACK	Yes
Threading Layer			
<code>mkl_intel_thread.lib</code>	OpenMP threading library for the Intel compilers		

File	Contents	Optional Component	
		Name	Installed by Default
mkl_tbb_thread.lib	Intel® Threading Building Blocks (Intel® TBB) threading library for the Intel compilers	Intel TBB threading support	Yes
mkl_tbb_threadd.lib	Intel® Threading Building Blocks (Intel® TBB) threading library for the Intel compilers compatible with mkl_sycl.lib	Intel TBB threading support	Yes
mkl_pgi_thread.lib	OpenMP threading library for the PGI* compiler	PGI* Compiler support	
mkl_sequential.lib	Sequential library		
Computational Layer			
mkl_core.lib	Kernel library for the Intel® 64 architecture		
Cluster Libraries			
mkl_scalapack_lp64.lib	ScaLAPACK routine library supporting the LP64 interface	Cluster support	
mkl_scalapack_ilp64.lib	ScaLAPACK routine library supporting the ILP64 interface	Cluster support	
mkl_cdft_core.lib	Cluster version of FFTs	Cluster support	
mkl_blacs_intelmpi_lp64.lib	LP64 version of BLACS routines supporting Intel® MPI Library	Cluster support	
mkl_blacs_intelmpi_ilp64.lib	ILP64 version of BLACS routines supporting Intel MPI Library	Cluster support	
mkl_blacs_msmpi_lp64.lib	LP64 version of BLACS routines supporting Microsoft* MPI	Cluster support	
mkl_blacs_msmpi_ilp64.lib	ILP64 version of BLACS routines supporting Microsoft* MPI	Cluster support	

Dynamic Libraries in the lib\intel64 Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
mkl_sycl_dll.lib	DPC++ library for the Intel DPC++ compilers		
mkl_sycl_d_dll.lib	DPC++ library for the Intel DPC++ compiler with debug runtime (/MDd)		
mkl_rt.lib	Single Dynamic Library to be used for linking		
Interface Layer			
mkl_intel_lp64_dll.lib	LP64 interface library for dynamic linking with the Intel compilers		
mkl_intel_ilp64_dll.lib	ILP64 interface library for dynamic linking with the Intel compilers		
Threading Layer			
mkl_intel_thread_dll.lib	OpenMP threading library for dynamic linking with the Intel compilers		
mkl_tbb_thread_dll.lib	Intel TBB threading library for the Intel compilers	Intel TBB threading support	Yes
mkl_tbb_threadd_dll.lib	Intel® Threading Building Blocks (Intel® TBB) threading library for the Intel compilers compatible with mkl_sycl_d_dll.lib	Intel® TBB threading support	Yes
mkl_pgi_thread_dll.lib	OpenMP threading library for dynamic linking with the PGI* compiler	PGI* Compiler support	
mkl_sequential_dll.lib	Sequential library for dynamic linking		
Computational Layer			
mkl_core_dll.lib	Core library for dynamic linking		
Cluster Libraries			
mkl_scalapack_lp64_dll.lib	ScaLAPACK routine library for dynamic linking supporting the LP64 interface	Cluster support	

File	Contents	Optional Component	
		Name	Installed by Default
mkl_scalapack_ilp64_dll.lib	ScaLAPACK routine library for dynamic linking supporting the ILP64 interface	Cluster support	
mkl_cdft_core_dll.lib	Cluster FFT library for dynamic linking	Cluster support	
mkl_blacs_lp64_dll.lib	LP64 version of BLACS interface library for dynamic linking	Cluster support	
mkl_blacs_ilp64_dll.lib	ILP64 version of BLACS interface library for dynamic linking	Cluster support	

Contents of the `redist\intel64` Directory

Some of the libraries in this directory are optional. However, some optional libraries are installed by default, while the rest are not. To get those libraries that are not installed by default, explicitly select the specified optional component during installation.

File	Contents	Optional Component	
		Name	Installed by Default
mkl_sycl.<version>.dll	DPC++ library for the Intel DPC++ compilers		
mkl_sycl_d.<version>.dll	DPC++ library for the Intel DPC++ compiler with debug runtime (/MDd)		
mkl_rt.<version>.dll	Single Dynamic Library		
Threading layer			
mkl_tbb_thread.<version>.dll	Dynamic Intel® Threading Building Blocks (Intel® TBB) threading library for the Intel compilers compatible with mkl_sycl_d.dll		
mkl_intel_thread.<version>.dll	Dynamic OpenMP threading library for the Intel compilers		
mkl_tbb_thread.<version>.dll	Dynamic Intel TBB threading library for the Intel compilers	Intel TBB threading support	Yes

File	Contents	Optional Component	
		Name	Installed by Default
mkl_pgi_thread.<version>.dll	Dynamic OpenMP threading library for the PGI* compiler	PGI* compiler support	
mkl_sequential.<version>.dll	Dynamic sequential library		
Computational layer			
mkl_core.<version>.dll	Core library containing processor-independent code and a dispatcher for dynamic loading of processor-specific code		
mkl_def.<version>.dll	Default kernel for the Intel® 64 architecture		
mkl_mc.<version>.dll	Kernel library for Intel® Supplemental Streaming SIMD Extensions 3 (Intel® SSE3) enabled processors		
mkl_mc3.<version>.dll	Kernel library for Intel® Streaming SIMD Extensions 4.2 (Intel® SSE4.2) enabled processors		
mkl_avx.<version>.dll	Kernel library optimized for Intel® Advanced Vector Extensions (Intel® AVX) enabled processors		
mkl_avx2.<version>.dll	Kernel library optimized for Intel® Advanced Vector Extensions 2 (Intel® AVX2) enabled processors		
mkl_avx512.<version>.dll	Kernel library optimized for Intel® Advanced Vector Extensions 512 (Intel® AVX-512) enabled processors		
mkl_vml_def.<version>.dll	Vector Mathematics (VM)/Vector Statistics (VS)/Data Fitting (DF) part of default kernel		
mkl_vml_mc.<version>.dll	VM/VS/DF for Intel® SSE3 enabled processors		

File	Contents	Optional Component	
		Name	Installed by Default
mk1_vml_mc2.<version>.dll	VM/VS/DF for 45nm Hi-k Intel® Core™2 and Intel Xeon® processor families		
mk1_vml_mc3.<version>.dll	VM/VS/DF for Intel® SSE4.2 enabled processors		
mk1_vml_avx.<version>.dll	VM/VS/DF optimized for Intel® AVX enabled processors		
mk1_vml_avx2.<version>.dll	VM/VS/DF optimized for Intel® AVX2 enabled processors		
mk1_vml_avx512.<version>.dll	VM/VS/DF optimized for Intel® AVX-512 enabled processors		
mk1_vml_cmpt.<version>.dll	VM/VS/DF library for conditional numerical reproducibility		
libimalloc.dll	Dynamic library to support renaming of memory functions		
Cluster Libraries			
mk1_scalapack_lp64.<version>.dll	ScaLAPACK routine library supporting the LP64 interface	Cluster support	
mk1_scalapack_ilp64.<version>.dll	ScaLAPACK routine library supporting the ILP64 interface	Cluster support	
mk1_cdft_core.<version>.dll	Cluster FFT dynamic library	Cluster support	
mk1_blacs_lp64.<version>.dll	LP64 version of BLACS routines	Cluster support	
mk1_blacs_ilp64.<version>.dll	ILP64 version of BLACS routines	Cluster support	
mk1_blacs_intelmpi_lp64.<version>.dll	LP64 version of BLACS routines for Intel® MPI Library	Cluster support	
mk1_blacs_intelmpi_ilp64.<version>.dll	ILP64 version of BLACS routines for Intel MPI Library	Cluster support	
mk1_blacs_msmpi_lp64.<version>.dll	LP64 version of BLACS routines for Microsoft* MPI	Cluster support	

File	Contents	Optional Component	
		Name	Installed by Default
mk1_blacs_msmapi_ilp64.<version>.dll	ILP64 version of BLACS routines for Microsoft* MPI	Cluster support	
Message Catalogs			
1033\mk1_msg.dll	Catalog of messages in English		

Index

A

- affinity mask 57
- aligning data, example 80
- architecture support 18

B

- BLAS
 - calling routines from C 67
 - Fortran 95 interface to 65
 - OpenMP* threaded routines 43
- building a custom DLL
 - in Visual Studio* IDE 38

C

- C interface to LAPACK, use of 67
- C, calling LAPACK, BLAS, CBLAS from 67
- C/C++, Intel(R) MKL complex types 68
- calling
 - BLAS functions from C 69
 - CBLAS interface from C 69
 - complex BLAS Level 1 function from C 69
 - complex BLAS Level 1 function from C++ 69
 - Fortran-style routines from C 67
- calling convention, cdecl 14
- CBLAS interface, use of 67
- Cluster FFT
 - environment variable for 99, 100
 - linking with 91
 - managing performance of 99, 100
- cluster software, Intel(R) MKL 90
- cluster software, linking with
 - commands 91
 - linking examples 95
- Cluster Sparse Solver, linking with 91
- code examples, use of 15
- coding
 - data alignment 80
 - techniques to improve performance 60
- compilation, Intel(R) MKL version-dependent 81
- compiler run-time libraries, linking with 34
- compiler support 14
- compiler-dependent function 66
- complex types in C and C++, Intel(R) MKL 68
- computation results, consistency 72
- computational libraries, linking with 33
- conditional compilation 81
- configuring
 - Intel(R) Visual Fortran 104
 - Microsoft Visual* C/C++ 103
- consistent results 72
- context-sensitive Help, for Intel(R) MKL, in Visual Studio* IDE 104
- conventions, notational 10
- custom DLL
 - building 34
 - composing list of functions 37

- custom DLL (*continued*)
 - specifying function names 38

D

- data alignment, example 80
- denormal number, performance 61
- direct call, to Intel(R) Math Kernel Library computational kernels 58
- directory structure
 - high-level 18
 - in-detail
- dispatch Intel(R) architectures, configure with an environment variable 101
- dispatch, new Intel(R) architectures, enable with an environment variable 101

E

- Enter index keyword 21
- environment variables
 - for threading control 50
 - setting for specific architecture and programming interface 14
 - to control dispatching for Intel(R) architectures 101
 - to control threading algorithm for ?gemm 53
 - to enable dispatching of new architectures 101
 - to manage behavior of function domains 98
 - to manage behavior of Intel(R) Math Kernel Library with 98
 - to manage performance of cluster FFT 99
- examples, linking
 - for cluster software 95
 - general 26

F

- FFT interface
 - OpenMP* threaded problems 43
- FFTW interface support 120
- Fortran 95 interface libraries 31
- function call information, enable printing 83

H

- header files, Intel(R) MKL 118
- heterogeneity
 - of Intel(R) Distribution for LINPACK* Benchmark 109
- heterogeneous cluster
 - support by Intel(R) Distribution for LINPACK* Benchmark for Clusters 113
- heterogeneous cores 58
- HT technology, configuration tip 56

I

- ILP64 programming, support for 29
- improve performance, for matrices of small sizes 58

include files, Intel(R) MKL 118
 information, for function call , enable printing 83
 installation, checking 12, 13
 Intel(R) Distribution for LINPACK* Benchmark
 heterogeneity of 109
 Intel(R) Distribution for LINPACK* Benchmark for Clusters
 heterogeneous support 113
 Intel(R) Hyper-Threading Technology, configuration tip 56
 Intel(R) Visual* Fortran project, linking with Intel(R)
 MKL 22
 Intel® Threading Building Blocks, functions threaded with 45
 IntelliSense*, with Intel(R) MKL, in Visual Studio* IDE 104
 interface
 Fortran 95, libraries 31
 LP64 and ILP64, use of 29
 interface libraries and modules, Intel(R) MKL 64
 interface libraries, linking with 29

K

kernel, in Intel(R) Math Kernel Library, direct call to 58

L

language interfaces support 117
 language-specific interfaces
 interface libraries and modules 64
 LAPACK
 C interface to, use of 67
 calling routines from C 67
 Fortran 95 interface to 65
 OpenMP* threaded routines 43
 performance of packed routines 60
 layers, Intel(R) MKL structure 19
 libraries to link with
 computational 33
 interface 29
 run-time 34
 system libraries 34
 threading 31
 link tool, command line 24
 linking
 Intel(R) Visual* Fortran project with Intel(R) MKL 22
 Microsoft Visual* C/C++ project with Intel(R) MKL 22
 linking examples
 cluster software 95
 general 26
 linking with
 compiler run-time libraries 34
 computational libraries 33
 interface libraries 29
 system libraries 34
 threading libraries 31
 linking, quick start 21, 39
 linking, Web-based advisor 24
 LINPACK benchmark 107

M

memory functions, redefining 62
 memory management 62
 memory renaming 62
 message-passing interface
 custom, usage 94
 Intel(R) Math Kernel Library interaction with 93
 support 90

Microsoft Visual* C/C++ project, linking with Intel(R)
 MKL 22
 mixed-language programming 66
 module, Fortran 95 65
 MPI
 custom, usage 94
 Intel(R) Math Kernel Library interaction with 93
 support 90
 multi-core performance 57

N

notational conventions 10
 number of threads
 changing at run time 48
 changing with OpenMP* environment variable 47
 Intel(R) MKL choice, particular cases 51
 setting for cluster 92
 techniques to set 47
 numerically reproducible results 72

O

OpenMP* threaded functions 43
 OpenMP* threaded problems 43

P

parallel performance 46
 parallelism, of Intel(R) MKL 43
 performance
 heterogeneous cores 58
 multi-core 57
 with denormals 61
 with subnormals 61
 performance improvement, for matrices of small sizes 58
 performance, of Intel(R) MKL, improve on specific
 processors 61

R

results, consistent, obtaining 72
 results, numerically reproducible, obtaining 72

S

ScaLAPACK, linking with 91
 SDL 23, 28
 Single Dynamic Library 23, 28
 structure
 high-level 18
 in-detail
 model 19
 support, technical 7
 supported architectures 18
 system libraries, linking with 34

T

technical support 7
 thread safety, of Intel(R) MKL 43
 threaded functions, with Intel® Threading Building Blocks 45
 threading control, Intel(R) MKL-specific 50
 threading libraries, linking with 31

U

unstable output, getting rid of 72

V

Vector Mathematics

 default mode, setting with environment variable 98

 environment variable to set default mode 98

verbose mode, of Intel(R) MKL 83

Visual Studio* IDE

 IntelliSense*, with Intel(R) MKL 104

 using Intel(R) MKL context-sensitive Help in 104