

AMCLIB User's Guide

ARM® Cortex® M4F



Contents

Chapter 1 Library	4
1.1 Introduction	4
1.1.1 Overview	4
1.1.2 Data types	4
1.1.3 API definition	4
1.1.4 Supported compilers	5
1.1.5 Library configuration	5
1.1.6 Special issues	5
1.2 Library integration into project (MCUXpresso IDE)	6
1.3 Library integration into project (Keil µVision)	8
1.4 Library integration into project (IAR Embedded Workbench)	15
 Chapter 2 Algorithms in detail	 22
2.1 AMCLIB_ACIMCtrlMTPA	22
2.1.1 Available versions	23
2.1.2 AMCLIB_ACIM_CTRL_MTPA_T_FLT type description	23
2.1.3 Declaration	23
2.1.4 Function use	24
2.2 AMCLIB_ACIMRotFluxObsrv	24
2.2.1 Available versions	26
2.2.2 AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT type description	27
2.2.3 Declaration	28
2.2.4 Function use	29
2.3 AMCLIB_ACIMSpeedMRAS	29
2.3.1 Available versions	30
2.3.2 AMCLIB_ACIMSpeedMRAS_T_FLT type description	31
2.3.3 Declaration	32
2.3.4 Function use	32
2.4 AMCLIB_AngleTrackObsrv	33
2.4.1 Available versions	35
2.4.2 AMCLIB_ANGLE_TRACK_OBSRV_T_F32	36
2.4.3 AMCLIB_ANGLE_TRACK_OBSRV_T_FLT	37
2.4.4 Declaration	38
2.4.5 Function use	38
2.5 AMCLIB_CtrlFluxWkng	39
2.5.1 Available versions	41
2.5.2 AMCLIB_CTRL_FLUX_WKNG_T_A32	42
2.5.3 AMCLIB_CTRL_FLUX_WKNG_T_FLT	43
2.5.4 Declaration	43
2.5.5 Function use	44
2.6 AMCLIB_PMSMBemfObsrvAB	45
2.6.1 Available versions	48
2.6.2 AMCLIB_BEMF_OBSRV_AB_T_A32 type description	49
2.6.3 AMCLIB_BEMF_OBSRV_AB_T_FLT type description	50
2.6.4 Declaration	52
2.6.5 Function use	52
2.7 AMCLIB_PMSMBemfObsrvDQ	53
2.7.1 Available versions	56
2.7.2 AMCLIB_BEMF_OBSRV_DQ_T_A32 type description	57
2.7.3 AMCLIB_BEMF_OBSRV_DQ_T_FLT type description	59

2.7.4 Declaration.....	60
2.7.5 Function use.....	60
2.8 AMCLIB_TrackObsrv.....	61
2.8.1 Available versions.....	62
2.8.2 AMCLIB_TRACK_OBSRV_T_F32.....	63
2.8.3 AMCLIB_TRACK_OBSRV_T_FLT.....	64
2.8.4 Declaration.....	64
2.8.5 Function use.....	64

Appendix A Library types.....	66
A.1 bool_t.....	66
A.2 uint8_t.....	66
A.3 uint16_t.....	67
A.4 uint32_t.....	68
A.5 int8_t.....	68
A.6 int16_t.....	69
A.7 int32_t.....	69
A.8 frac8_t.....	70
A.9 frac16_t.....	71
A.10 frac32_t.....	71
A.11 acc16_t.....	72
A.12 acc32_t.....	73
A.13 float_t.....	73
A.14 GMCLIB_3COOR_T_F16.....	76
A.15 GMCLIB_3COOR_T_FLT.....	76
A.16 GMCLIB_2COOR_ALBE_T_F16.....	77
A.17 GMCLIB_2COOR_ALBE_T_FLT.....	77
A.18 GMCLIB_2COOR_DQ_T_F16.....	78
A.19 GMCLIB_2COOR_DQ_T_F32.....	78
A.20 GMCLIB_2COOR_DQ_T_FLT.....	78
A.21 GMCLIB_2COOR_SINCOS_T_F16.....	79
A.22 GMCLIB_2COOR_SINCOS_T_FLT.....	79
A.23 FALSE.....	80
A.24 TRUE.....	80
A.25 FRAC8.....	80
A.26 FRAC16.....	81
A.27 FRAC32.....	81
A.28 ACC16.....	81
A.29 ACC32.....	82

Chapter 1

Library

1.1 Introduction

1.1.1 Overview

This user's guide describes the Advanced Motor Control Library (AMCLIB) for the family of ARM Cortex M4F core-based microcontrollers. This library contains optimized functions.

1.1.2 Data types

AMCLIB supports several data types: (un)signed integer, fractional, and accumulator, and floating point. The integer data types are useful for general-purpose computation; they are familiar to the MPU and MCU programmers. The fractional data types enable powerful numeric and digital-signal-processing algorithms to be implemented. The accumulator data type is a combination of both; that means it has the integer and fractional portions. The floating-point data types are capable of storing real numbers in wide dynamic ranges. The type is represented by binary digits and an exponent. The exponent allows scaling the numbers from extremely small to extremely big numbers. Because the exponent takes part of the type, the overall resolution of the number is reduced when compared to the fixed-point type of the same size.

The following list shows the integer types defined in the libraries:

- **Unsigned 16-bit integer**—<0 ; 65535> with the minimum resolution of 1
- **Signed 16-bit integer**—<-32768 ; 32767> with the minimum resolution of 1
- **Unsigned 32-bit integer**—<0 ; 4294967295> with the minimum resolution of 1
- **Signed 32-bit integer**—<-2147483648 ; 2147483647> with the minimum resolution of 1

The following list shows the fractional types defined in the libraries:

- **Fixed-point 16-bit fractional**—<-1 ; $1 - 2^{-15}$ > with the minimum resolution of 2^{-15}
- **Fixed-point 32-bit fractional**—<-1 ; $1 - 2^{-31}$ > with the minimum resolution of 2^{-31}

The following list shows the accumulator types defined in the libraries:

- **Fixed-point 16-bit accumulator**—<-256.0 ; $256.0 - 2^{-7}$ > with the minimum resolution of 2^{-7}
- **Fixed-point 32-bit accumulator**—<-65536.0 ; $65536.0 - 2^{-15}$ > with the minimum resolution of 2^{-15}

The following list shows the floating-point types defined in the libraries:

- **Floating point 32-bit single precision**—< $-3.40282 \cdot 10^{38}$; $3.40282 \cdot 10^{38}$ > with the minimum resolution of 2^{-23}

1.1.3 API definition

AMCLIB uses the types mentioned in the previous section. To enable simple usage of the algorithms, their names use set prefixes and postfixes to distinguish the functions' versions. See the following example:

```
f32Result = MLIB_Mac_F32lss(f32Accum, f16Mult1, f16Mult2);
```

where the function is compiled from four parts:

- **MLIB**—this is the library prefix
- **Mac**—the function name—Multiply-Accumulate
- **F32**—the function output type

- *lss*—the types of the function inputs; if all the inputs have the same type as the output, the inputs are not marked

The input and output types are described in the following table:

Table 1. Input/output types

Type	Output	Input
<code>frac16_t</code>	F16	s
<code>frac32_t</code>	F32	l
<code>acc32_t</code>	A32	a
<code>float_t</code>	FLT	f

1.1.4 Supported compilers

AMCLIB for the ARM Cortex M4F core is written in C language or assembly language with C-callable interface depending on the specific function. The library is built and tested using the following compilers:

- MCUXpresso IDE
- IAR Embedded Workbench
- Keil μ Vision

For the MCUXpresso IDE, the library is delivered in the *amclib.a* file.

For the Kinetis Design Studio, the library is delivered in the *amclib.a* file.

For the IAR Embedded Workbench, the library is delivered in the *amclib.a* file.

For the Keil μ Vision, the library is delivered in the *amclib.lib* file.

The interfaces to the algorithms included in this library are combined into a single public interface include file, *amclib.h*. This is done to lower the number of files required to be included in your application.

1.1.5 Library configuration

AMCLIB for the ARM Cortex M4F core is written in C language or assembly language with C-callable interface depending on the specific function. Some functions from this library are inline type, which are compiled together with project using this library. The optimization level for inline function is usually defined by the specific compiler setting. It can cause an issue especially when high optimization level is set. Therefore the optimization level for all inline assembly written functions is defined by compiler pragmas using macros. The configuration header file *RTCESL_cfg.h* is located in: *specific library folder\MLIB\Include*. The optimization level can be changed by modifying the macro value for specific compiler. In case of any change the library functionality is not guaranteed.

1.1.6 Special issues

1. The equations describing the algorithms are symbolic. If there is positive 1, the number is the closest number to 1 that the resolution of the used fractional type allows. If there are maximum or minimum values mentioned, check the range allowed by the type of the particular function version.
2. The library functions that round the result (the API contains Rnd) round to nearest (half up).
3. This RTCESL requires the DSP extension for some saturation functions. If the core does not support the DSP extension feature the assembler code of the RTCESL will not be buildable. For example the core1 of the LPC55s69 has no DSP extension.

1.2 Library integration into project (MCUXpresso IDE)

This section provides a step-by-step guide on how to quickly and easily include AMCLIB into any MCUXpresso SDK example or new SDK project using MCUXpresso IDE. The SDK based project uses RTCESL from SDK package.

Adding RTCESL component to project

The MCUXpresso SDK package is necessary to add any example or new project and RTCESL component. In case the package has not been downloaded go to mcuxpresso.nxp.com, build the final MCUXpresso SDK package for required board and download it.

After package is downloaded, open the MCUXpresso IDE and drag&drop the SDK package in zip format to the Installed SDK window of the MCUXpresso IDE. After SDK package is dropped the message accepting window appears as can be show in following figure.

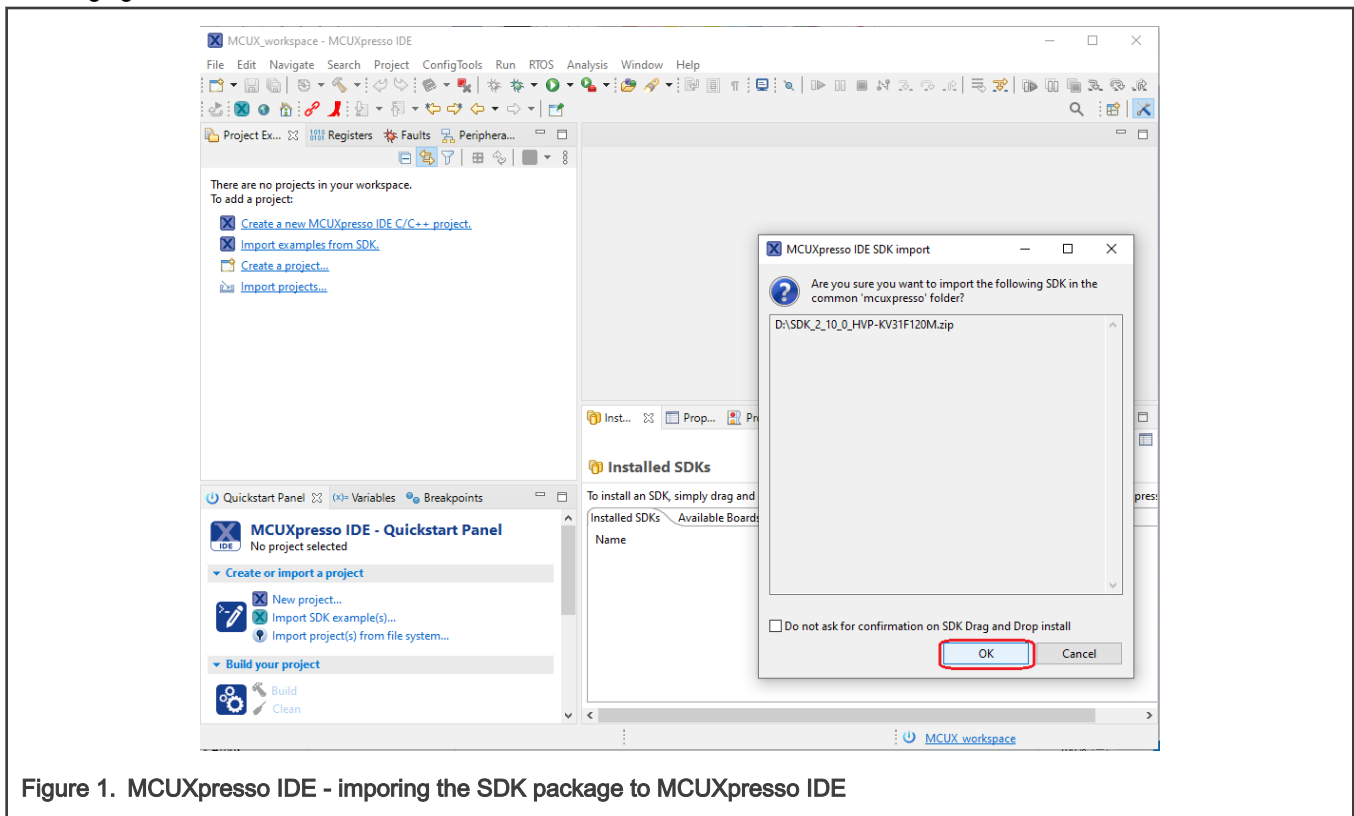


Figure 1. MCUXpresso IDE - importing the SDK package to MCUXpresso IDE

Click OK to confirm the SDK package import. Find the Quickstart panel in left bottom part of the MCUXpresso IDE and click New project... item or Import SDK example(s)... to add rtcesl component to the project.

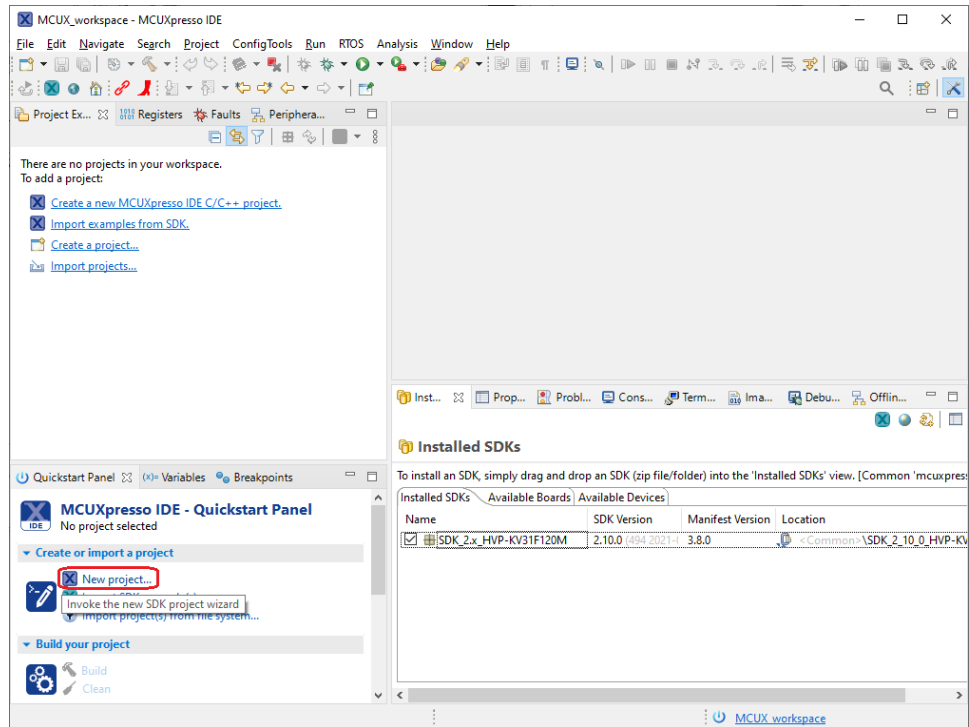


Figure 2. MCUXpresso IDE - create new project or Import SDK example(s)

Then select your board, and click Next button.

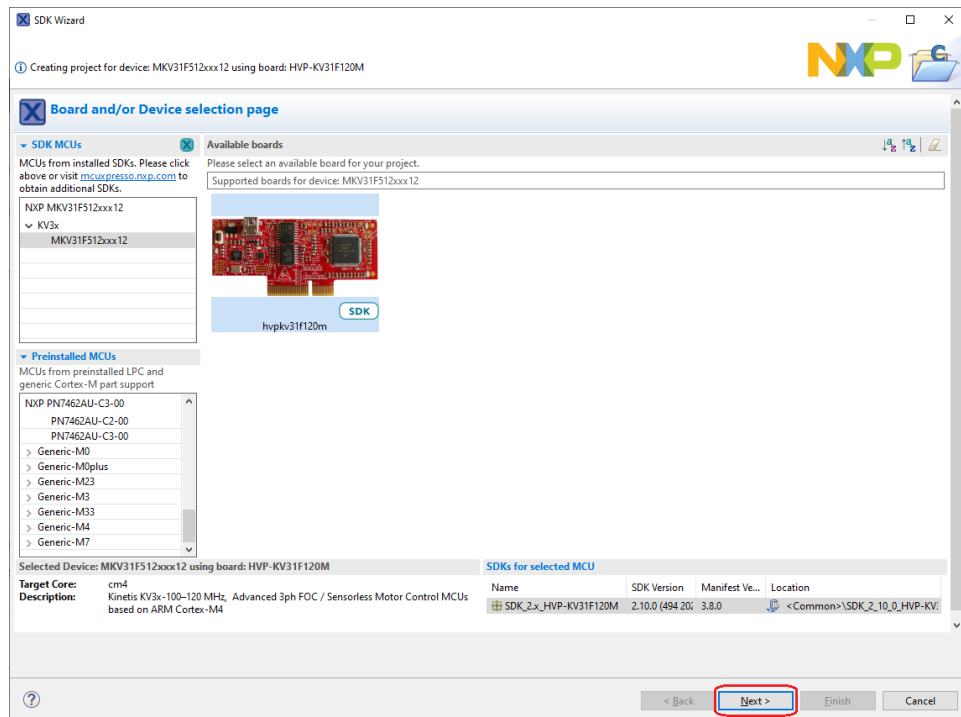


Figure 3. MCUXpresso IDE - selecting the board

Find the Middleware tab in the Components part of the window and click on the checkbox to be the rtcesl component ticked. Last step is to click the Finish button and wait for project creating with all RTCESL libraries and include paths.

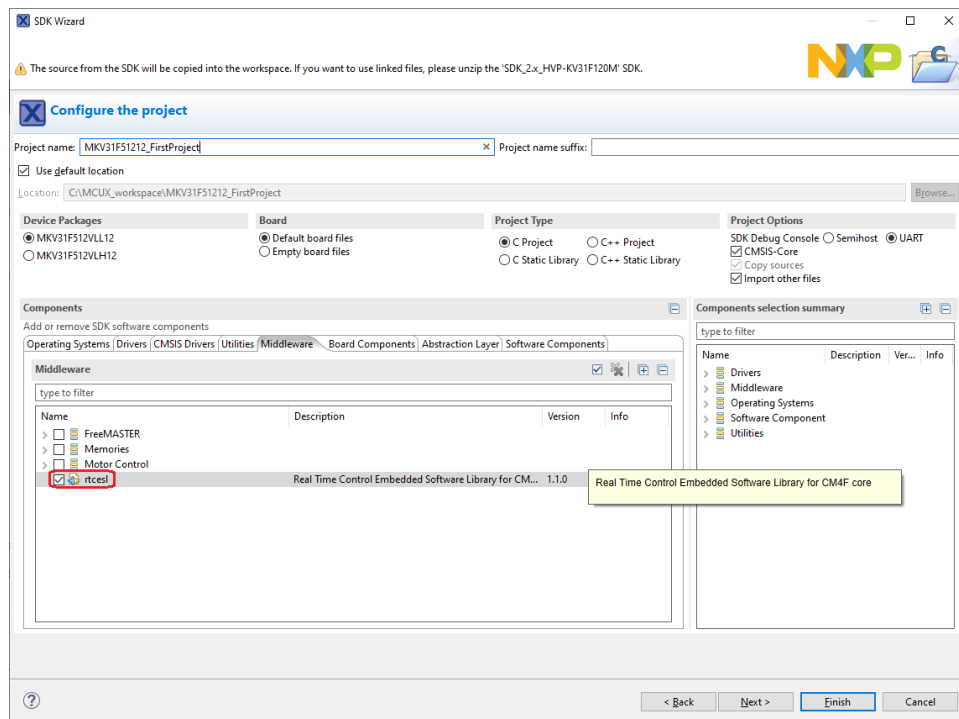


Figure 4. MCUXpresso IDE - selecting rtcsel component

Type the `#include` syntax into the code where you want to call the library functions. In the left-hand dialog, open the required .c file. After the file opens, include the following lines into the `#include` section:

```
#include "mlib_FP.h"
#include "gflib_FP.h"
#include "gdflib_FP.h"
#include "gmclib_FP.h"
#include "amclib_FP.h"
```

When you click the Build icon (hammer), the project is compiled without errors.

1.3 Library integration into project (Keil μ Vision)

This section provides a step-by-step guide on how to quickly and easily include AMCLIB into an empty project or any MCUXpresso SDK example or demo application projects using Keil μ Vision. This example uses the default installation path (C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL). If you have a different installation path, use that path instead. If any MCUXpresso SDK project is intended to use (for example hello_world project) go to [Linking the files into the project](#) chapter otherwise read next chapter.

NXP pack installation for new project (without MCUXpresso SDK)

This example uses the NXP MKV46F256xxx15 part, and the default installation path (C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL) is supposed. If the compiler has never been used to create any NXP MCU-based projects before, check whether the NXP MCU pack for the particular device is installed. Follow these steps:

1. Launch Keil μ Vision.
2. In the main menu, go to Project > Manage > Pack Installer....
3. In the left-hand dialog (under the Devices tab), expand the All Devices > Freescale (NXP) node.
4. Look for a line called "KVxx Series" and click it.

5. In the right-hand dialog (under the Packs tab), expand the Device Specific node.
6. Look for a node called "Keil::Kinetis_KVxx_DFP." If there are the Install or Update options, click the button to install/update the package. See [Figure 5](#).
7. When installed, the button has the "Up to date" title. Now close the Pack Installer.

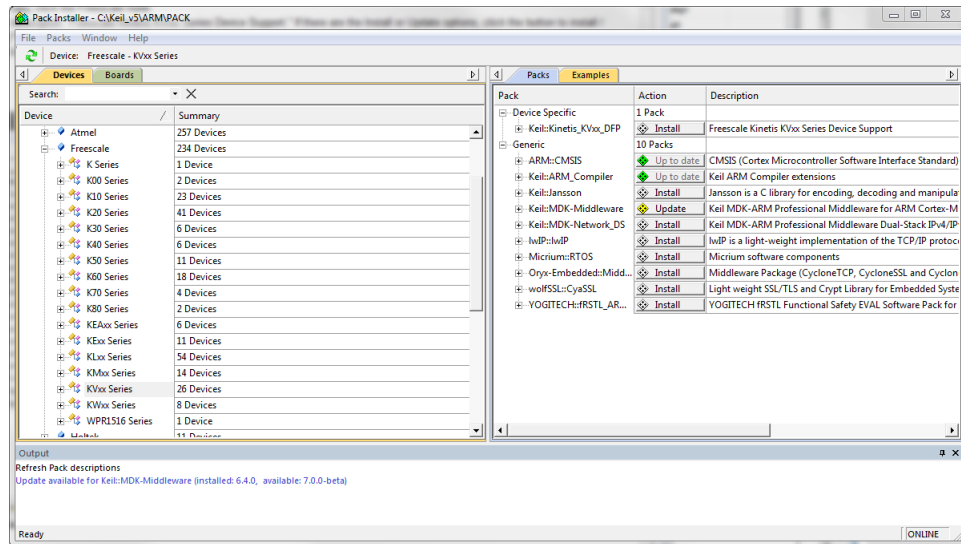


Figure 5. Pack Installer

New project (without MCUXpresso SDK)

To start working on an application, create a new project. If the project already exists and is opened, skip to the next section. Follow these steps to create a new project:

1. Launch Keil μ Vision.
2. In the main menu, select Project > New μ Vision Project..., and the Create New Project dialog appears.
3. Navigate to the folder where you want to create the project, for example C:\KeilProjects\MyProject01. Type the name of the project, for example MyProject01. Click Save. See [Figure 6](#).

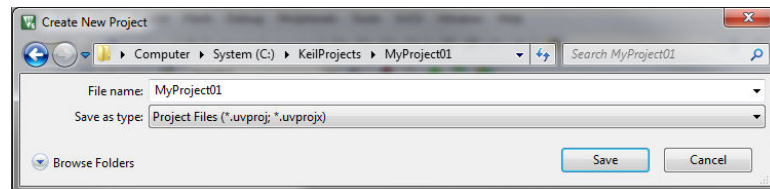


Figure 6. Create New Project dialog

4. In the next dialog, select the Software Packs in the very first box.
5. Type 'kv4' into the Search box, so that the device list is reduced to the KV4x devices.
6. Expand the KV4x node.
7. Click the MKV46F256xxx15 node, and then click OK. See [Figure 7](#).

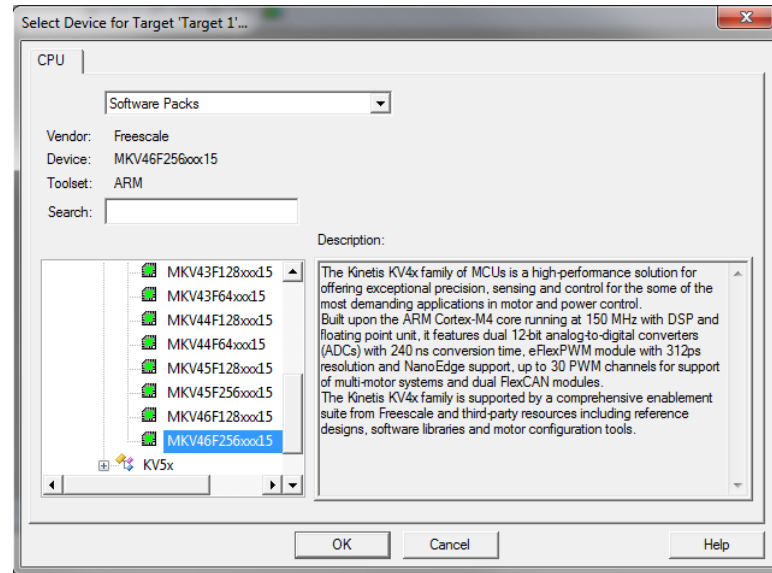


Figure 7. Select Device dialog

8. In the next dialog, expand the Device node, and tick the box next to the Startup node. See Figure 8.
9. Expand the CMSIS node, and tick the box next to the CORE node.

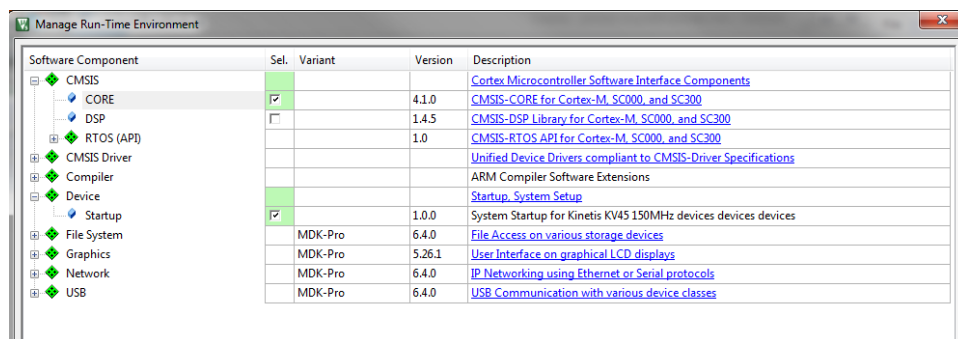
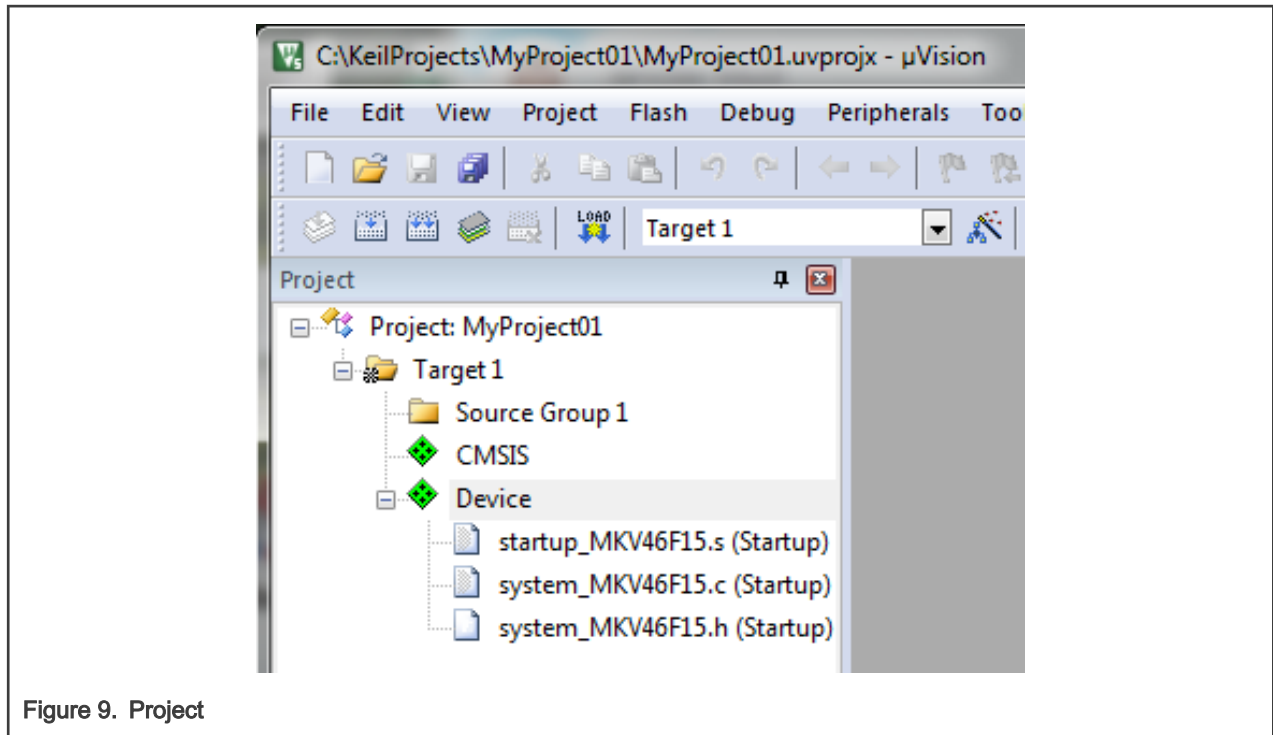
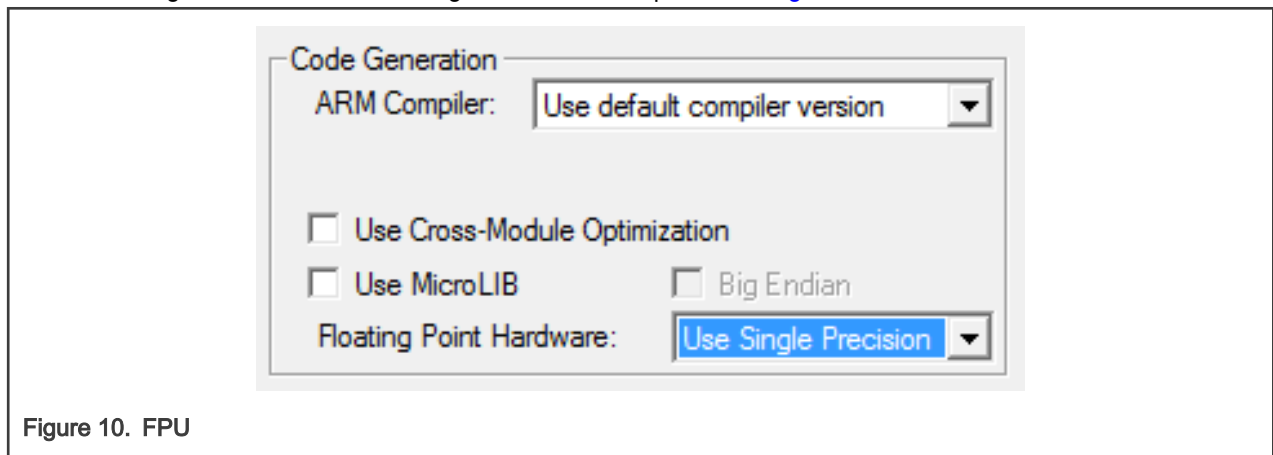


Figure 8. Manage Run-Time Environment dialog

10. Click OK, and a new project is created. The new project is now visible in the left-hand part of Keil µVision. See Figure 9.



11. In the main menu, go to Project > Options for Target 'Target1'..., and a dialog appears.
12. Select the Target tab.
13. Select Use Single Precision in the Floating Point Hardware option. See [Figure 9](#).



Linking the files into the project

AMCLIB requires MLIB and GDFLIB and GFLIB and GMCLIB to be included too. The following steps show how to include all dependent modules.

To include the library files in the project, create groups and add them.

1. Right-click the Target 1 node in the left-hand part of the Project tree, and select Add Group... from the menu. A new group with the name New Group is added.
2. Click the newly created group, and press F2 to rename it to RTCESL.
3. Right-click the RTCESL node, and select Add Existing Files to Group 'RTCESL'... from the menu.

4. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\MLIB\Include, and select the *mllib_FP.h* file. If the file does not appear, set the Files of type filter to Text file. Click Add. See [Figure 11](#).

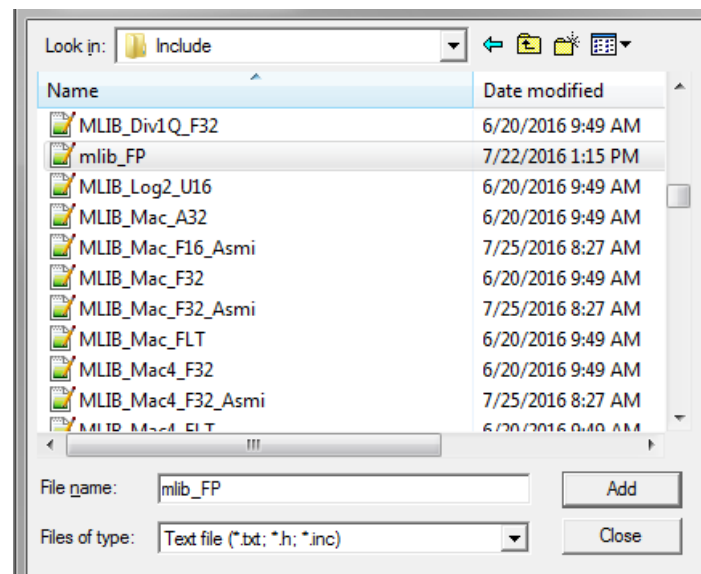


Figure 11. Adding .h files dialog

5. Navigate to the parent folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\MLIB, and select the *mllib.lib* file. If the file does not appear, set the Files of type filter to Library file. Click Add. See [Figure 12](#).

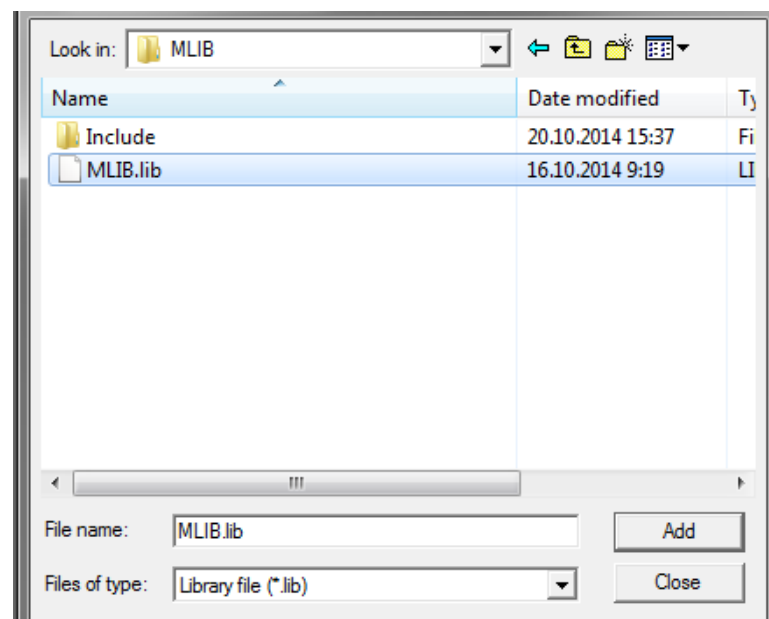


Figure 12. Adding .lib files dialog

6. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GFLIB\Include, and select the *gflib_FP.h* file. If the file does not appear, set the Files of type filter to Text file. Click Add.
7. Navigate to the parent folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GFLIB, and select the *gflib.lib* file. If the file does not appear, set the Files of type filter to Library file. Click Add.
8. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GDFLIB\Include, and select the *gdflib_FP.h* file. If the file does not appear, set the Files of type filter to Text file. Click Add.

9. Navigate to the parent folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GDFLIB, and select the *gdfplib.lib* file. If the file does not appear, set the Files of type filter to Library file. Click Add.
10. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GMCLIB\Include, and select the *gmclib_FP.h* file. If the file does not appear, set the Files of type filter to Text file. Click Add.
11. Navigate to the parent folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GMCLIB, and select the *gmclib.lib* file. If the file does not appear, set the Files of type filter to Library file. Click Add.
12. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\AMCLIB\Include, and select the *amclib_FP.h* file. If the file does not appear, set the Files of type filter to Text file. Click Add.
13. Navigate to the parent folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\AMCLIB, and select the *amclib.lib* file. If the file does not appear, set the Files of type filter to Library file. Click Add.
14. Now, all necessary files are in the project tree; see [Figure 13](#). Click Close.

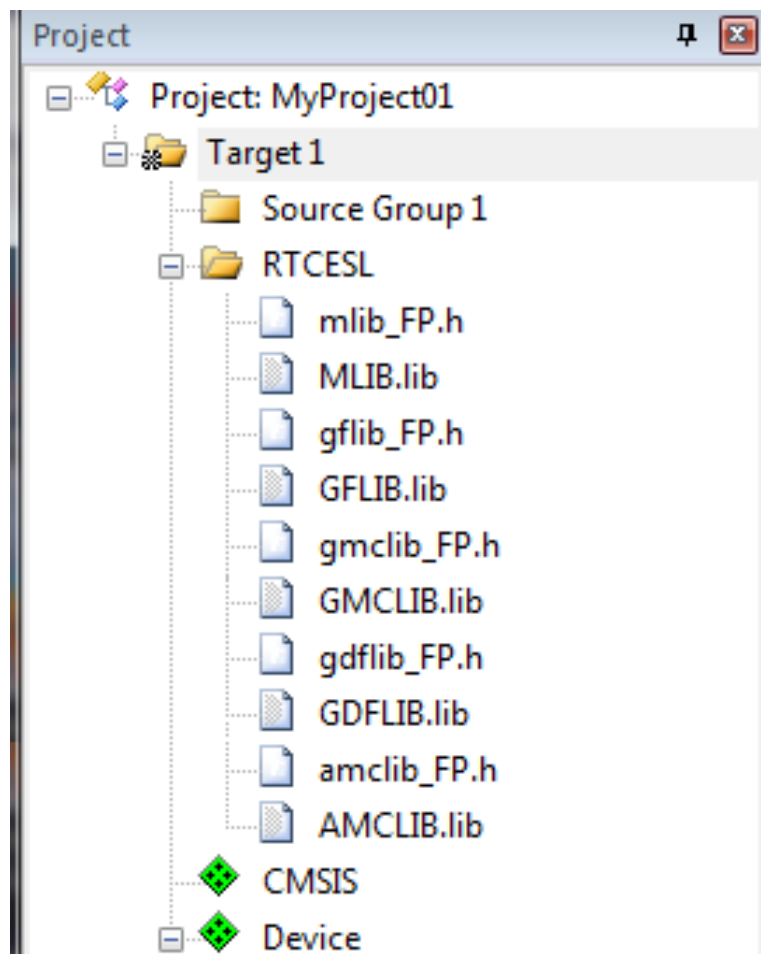


Figure 13. Project workspace

Library path setup

The following steps show the inclusion of all dependent modules.

1. In the main menu, go to Project > Options for Target 'Target1'..., and a dialog appears.
2. Select the C/C++ tab. See [Figure 14](#).

3. In the Include Paths text box, type the following paths (if there are more paths, they must be separated by ';') or add them by clicking the ... button next to the text box:
 - "C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\MLIB\Include"
 - "C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GFLIB\Include"
 - "C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GDFLIB\Include"
 - "C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\GMCLIB\Include"
 - "C:\NXP\RTCESL\CM4F_RTCESL_4.7_KEIL\AMCLIB\Include"
4. Click OK.
5. Click OK in the main dialog.

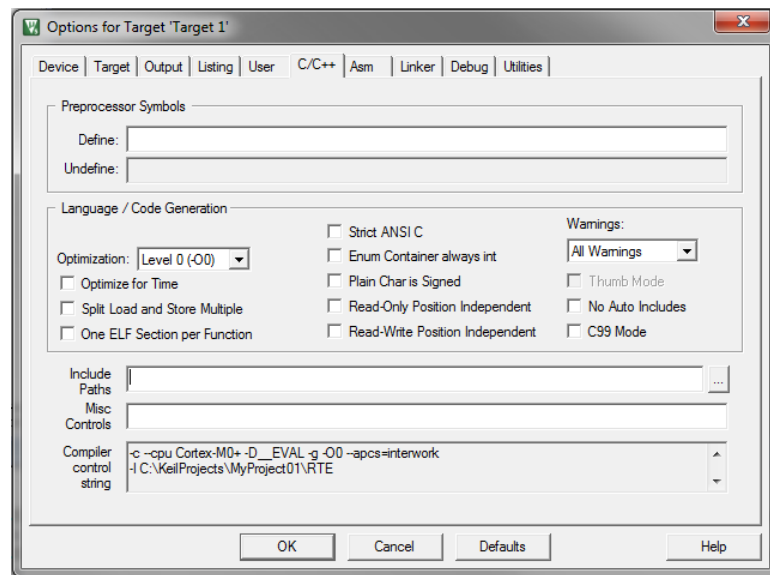


Figure 14. Library path addition

Type the `#include` syntax into the code. Include the library into a source file. In the new project, it is necessary to create a source file:

1. Right-click the Source Group 1 node, and Add New Item to Group 'Source Group 1'... from the menu.
2. Select the C File (.c) option, and type a name of the file into the Name box, for example '*main.c*'. See [Figure 15](#).

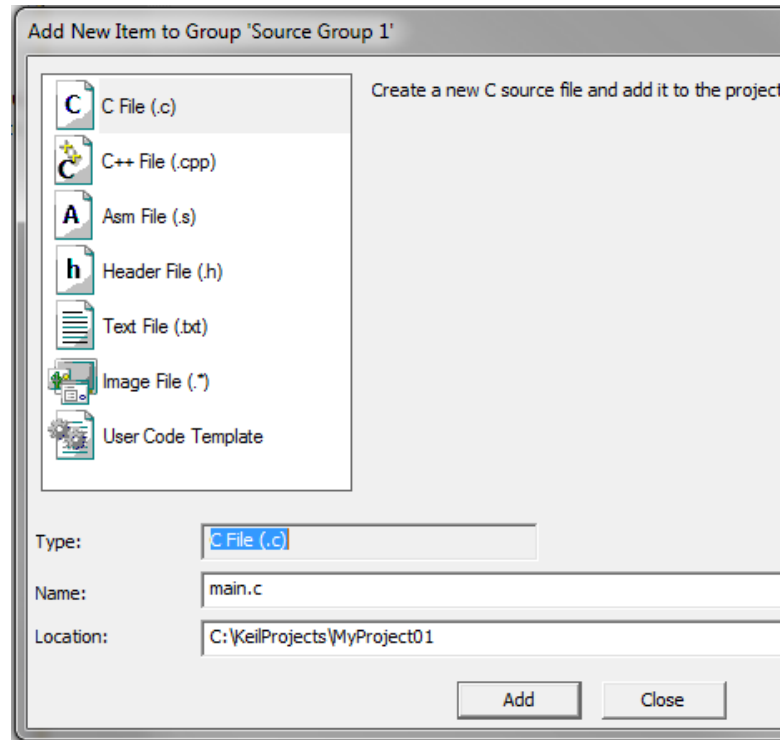


Figure 15. Adding new source file dialog

3. Click Add, and a new source file is created and opened up.
4. In the opened source file, include the following lines into the `#include` section, and create a main function:

```
#include "mlib_FP.h"
#include "gflib_FP.h"
#include "gdfliB_FP.h"
#include "gmclib_FP.h"
#include "amclib_FP.h"

int main(void)
{
    while(1);
}
```

When you click the Build (F7) icon, the project will be compiled without errors.

1.4 Library integration into project (IAR Embedded Workbench)

This section provides a step-by-step guide on how to quickly and easily include the AMCLIB into an empty project or any MCUXpresso SDK example or demo application projects using IAR Embedded Workbench. This example uses the default installation path (C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR). If you have a different installation path, use that path instead. If any MCUXpresso SDK project is intended to use (for example hello_world project) go to [Linking the files into the project](#) chapter otherwise read next chapter.

New project (without MCUXpresso SDK)

This example uses the NXP MKV46F256xxx15 part, and the default installation path (C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR) is supposed. To start working on an application, create a new project. If the project already exists and is opened, skip to the next section. Perform these steps to create a new project:

1. Launch IAR Embedded Workbench.
2. In the main menu, select Project > Create New Project... so that the "Create New Project" dialog appears. See [Figure 16](#).

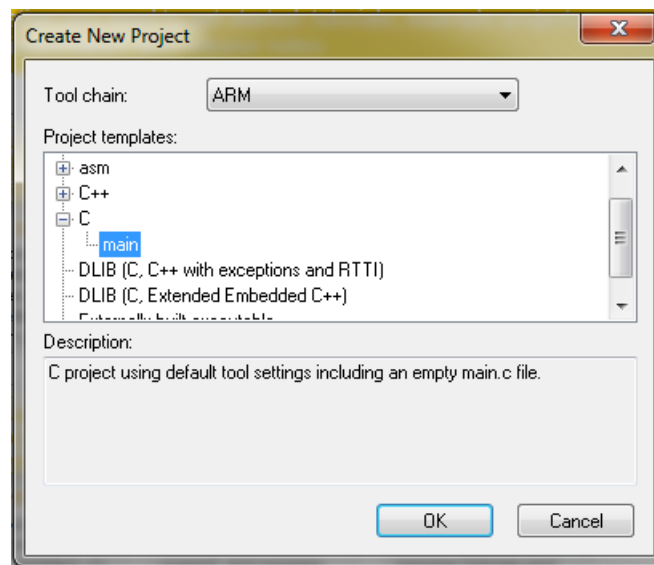


Figure 16. Create New Project dialog

3. Expand the C node in the tree, and select the "main" node. Click OK.
4. Navigate to the folder where you want to create the project, for example, C:\IARProjects\MyProject01. Type the name of the project, for example, MyProject01. Click Save, and a new project is created. The new project is now visible in the left-hand part of IAR Embedded Workbench. See [Figure 17](#).

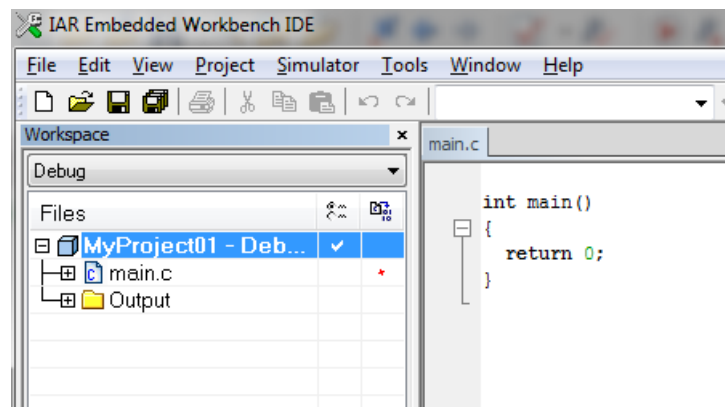


Figure 17. New project

5. In the main menu, go to Project > Options..., and a dialog appears.
6. In the Target tab, select the Device option, and click the button next to the dialog to select the MCU. In this example, select NXP > KV4x > NXP MKV46F256xxx15. Select VFPv4 single precision in the FPU option. Click OK. See [Figure 18](#).

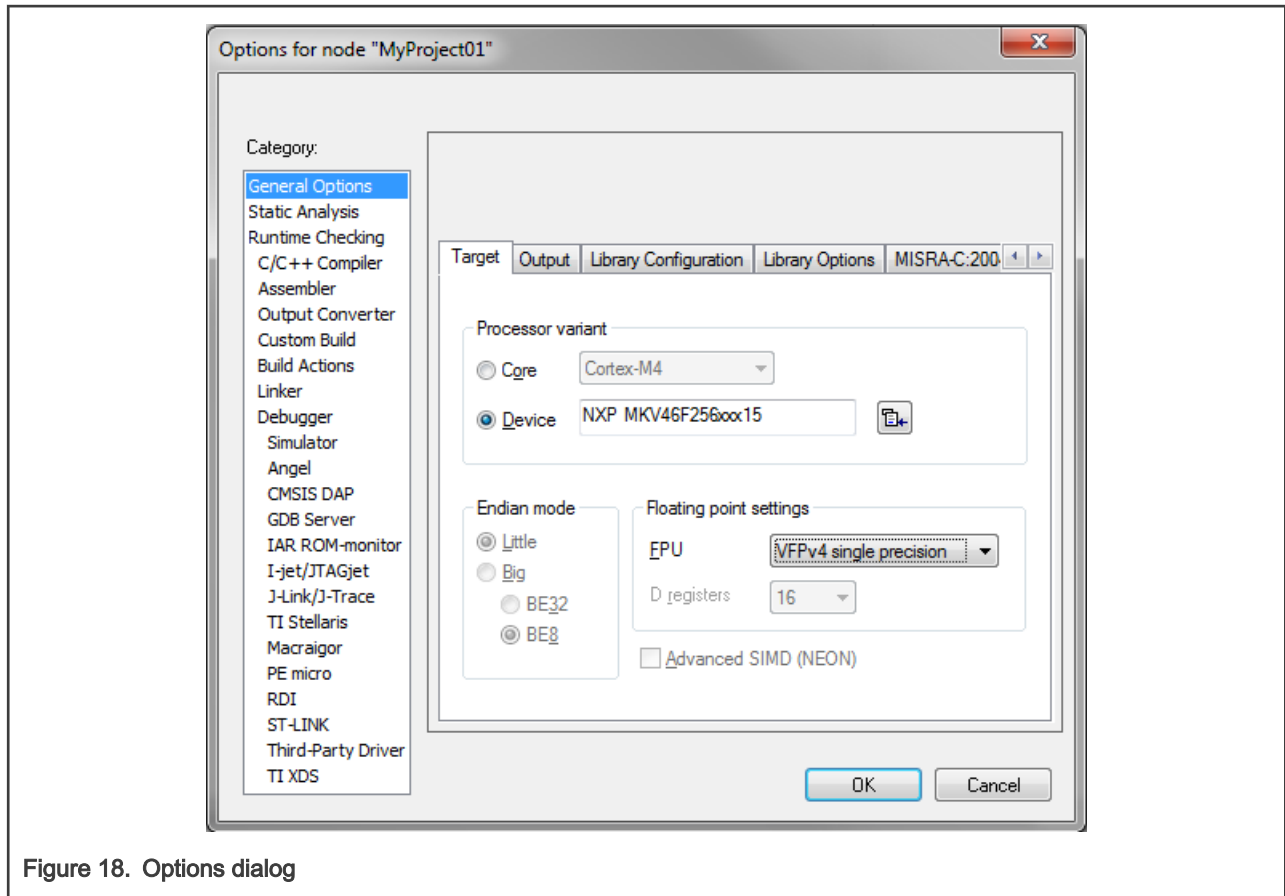


Figure 18. Options dialog

Library path variable

To make the library integration easier, create a variable that will hold the information about the library path.

1. In the main menu, go to Tools > Configure Custom Argument Variables..., and a dialog appears.
2. Click the New Group button, and another dialog appears. In this dialog, type the name of the group PATH, and click OK. See [Figure 19](#).

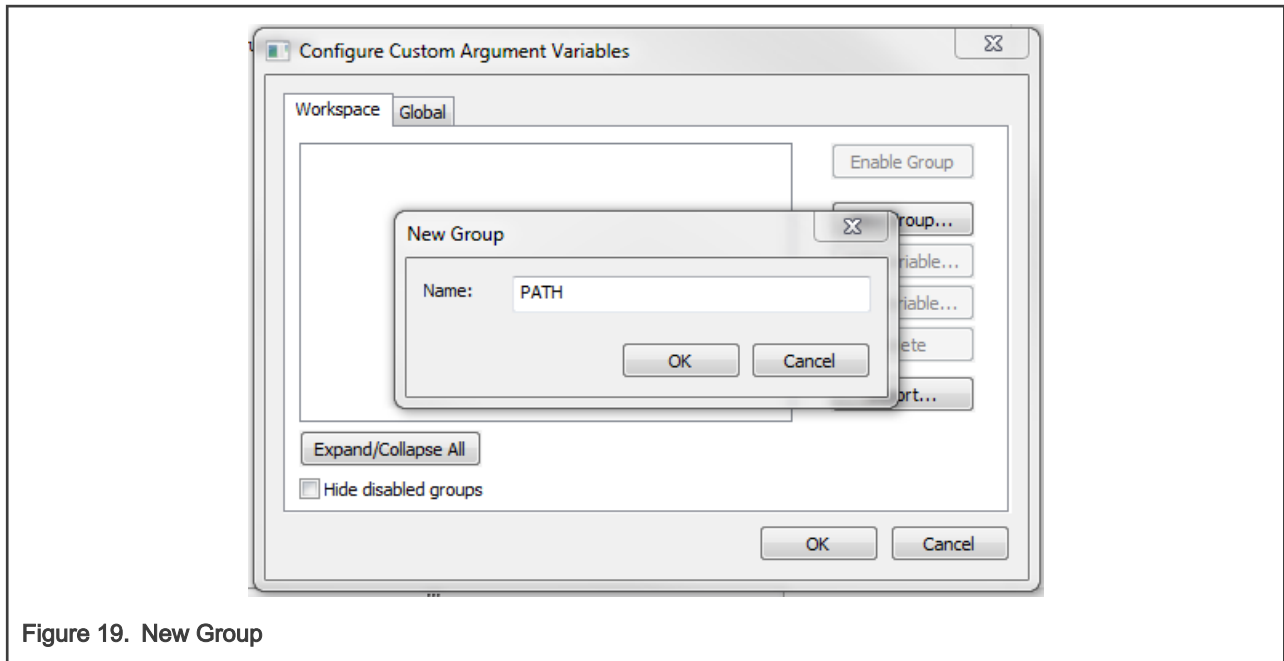


Figure 19. New Group

3. Click on the newly created group, and click the Add Variable button. A dialog appears.
4. Type this name: RTCESL_LOC
5. To set up the value, look for the library by clicking the '...' button, or just type the installation path into the box: C:\NXP\RTCESL\CM4F_RTCESEL_4.7_IAR. Click OK.
6. In the main dialog, click OK. See [Figure 20](#).

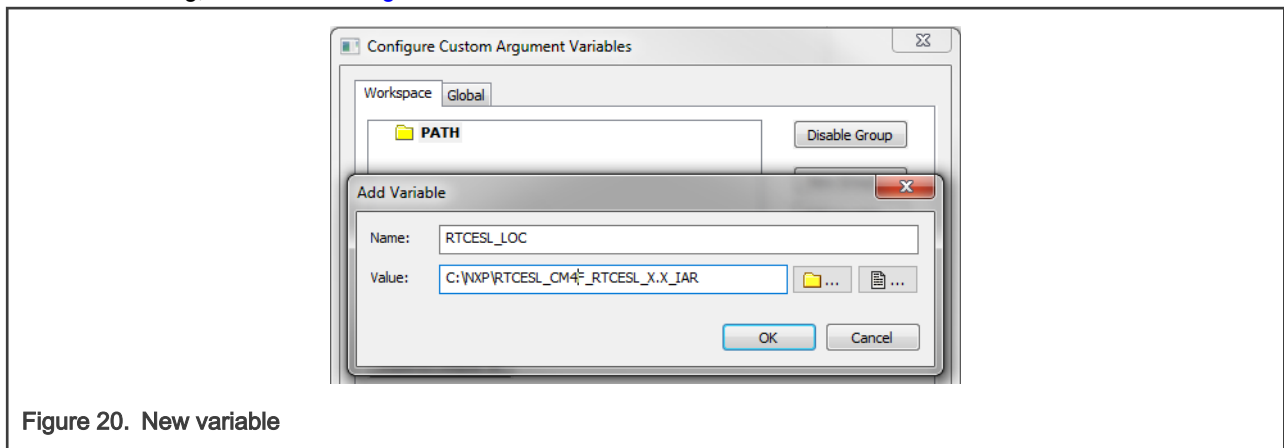


Figure 20. New variable

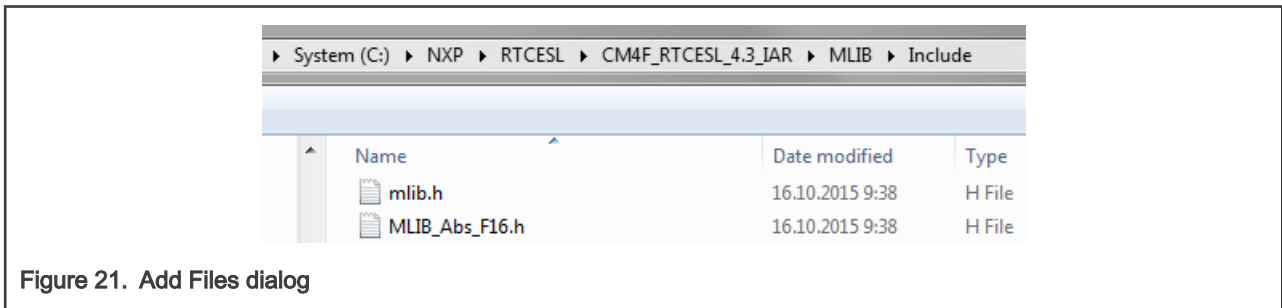
Linking the files into the project

AMCLIB requires MLIB and GDFLIB and GFLIB and GMCLIB to be included too. The following steps show the inclusion of all dependent modules.

To include the library files into the project, create groups and add them.

1. Go to the main menu Project > Add Group...
2. Type RTCESL, and click OK.
3. Click on the newly created node RTCESL, go to Project > Add Group..., and create a MLIB subgroup.
4. Click on the newly created node MLIB, and go to the main menu Project > Add Files... See [Figure 22](#).

5. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\MLIB\Include, and select the *mllib_FP.h* file. (If the file does not appear, set the file-type filter to Source Files.) Click Open. See [Figure 21](#).
6. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\MLIB, and select the *mllib.a* file. If the file does not appear, set the file-type filter to Library / Object files. Click Open.



7. Click on the RTCESL node, go to Project > Add Group..., and create a GFLIB subgroup.
8. Click on the newly created node GFLIB, and go to the main menu Project > Add Files....
9. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GFLIB\Include, and select the *gflib_FP.h* file. (If the file does not appear, set the file-type filter to Source Files.) Click Open.
10. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GFLIB, and select the *gflib.a* file. If the file does not appear, set the file-type filter to Library / Object files. Click Open.
11. Click on the RTCESL node, go to Project > Add Group..., and create a GDFLIB subgroup.
12. Click on the newly created node GDFLIB, and go to the main menu Project > Add Files....
13. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GDFLIB\Include, and select the *gdflib_FP.h* file. (If the file does not appear, set the file-type filter to Source Files.) Click Open.
14. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GDFLIB, and select the *gdflib.a* file. If the file does not appear, set the file-type filter to Library / Object files. Click Open.
15. Click on the RTCESL node, go to Project > Add Group..., and create a GMCLIB subgroup.
16. Click on the newly created node GMCLIB, and go to the main menu Project > Add Files....
17. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GMCLIB\Include, and select the *gmclib_FP.h* file. If the file does not appear, set the file-type filter to Source Files. Click Open.
18. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\GMCLIB, and select the *gmclib.a* file. If the file does not appear, set the file-type filter to Library / Object files. Click Open.
19. Click on the RTCESL node, go to Project > Add Group..., and create an AMCLIB subgroup.
20. Click on the newly created node AMCLIB, and go to the main menu Project > Add Files....
21. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\AMCLIB\Include, and select the *amclib_FP.h* file. If the file does not appear, set the file-type filter to Source Files. Click Open.
22. Navigate into the library installation folder C:\NXP\RTCESL\CM4F_RTCESL_4.7_IAR\AMCLIB, and select the *amclib.a* file. If the file does not appear, set the file-type filter to Library / Object files. Click Open.
23. Now you will see the files added in the workspace. See [Figure 22](#).

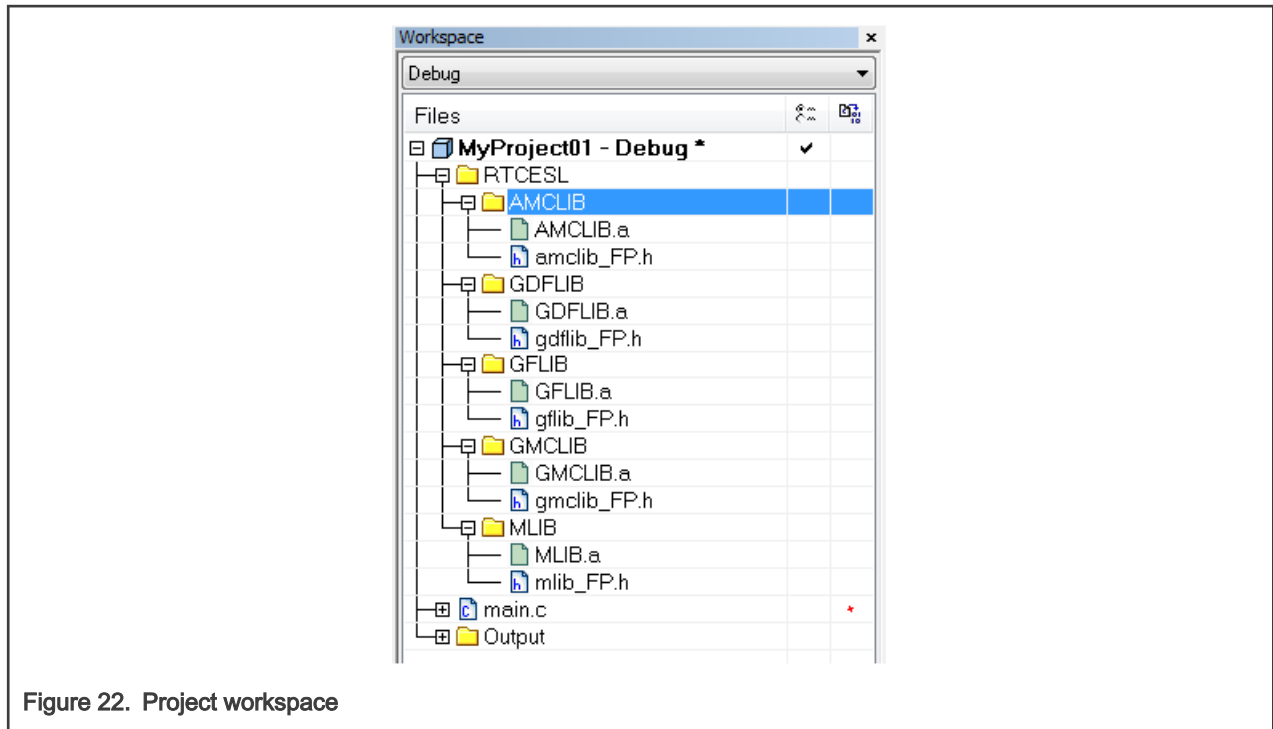


Figure 22. Project workspace

Library path setup

The following steps show the inclusion of all dependent modules:

1. In the main menu, go to Project > Options..., and a dialog appears.
2. In the left-hand column, select C/C++ Compiler.
3. In the right-hand part of the dialog, click on the Preprocessor tab (it can be hidden in the right; use the arrow icons for navigation).
4. In the text box (at the Additional include directories title), type the following folder (using the created variable):
 - \$RTCESL_LOC\$MLIB\Include
 - \$RTCESL_LOC\$GFLIB\Include
 - \$RTCESL_LOC\$GDFLIB\Include
 - \$RTCESL_LOC\$GMCLIB\Include
 - \$RTCESL_LOC\$AMCLIB\Include
5. Click OK in the main dialog. See [Figure 23](#).

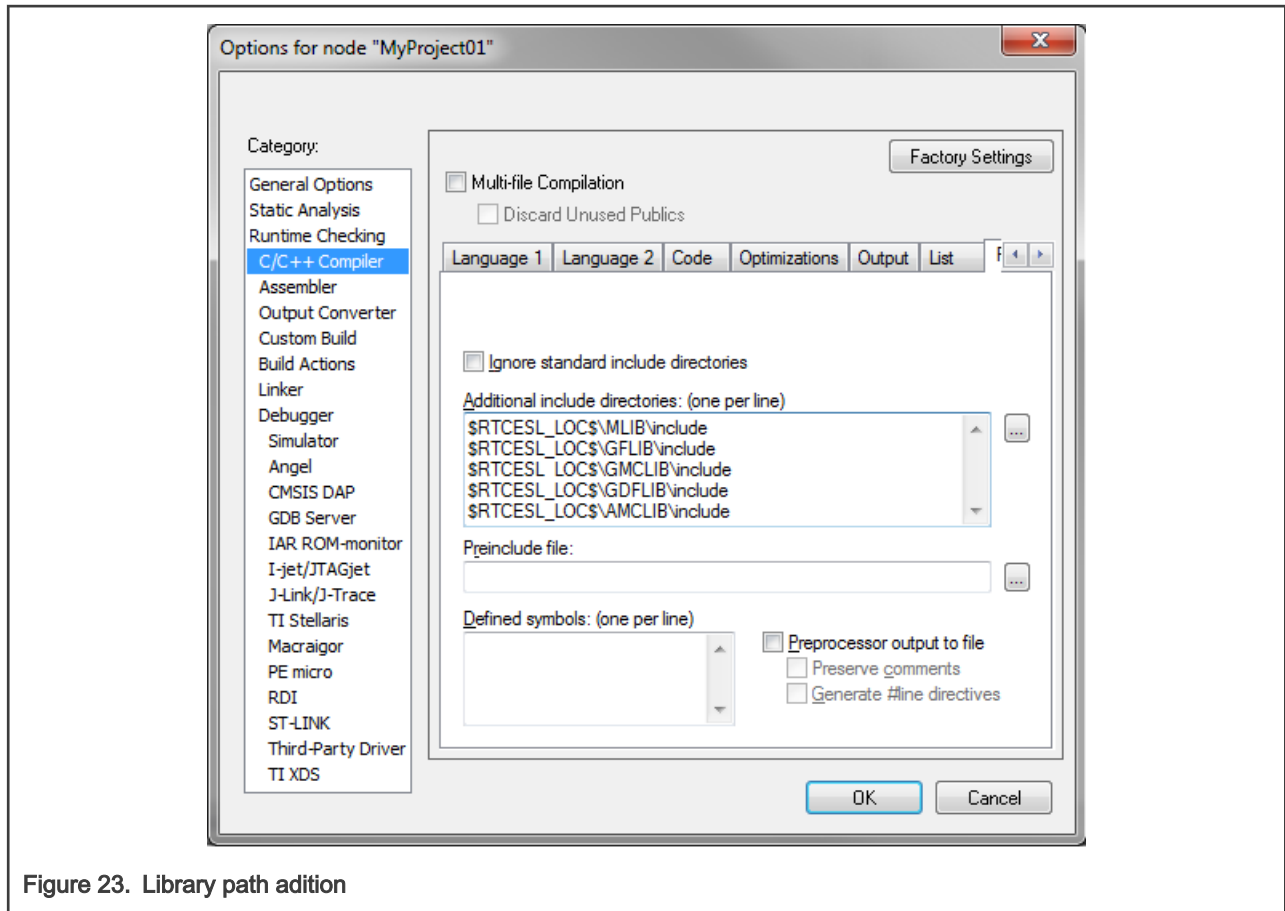


Figure 23. Library path addition

Type the `#include` syntax into the code. Include the library included into the *main.c* file. In the workspace tree, double-click the *main.c* file. After the *main.c* file opens up, include the following lines into the `#include` section:

```
#include "mlib_FP.h"
#include "gflib_FP.h"
#include "gdflib_FP.h"
#include "gmclib_FP.h"
#include "amclib_FP.h"
```

When you click the Make icon, the project will be compiled without errors.

Chapter 2

Algorithms in detail

2.1 AMCLIB_ACIMCtrlMTPA

The [AMCLIB_ACIMCtrlMTPA](#) function enables to minimize the ACIM losses by applying the Max Torque per Ampere (MTPA) strategy. The principle is derived from the ACIM torque equation:

$$T(\theta_I) = \frac{3}{2} \cdot P_P \cdot \frac{L_m^2}{L_r} \cdot i_{sd}(\theta_I) \cdot i_{sq}(\theta_I) = \frac{3}{4} \cdot P_P \cdot \frac{L_m^2}{L_r} \cdot |i_{sdq}| \cdot \sin(2 \cdot \theta_I)$$

where:

- i_{sd} is the D component of the stator current vector
- i_{sq} is the Q component of the stator current vector
- i_{sdq} is the stator current vector
- θ_I is the angle of stator the current vector
- L_r is the rotor equivalent inductance
- L_m is the mutual equivalent inductance
- P_P is the motor pole pair number constant
- T is the motor mechanic torque

Motor torque depends on the angle of the stator current vector. Maximum efficiency (minimum stator joule losses) can be calculated when motor torque differential is equal zero:

$$\frac{dT(\theta_I)}{d\theta_I} = \frac{3}{4} \cdot P_P \cdot \frac{L_m^2}{L_r} \cdot |i_{sdq}| \cdot \cos(2 \cdot \theta_I) = 0 \Rightarrow \theta_I = \frac{\pi}{4}$$

It is clear that the stator current components must be the same values to achieve the $\theta_I = \pi/4$ angle. The MTPA stator current vector trajectory in consideration of the i_{sd} limits given by the minimal field excitation and current limitations is shown in [Figure 1](#).

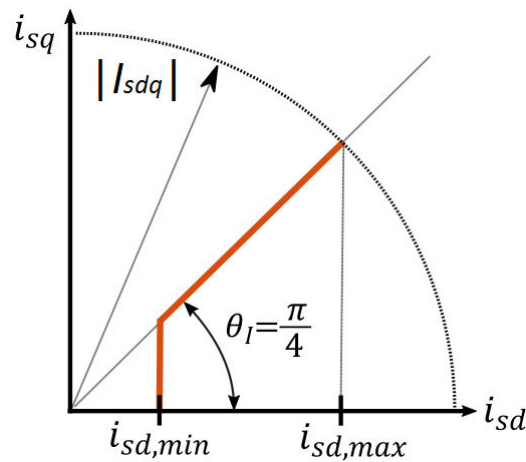


Figure 24. Minimal losses stator current vector trajectory with limits

2.1.1 Available versions

The available versions of the [AMCLIB_ACIMCtrlMTPA](#) function are shown in the following table:

Table 2. Init function versions

Function name	Input type		Parameters	Result type
	IdMin	IdMax		
AMCLIB_ACIMCtrlMTPAInit_FLT	float_t	float_t	AMCLIB_ACIM_CTRL_MTPA_T_FLT*	void
The input arguments are the 32-bit single precision floating-point values that contain the limits for i_{sd} . They both are positive values (the minimum must be lower than the maximum) and the pointers to a structure that contains the parameters defined in AMCLIB_ACIM_CTRL_MTPA_T_FLT type description .				

Table 3. Function version

Function name	Input type	Parameters	Result type
AMCLIB_ACIMCtrlMTPA_FLT	float_t	AMCLIB_ACIM_CTRL_MTPA_T_FLT*	float_t
The input arguments are the 32-bit single precision floating-point values that contain the limits for i_{sd} . They both are positive values (the minimum must be lower than the maximum) and the pointers to a structure that contains the parameters defined in AMCLIB_ACIM_CTRL_MTPA_T_FLT type description .			

2.1.2 AMCLIB_ACIM_CTRL_MTPA_T_FLT type description

Variable name	Data type	Description
fltIdExpParam	GDFLIB_FILTER_EXP_T_FLT	The exponential filter structure for the i_{sd} current filtration. Set by the user.
fltLowerLim	float_t	The minimal output limit of i_{sd} . Usually determined from the minimum ACIM rotor flux excitation, as shown in Figure 1 . Set by the user, must be a positive value lower than the upper limit.
fltUpperLim	float_t	The maximal output limit of i_{sd} . Usually determined from the maximum (typically nominal) ACIM current, as shown in Figure 1 . Set by the user, must be a positive value higher than the lower limit.

2.1.3 Declaration

The available AMCLIB_ACIMCtrlMTPAInit functions have the following declarations:

```
void AMCLIB_ACIMCtrlMTPAInit_FLT(float\_tfltMin,float\_tfltMax,AMCLIB\_ACIM\_CTRL\_MTPA\_T\_FLT *psCtrl)
```

The available AMCLIB_ACIMCtrlMTPA functions have the following declarations:

```
float\_t AMCLIB_ACIMCtrlMTPA_FLT(float\_tfltIq,AMCLIB\_ACIM\_CTRL\_MTPA\_T\_FLT *psCtrl)
```

2.1.4 Function use

The use of the `AMCLIB_ACIMCtrlMTPA` function is shown in the following examples:

Floating-point version:

```
#include "amclib.h"

static AMCLIB_ACIM_CTRL_MTPA_T_FLT sMTPAParam;
static float_t fltIsd;
static float_t fltIsq;
static float_t fltIDMin;
static float_t fltIDMax;

void Isr(void);

void main (void)
{
    /* Structure parameter setting */
    sMTPAParam.sCtrl.fltIdExpParam.fltA = 0.05F;
    fltIDMin = 0.1F;
    fltIDMax = 2.2F;

    /* Initialization of the ACIMCtrlMTPA's structure */
    AMCLIB_ACIMCtrlMTPAInit_FLT (fltIDMin, fltIDMax, &sMTPAParam);

    /* Assign Isq value */
    fltIsq = -0.6F;
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Calculating required Isd by MTPA algorithm */
    fltIsd = AMCLIB_ACIMCtrlMTPA_FLT(fltIsq, &sMTPAParam);
}
```

2.2 AMCLIB_ACIMRotFluxObsrv

The `AMCLIB_ACIMRotFluxObsrv` function calculates the ACIM flux estimate and its position (angle) from the available measured signals (currents and voltages). In the case of ACIM FOC, the rotor flux position (angle) is needed to perform the Park transformation.

The closed-loop flux observer is formed from the two most desirable open-loop estimators, which are referred to as the voltage model and the current model (as shown in [Figure 1](#)). The current model is used for low-speed operation and the voltage model is used for high-speed operation. A smooth transition between these two models is ensured by the PI controller.

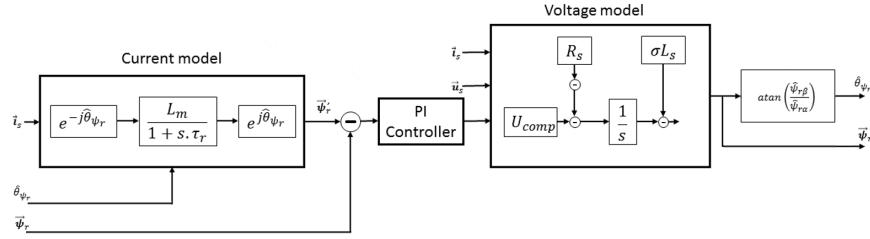


Figure 25. ACIM rotor flux observer block diagram

The voltage model (stator model) is used to estimate the stator flux-linkage vector or the rotor flux-linkage vector without a speed signal. The voltage model is derived by integrating the stator voltage equation in the stator stationary coordinates as:

$$\begin{aligned}\vec{u}_s &= R_s \cdot \vec{i}_s + \frac{d\vec{\psi}_s}{dt} \\ \vec{\psi}_s &= \int (\vec{u}_s - R_s \cdot \vec{i}_s) dt \\ \vec{\psi}_r &= \frac{L_r}{L_m} (\vec{\psi}_s - L_s \cdot \sigma \cdot \vec{i}_s)\end{aligned}$$

Expressed in discrete form as:

$$\begin{aligned}\psi_{s\alpha}(k) &= \frac{\tau_l}{\tau_l + T_s} [\psi_{s\alpha}(k-1) + T_s \cdot (u_{s\alpha}(k) - R_s \cdot i_{s\alpha}(k))] \\ \psi_{s\beta}(k) &= \frac{\tau_l}{\tau_l + T_s} [\psi_{s\beta}(k-1) + T_s \cdot (u_{s\beta}(k) - R_s \cdot i_{s\beta}(k))] \\ \psi_{r\alpha}(k) &= \frac{L_r}{L_m} (\psi_{s\alpha}(k) - L_s \cdot \sigma \cdot i_{s\alpha}(k)) \\ \psi_{r\beta}(k) &= \frac{L_r}{L_m} (\psi_{s\beta}(k) - L_s \cdot \sigma \cdot i_{s\beta}(k))\end{aligned}$$

where:

- u_s is the stator voltage vector
- i_s is the stator current vector
- Ψ_s is the stator flux-linkage vector
- Ψ_r is the rotor flux-linkage vector
- ω_r is the rotor electrical angular speed
- ω_s is the electrical angular slip speed
- R_s is the stator resistance
- R_r is the rotor equivalent resistance
- L_s is the stator equivalent inductance
- L_r is the rotor equivalent inductance
- L_m is the mutual equivalent inductance
- τ_r is the motor electrical time constant
- T_s is the sample time
- σ is the motor leakage coefficient

These equations show that the rotor flux linkage is basically the difference between the stator flux-linkage and the leakage flux. The rotor flux equation is used to estimate the respective flux-linkage vector, corresponding angle. The argument Ψ_r of the rotor flux-linkage vector is the rotor field angle θ_{ψ_r} calculated as:

$$\theta_{\psi_r} = \text{atan}\left(\frac{\psi_{r\beta}}{\psi_{r\alpha}}\right)$$

The voltage model (stator model) is sufficiently robust and accurate at higher stator frequencies. Two basic deficiencies can degrade this model as the speed reduces: the integration problem, and model's sensitivity to stator resistance mismatch.

The current model (rotor model) is derived from the differential equation of the rotor winding. The stator coordinate implementation is:

$$\frac{d\vec{\psi}_r}{dt} = \frac{L_m}{\tau_r} \vec{i}_s - \frac{1}{\tau_r} \vec{\psi}_r - j\omega_{slip} \cdot \vec{\psi}_r$$

When applying field-oriented control assumptions (such as $\Psi_{rq} = 0$), then the rotor flux estimated by the current model in the synchronous rotating frame is:

$$\frac{d\vec{\psi}_{rd}}{dt} = -\frac{1}{\tau_r} \vec{\psi}_{rd} + \frac{L_m}{\tau_r} \vec{i}_{sd}$$

In discrete form:

$$\psi_{rd}(k) = \frac{\tau_r}{\tau_r + T_s} \left[\psi_{rd}(k-1) + T_s \frac{L_m}{\tau_r} i_{sd}(k) \right]$$

The accuracy of the rotor model depends on correct model parameters. It is the rotor time constant in particular that determines the accuracy of the estimated field angle (the most critical variable in a vector-controlled drive).

2.2.1 Available versions

The available versions of the [AMCLIB_ACIMRotFluxObsrv](#) function are shown in the following table:

Table 4. Init version

Function name	Parameters	Result type
AMCLIB_ACIMRotFluxObsrvInit_FLT	AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT *	void
	The initialization does not have any input.	

Table 5. Function version

Function name	Input/output type		Result type
AMCLIB_ACIMRotFluxObsrv_FLT	Input	GMCLIB_2COOR_ALBE_T_FLT *	void
		GMCLIB_2COOR_ALBE_T_FLT *	
	Parameters	AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT *	
	Rotor flux observer with a 32-bit single precision floating-point inputs: stator current and voltage in alpha-beta coordinates. All are within the full range. The function does not return anything. All calculated variables are stored in the AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT structure.		

2.2.2 AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT type description

Variable name		Data type	Description
sPsiRotRDQ		GMCLIB_2COOR_DQ_T_FLT	The output rotor flux estimated structure calculated from the current model. The structure consists of the D and Q rotor flux components stored for the next steps. The quadrature component is forced to zero value - required by FOC.
sPsiRotSAIBe		GMCLIB_2COOR_ALBE_T_FLT	The output rotor flux estimated structure calculated from the voltage model. The structure consists of the alpha and beta rotor flux components stored for the next steps.
sPsiStatSAIBe		GMCLIB_2COOR_ALBE_T_FLT	The output stator flux estimated structure calculated from the voltage model. The structure consists of the alpha and beta stator flux components stored for the next steps.
fltTorque		float_t	The output estimated motor torque calculated as: $T = \frac{3 \cdot P_p \cdot L_m \cdot (\Psi_{r\alpha} \cdot I_{s\beta} - \Psi_{r\beta} \cdot I_{s\alpha})}{2 \cdot I_{max}}$ The variable is a 32-bit single precision floating-point type value.
a32RotFluxPos		acc32_t	The output rotor flux estimated electric position (angle) - a 32-bit accumulator is normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π).
sCtrl	fltCompAlphaInteg_1	float_t	The state variable in the alpha part of the controller; integral part at step k-1.
	fltCompBetaInteg_1	float_t	The state variable in the beta part of the controller; integral part at step k-1.
	fltCompAlphaErr_1	float_t	The state variable in the alpha part of the controller; error part at step k-1.
	fltCompBetaErr_1	float_t	The state variable in the beta part of the controller; error part at step k-1.
	fltPGain	float_t	The proportional gain Kp for the stator model PI correction. Set by the user.
	fltIGain	float_t	The integration gain Ki for the stator model PI correction. Set by the user.
fltPsiRA1Gain		float_t	The gain is defined as: $\frac{\tau_r}{\tau_r + T_s} \text{ where: } \tau_r = \frac{L_r}{R_r}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltPsiRB1Gain		float_t	The coefficient gain is defined as: $\frac{L_m \cdot T_s}{\tau_r} \text{ where: } \tau_r = \frac{L_r}{R_r}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.

Table continues on the next page...

Table continued from the previous page...

Variable name	Data type	Description
fltPsiSA1Gain	float_t	The gain is defined as: $\frac{1}{1 + T_s \cdot 2\pi \cdot f_{\text{integ}}}$ The f_{integ} is a cut-off frequency of a low-pass filter approximation of a pure integrator. The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltPsiSA2Gain	float_t	The coefficient gain is defined as: $\frac{T_s}{1 + T_s \cdot 2\pi \cdot f_{\text{integ}}}$ The f_{integ} is a cut-off frequency of a low-pass filter approximation of a pure integrator. The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltKrInvGain	float_t	The gain is defined as: $\frac{L_r}{L_m}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltKrLsTotLeakGain	float_t	The coefficient gain is defined as: $\frac{L_s \cdot L_r - L_m^2}{L_m}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltRsEst	float_t	The stator resistance parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltTorqueGain	float_t	The torque constant coefficient gain is defined as: $\frac{3 \cdot P_p \cdot L_m}{2 \cdot L_r}$ The P_p is a number of motor pole-pairs. The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.

2.2.3 Declaration

The available AMCLIB_ACIMRotFluxObsrvInit function has the following declarations:

```
void AMCLIB_ACIMRotFluxObsrvInit_FLT(AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT *psCtrl)
```

The available AMCLIB_ACIMRotFluxObsrv function has the following declarations:

```
void AMCLIB_ACIMRotFluxObsrv_FLT(const GMCLIB_2COORD_ALBE_T_FLT *psISAlBe, const
GMCLIB_2COORD_ALBE_T_FLT *psUSAlBe, AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT *psCtrl)
```

2.2.4 Function use

The use of the `AMCLIB_ACIMRotFluxObsrv` function is shown in the following examples:

Floating-point version:

```
#include "amclib.h"

static GMCLIB_2COOR_ALBE_T_FLT sIsAlBe, sUsAlBe;
static AMCLIB_ACIM_ROT_FLUX_OBSRV_T_FLT sRfoParam;

void Isr(void);

void main (void)
{
    sRfoParam.sCtrl.fltPGain      = 32750.0F;
    sRfoParam.sCtrl.fltIGain      = 12500.0F;
    sRfoParam.fltKrInvGain        = 1.0851063829787235F;
    sRfoParam.fltKrLsTotLeakGain  = 0.08340425531914897F;
    sRfoParam.fltPsiRA1Gain       = 0.995151077592515F;
    sRfoParam.fltPsiRB1Gain       = 0.002278993531517996F;
    sRfoParam.fltPsiSA1Gain       = 0.9981185907806752F;
    sRfoParam.fltPsiSA2Gain       = 0.00009981185907806752F;
    sRfoParam.fltRsEst            = 26.06F;

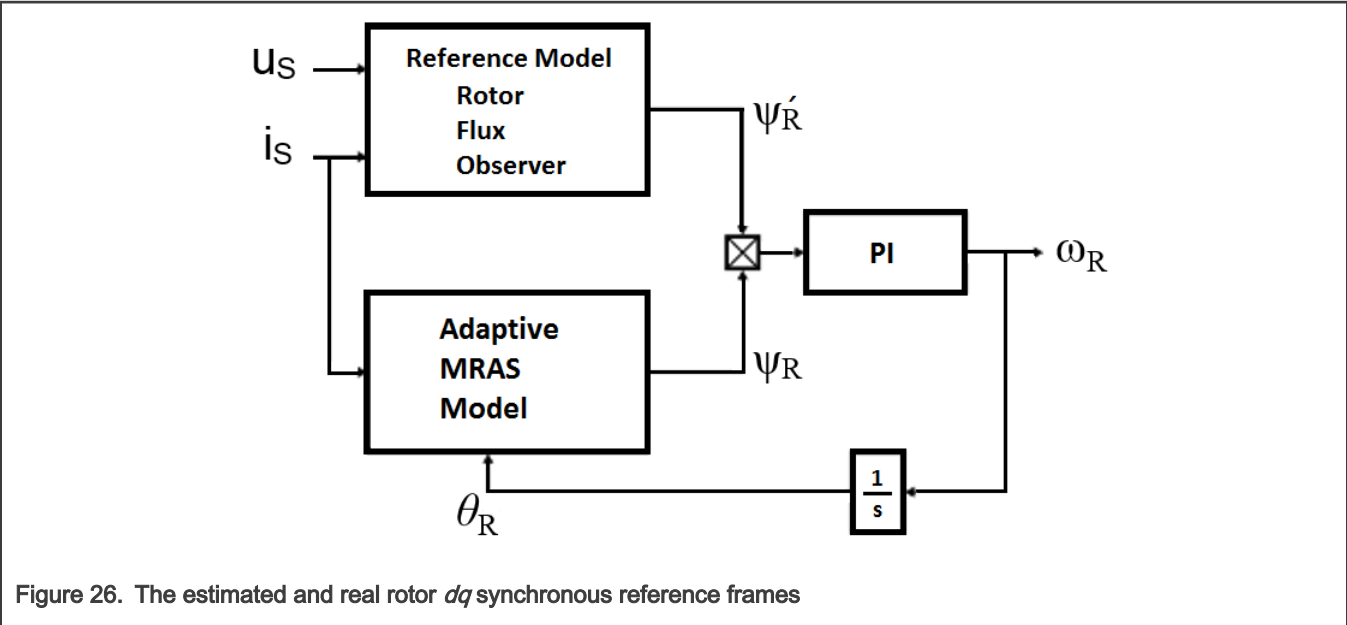
    /* Initialization of the RFO's structure */
    AMCLIB_ACIMRotFluxObsrvInit_FLT (&sRfoParam);

    sIsAlBe.fltAlpha = 0.05F;
    sIsAlBe.fltBeta  = 0.1F;
    sUsAlBe.fltAlpha = 0.2F;
    sUsAlBe.fltBeta  = -0.1F;
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Rotor flux observer calculation */
    AMCLIB_ACIMRotFluxObsrv_FLT(&sIsAlBe, &sUsAlBe, &sRfoParam);
}
```

2.3 AMCLIB_ACIMSpeedMRAS

The `AMCLIB_ACIMSpeedMRAS` function is based on the model reference approach (MRAS), and it uses the redundancy of two machine models of different structures that estimate the same state variable based on different sets of input variables. It means that the rotor speed can be obtained using an estimator with MRAS principle, in which the error vector is formed from the outputs of two models (both dependent on different motor parameters) - as shown in [Figure 1](#).



The closed-loop flux observer provides a stationary-axis-based rotor flux Ψ_R from RFO as a reference for the MRAS model, whereas the adaptive model of MRAS is the current-mode flux observer, which provides adjustable stationary-axis-based rotor flux:

$$\frac{d\vec{\psi}_r^{MRAS}}{dt} = -\frac{1}{\tau_r} \cdot \vec{\psi}_r^{MRAS} + \frac{L_m}{\tau_r} \vec{i}_s$$

where:

- i_s is the stator current vector
- Ψ_r is the rotor flux-linkage vector
- ω_r is the rotor electrical angular speed
- τ_r is the rotor electrical time constant
- L_m is the mutual equivalent inductance

The phase angle between the two estimated rotor flux vectors is used to correct the adaptive model, according to:

$$e_{MRAS} = \vec{\psi}_{ra}^{RFO} \cdot \vec{\psi}_{r\beta}^{MRAS} - \vec{\psi}_{r\beta}^{RFO} \cdot \vec{\psi}_{ra}^{MRAS}$$

The estimated speed ω_R is adjusted by a PI regulator.

2.3.1 Available versions

The available versions of the [AMCLIB_ACIMSpeedMRAS](#) function are shown in the following table:

Table 6. Init version

Function name	Parameters	Result type
AMCLIB_ACIMSpeedMRASInit_FLT	AMCLIB_ACIMSpeedMRAS_T_FLT *	void
The initialization does not have an input.		

Table 7. Function version

Function name	Input/output type		Result type
AMCLIB_ACIMSpeedMRAS_FLT	Input	GMCLIB_2COOR_ALBE_T_FLT *	void
		GMCLIB_2COOR_ALBE_T_FLT *	
		acc32_t	
	Parameters	AMCLIB_ACIMSpeedMRAS_T_FLT *	
	The AMCLIB_ACIMSpeedMRAS_FLT function with a 32-bit single precision floating-point inputs: stator current and voltage in alpha-beta coordinates.		

2.3.2 AMCLIB_ACIMSpeedMRAS_T_FLT type description

Variable name		Data type	Description
sSpeedIIR1Param		GDFLIB_FILTER_IIR1_T_FLT	The IIR1 filter structure for estimated speed filtration. Set by the user.
sPsiRotRDQ		GMCLIB_2COOR_DQ_T_FLT	The output rotor flux estimated structure from the current model. The structure consists of the D and Q rotor flux components stored for the next step.
fltSpeed		float_t	The output rotor estimated electrical speed.
fltSpeedEIIR1		float_t	The output rotor estimated electrical speed filtered.
fltSpeedMeIIR1		float_t	The output rotor estimated mechanical speed filtered.
a32RotPos		acc32_t	The output rotor estimated electric position (angle) - a 32-bit accumulator is normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <- π ; π).
sCtrl	fltSpeedInteg_1	float_t	The speed integral part - state variable at step k-1 of the electrical speed controller.
	fltSpeedErr_1	float_t	The speed error - state variable at step k-1 of the electrical speed controller.
	fltPGain	float_t	The MRAS proportional gain coefficient. Set by the user.
	fltIGain	float_t	The MRAS integral gain coefficient. Set by the user.
fltPsiRA1Gain		float_t	The coefficient gain is defined as: $\frac{\tau_r}{\tau_r + T_s} \text{ where: } \tau_r = \frac{L_r}{R_r}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltPsiRB1Gain		float_t	The coefficient gain is defined as: $\frac{L_m \cdot T_s}{\tau_r} \text{ where: } \tau_r = \frac{L_r}{R_r}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.

Table continues on the next page...

Table continued from the previous page...

Variable name	Data type	Description
fltTs	float_t	The sample time constant - the time between the steps. Set by the user.
fltSpeedMeGain	float_t	The speed gain coefficient, defined as: $\frac{60}{2\pi \cdot P_p}$ Where P_p is the number of motor pole-pairs. The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.

2.3.3 Declaration

The available AMCLIB_ACIMSpeedMRASInit function have the following declarations:

```
void AMCLIB_ACIMSpeedMRASInit_FLT(AMCLIB_ACIM_SPEED_MRAS_T_FLT *psCtrl)
```

The available AMCLIB_ACIMSpeedMRAS function have the following declarations:

```
void AMCLIB_ACIMSpeedMRAS_FLT(const GMCLIB_2COOR_ALBE_T_FLT *psISAlBe, const GMCLIB_2COOR_ALBE_T_FLT *psPsiRAlBe, acc32_t a32RotPos, AMCLIB_ACIM_SPEED_MRAS_T_FLT *psCtrl)
```

2.3.4 Function use

The use of the AMCLIB_ACIMSpeedMRAS function is shown in the following examples:

Floating-point version:

```
#include "amclib.h"

static GMCLIB_2COOR_ALBE_T_FLT sIsAlBe, sPsiRAlBe;
static AMCLIB_ACIM_SPEED_MRAS_T sMrasParam;
static acc32_t a32RotPosIn;

void Isr(void);

void main (void)
{
    sMrasParam.sCtrl.fltIGain = 12500.0F;
    sMrasParam.sCtrl.fltPGain = 32750.0F;
    sMrasParam.fltPsiRAlGain = 0.995151077592515F;
    sMrasParam.fltPsiRB1Gain = 0.002278993531517996F;
    sMrasParam.fltTs = 0.0001F;

    /* Initialization of the MRAS's structure */
    AMCLIB_ACIMSpeedMRASInit_FLT (&sMrasParam);

    sIsAlBe.fltAlpha = 0.05F;
    sIsAlBe.fltBeta = 0.1F;
```

```

sPsiRAlBe.fltAlpha = 0.2F;
sPsiRAlBe.fltBeta  = -0.1F;
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Speed estimation calculation based on MRAS */
    AMCLIB_ACIMSpeedMRAS_FLT(&sIsAlBe, &sPsiRAlBe, a32RotPosIn, &sMrasParam);
}

```

2.4 AMCLIB_AngleTrackObsrv

The [AMCLIB_TrackObsrv](#) function calculates an angle-tracking observer for determination of angular speed and position of the input signal. It requires two input arguments as sine and cosine samples. The practical implementation of the angle-tracking observer algorithm is described below.

The angle-tracking observer compares values of the input signals - $\sin(\theta)$, $\cos(\theta)$ with their corresponding estimations. As in any common closed-loop systems, the intent is to minimize the observer error towards zero value. The observer error is given here by subtracting the estimated resolver rotor angle from the actual rotor angle.

The tracking-observer algorithm uses the phase-locked loop mechanism. It is recommended to call this function at every sampling period. It requires a single input argument as phase error. A phase-tracking observer with standard PI controller used as the loop compensator is shown in [Figure 1](#).

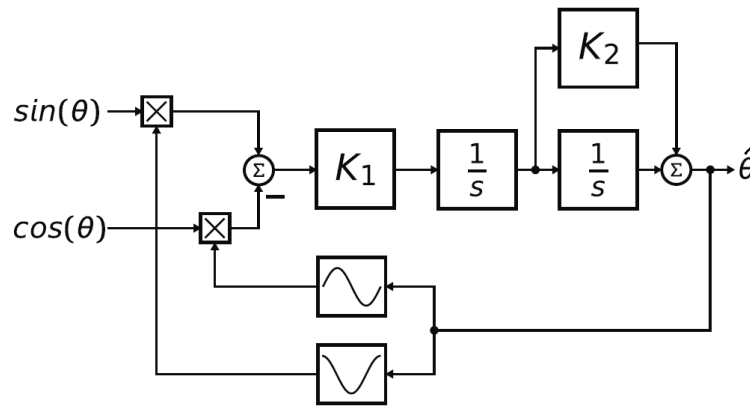


Figure 27. Block diagram of proposed PLL scheme for position estimation

Note that the mathematical expression of the observer error is known as the formula of the difference between two angles:

$$\sin(\theta - \hat{\theta}) = \sin(\theta) \cdot \cos(\hat{\theta}) - \cos(\theta) \cdot \sin(\hat{\theta})$$

If the deviation between the estimated and the actual angle is very small, then the observer error may be expressed using the following equation:

$$\sin(\theta - \hat{\theta}) \approx \theta - \hat{\theta}$$

The primary benefit of the angle-tracking observer utilization, in comparison with the trigonometric method, is its smoothing capability. This filtering is achieved by the integrator and the proportional and integral controllers, which are connected in series and closed by a unit feedback loop. This block diagram tracks the actual rotor angle and speed, and continuously updates their estimations. The angle-tracking observer transfer function is expressed as follows:

$$\frac{\hat{\theta}(s)}{\theta(s)} = \frac{K_I(1+sK_2)}{s^2+sK_IK_2+K_I}$$

The characteristic polynomial of the angle-tracking observer corresponds to the denominator of the following transfer function:

$$s^2+sK_IK_2+K_I$$

Appropriate dynamic behavior of the angle-tracking observer is achieved by the placement of the poles of characteristic polynomial. This general method is based on matching the coefficients of characteristic polynomial with the coefficients of a general second-order system.

The analog integrators in the previous figure (marked as $1/s$) are replaced by an equivalent of the discrete-time integrator using the backward Euler integration method. The discrete-time block diagram of the angle-tracking observer is shown in the following figure:

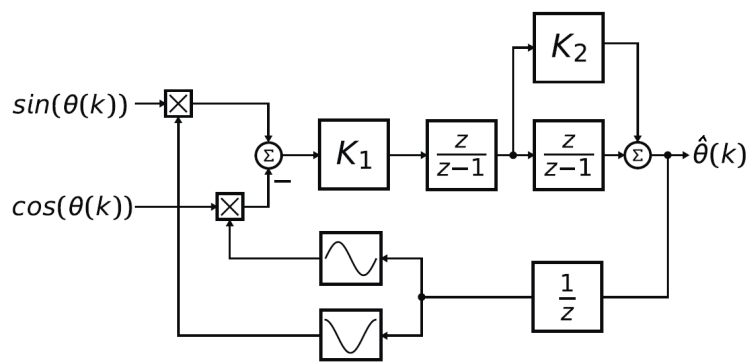


Figure 28. Block scheme of discrete-time tracking observer

The essential equations for implementing the angle-tracking observer (according to this block scheme) are as follows:

$$e(k) = \sin(\theta(k)) \cdot \cos(\hat{\theta}(k-1)) - \cos(\theta(k)) \cdot \sin(\hat{\theta}(k-1))$$

$$\omega(k) = T_s \cdot K_I \cdot e(k) + \omega(k-1)$$

$$a_2(k) = T_s \cdot \omega(k) + a_2(k-1)$$

$$\hat{\theta}(k) = K_2 \cdot \omega(k) + a_2(k)$$

where:

- K_I is the integral gain of the I controller
- K_2 is the proportional gain of the PI controller
- T_s is the sampling period [s]
- $e(k)$ is the position error in step k
- $\omega(k)$ is the rotor speed [rad / s] in step k
- $\omega(k-1)$ is the rotor speed [rad / s] in step k - 1
- $a(k)$ is the integral output of the PI controller [rad / s] in step k
- $a(k-1)$ is the integral output of the PI controller [rad / s] in step k - 1
- $\theta(k)$ is the rotor angle [rad] in step k

- $\theta(k-1)$ is the rotor angle [rad] in step $k-1$
- $\theta(k)$ is the estimated rotor angle [rad] in step k
- $\theta(k-1)$ is the estimated rotor angle [rad] in step $k-1$

In the fractional arithmetic, [AMCLIB_AngleTrackObsrv_Eq5](#) to [AMCLIB_AngleTrackObsrv_Eq8](#) are as follows:

$\omega_{sc}(k) \cdot \omega_{max} = T_s \cdot K_I \cdot e(k) + \omega_{sc}(k-1) \cdot \omega_{max}$
$a_{2sc}(k) \cdot \theta_{max} = T_s \cdot \omega_{sc}(k) \cdot \omega_{max} + a_{2sc}(k-1) \cdot \theta_{max}$
$\hat{\theta}_{sc}(k) \cdot \theta_{max} = K_2 \cdot \omega_{sc}(k) \cdot \omega_{max} + a_{2sc}(k) \cdot \theta_{max}$

where:

- $e_{sc}(k)$ is the scaled position error in step k
- $\omega_{sc}(k)$ is the scaled rotor speed [rad / s] in step k
- $\omega_{sc}(k-1)$ is the scaled rotor speed [rad / s] in step $k-1$
- $a_{sc}(k)$ is the integral output of the PI controller [rad / s] in step k
- $a_{sc}(k-1)$ is the integral output of the PI controller [rad / s] in step $k-1$
- $\theta_{sc}(k)$ is the scaled rotor angle [rad] in step k
- $\theta_{sc}(k-1)$ is the scaled rotor angle [rad] in step $k-1$
- $\theta_{sc}(k)$ is the scaled rotor angle [rad] in step k
- $\theta_{sc}(k-1)$ is the scaled rotor angle [rad] in step $k-1$
- ω_{max} is the maximum speed
- θ_{max} is the maximum rotor angle (typically π)

2.4.1 Available versions

The function is available in the following versions:

- Fractional output - the output is the fractional portion of the result; the result is within the range $<-1; 1)$.
- Accumulator output with floating point inputs - the output is the accumulator type, where the inputs for the calculation are the floating-point types within the range $<-1.0; 1.0>$.

The available versions of the [AMCLIB_AngleTrackObsrv](#) function are shown in the following table:

Table 8. Init versions

Function name	Init angle	Parameters	Result type
AMCLIB_AngleTrackObsrvInit_F16	frac16_t	AMCLIB_ANGLE_TRACK_OBSRV_T_F32 *	void
	The input is a 16-bit fractional value of the angle normalized to the range $<-1; 1)$ that represents an angle in (radians) within the range $<-\pi; \pi)$.		
AMCLIB_AngleTrackObsrvInit_A32af	acc32_t	AMCLIB_ANGLE_TRACK_OBSRV_T_FLT *	void
	The input is a 32-bit accumulator value of the angle divided by π .		

Table 9. Function versions

Function name	Input type	Parameters	Result type
AMCLIB_AngleTrackObsrv_F16	GMCLIB_2COOR_SINCOS_T_F16 *	AMCLIB_ANGLE_TRACK_OBSRV_T_F32 *	frac16_t
	Angle-tracking observer with a two-component (sin/cos) 16-bit fractional position input within the range <-1 ; 1). The output from the observer is a 16-bit fractional position normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π).		
AMCLIB_AngleTrackObsrv_A32ff	GMCLIB_2COOR_SINCOS_T_FLT *	AMCLIB_ANGLE_TRACK_OBSRV_T_FLT *	acc32_t
	Tracking observer with a two-component (sin/cos) 32-bit accumulator position input within the range <-1.0 ; 1.0>. The output from the observer is a 32-bit accumulator position normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π).		

2.4.2 AMCLIB_ANGLE_TRACK_OBSRV_T_F32

Variable name	Input type	Description
f32Speed	frac32_t	Estimated speed as the output of the first numerical integrator. The parameter is within the range <-1 ; 1). Controlled by the AMCLIB_AngleTrackObsrv_F16 algorithm; cleared by the AMCLIB_AngleTrackObsrvInit_F16 function.
f32A2	frac32_t	Output of the second numerical integrator. The parameter is within the range <-1 ; 1). Controlled by the AMCLIB_AngleTrackObsrv_F16 and AMCLIB_AngleTrackObsrvInit_F16 algorithms.
f16Theta	frac16_t	Estimated position as the output of the observer. The parameter is normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π). Controlled by the AMCLIB_AngleTrackObsrv_F16 and AMCLIB_AngleTrackObsrvInit_F16 algorithms.
f16SinEstim	frac16_t	Sine of the estimated position as the output of the actual step. Keeps the sine of the position for the next step. The parameter is within the range <-1 ; 1). Controlled by the AMCLIB_AngleTrackObsrv_F16 and AMCLIB_AngleTrackObsrvInit_F16 algorithms.
f16CosEstim	frac16_t	Cosine of the estimated position as the output of the actual step. Keeps the cosine of the position for the next step. The parameter is within the range <-1 ; 1). Controlled by the AMCLIB_AngleTrackObsrv_F16 and AMCLIB_AngleTrackObsrvInit_F16 algorithms.
f16K1Gain	frac16_t	Observer K1 gain is set up according to Equation 9 as: $T_s \cdot K_I \cdot \frac{1}{\omega_{max}} \cdot 2^{-K1sh}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16K1GainSh	int16_t	Observer K2 gain shift takes care of keeping the f16K1Gain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(T_s \cdot K_I \cdot \frac{1}{\omega_{max}}) - \log_2 1 < K1sh \leq \log_2(T_s \cdot K_I \cdot \frac{1}{\omega_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.

Table continues on the next page...

Table continued from the previous page...

Variable name	Input type	Description
f16K2Gain	frac16_t	Observer K2 gain is set up according to Equation 11 as: $K_2 \cdot \frac{\omega_{max}}{\theta_{max}} \cdot 2^{-K2sh}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16K2GainSh	int16_t	Observer K2 gain shift takes care of keeping the f16K2Gain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(K_2 \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 1 < K2sh \leq \log_2(K_2 \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.
f16A2Gain	frac16_t	Observer A2 gain for the output position is set up according to Equation 10 as: $T_s \cdot \frac{\omega_{max}}{\theta_{max}} \cdot 2^{-A2sh}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16A2GainSh	int16_t	Observer A2 gain shift for the position integrator takes care of keeping the f16A2Gain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(T_s \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 1 < A2sh \leq \log_2(T_s \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.

2.4.3 AMCLIB_ANGLE_TRACK_OBSRV_T_FLT

Variable name	Input type	Description
fltSpeed	float_t	Estimated speed as the output of the first numerical integrator. The parameter is within the range <-32768.0; 32767.99998). Controlled by the AMCLIB_AngleTrackObsrv_A32ff algorithm; cleared by AMCLIB_AngleTrackObsrvInit_A32af function.
f32A2	frac32_t	Output of the second numerical integrator. The parameter is within the range <-1 ; 1). Controlled by the AMCLIB_AngleTrackObsrv_A32ff and AMCLIB_AngleTrackObsrvInit_A32af algorithms.
a32Theta	acc32_t	Estimated position as the output of the observer. The parameter is normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π). Controlled by the AMCLIB_AngleTrackObsrv_A32ff and AMCLIB_AngleTrackObsrvInit_A32af algorithms.
fltSinEstim	float_t	Sine of the estimated position as the output of the actual step. Keeps the sine of the position for the next step. The parameter is within the range <-1 ; 1>. Controlled by the AMCLIB_AngleTrackObsrv_A32ff and AMCLIB_AngleTrackObsrvInit_A32af algorithms.
fltCosEstim	float_t	Cosine of the estimated position as the output of the actual step. Keeps the cosine of the position for the next step. The parameter is within the range <-1 ; 1>. Controlled by the AMCLIB_AngleTrackObsrv_A32ff and AMCLIB_AngleTrackObsrvInit_A32af algorithms.

Table continues on the next page...

Table continued from the previous page...

Variable name	Input type	Description
fltK1Gain	float_t	Observer K1 gain is set up according to Equation 6 as: K_1T_s . The parameter is a 32-bit single precision floating-point value in range (0; 16383.99999). Set by the user.
fltK2Gain	float_t	Observer K2 gain is set up according to Equation 8 as: K_2 . The parameter is a 32-bit single precision floating-point value in range (0; 65535.9999689999). Set by the user.
fltA2Gain	float_t	Observer A2 gain for the output position is set up according to Equation 7 as: T_s . The parameter is a 32-bit single precision floating-point value in range (0; 65535.9999689999). Set by the user.

2.4.4 Declaration

The available AMCLIB_AngleTrackObsrvInit functions have the following declarations:

```
void AMCLIB_AngleTrackObsrvInit_Fl6(frac16_t f16ThetaInit, AMCLIB_ANGLE_TRACK_OBSRV_T_F32 *psCtrl)

void AMCLIB_AngleTrackObsrvInit_A32ff(acc32_t a32ThetaInit, AMCLIB_ANGLE_TRACK_OBSRV_T_FLT *psCtrl)
```

The available AMCLIB_AngleTrackObsrv functions have the following declarations:

```
frac16_t AMCLIB_AngleTrackObsrv_Fl6(const GMCLIB_2COOR_SINCOS_T_F16 *psAnglePos,
AMCLIB_ANGLE_TRACK_OBSRV_T_F32 *psCtrl)

acc32_t AMCLIB_AngleTrackObsrv_A32ff(const GMCLIB_2COOR_SINCOS_T_FLT *psAnglePos,
AMCLIB_ANGLE_TRACK_OBSRV_T_FLT *psCtrl)
```

2.4.5 Function use

The use of the AMCLIB_AngleTrackObsrvInit and AMCLIB_AngleTrackObsrv functions is shown in the following example:

```
#include "amclib.h"

static AMCLIB_ANGLE_TRACK_OBSRV_T_F32 sAto;
static GMCLIB_2COOR_SINCOS_T_F16 sAnglePos;
static frac16_t f16PositionEstim, f16PositionInit;

void Isr(void);

void main(void)
{
    sAto.f16K1Gain = FRAC16(0.6434);
    sAto.i16K1GainSh = -9;
    sAto.f16K2Gain = FRAC16(0.6801);
```

```

sAto.il6K2GainSh = -2;
sAto.fl6A2Gain   = FRAC16(0.6400);
sAto.il6A2GainSh = -4;

fl6PositionInit = FRAC16(0.0);

AMCLIB_AngleTrackObsrvInit_Fl6(fl6PositionInit, &sAto);

sAnglePos.fl6Sin = FRAC16(0.0);
sAnglePos.fl6Cos = FRAC16(1.0);
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Angle tracking observer calculation */
    fl6PositionEstim = AMCLIB_AngleTrackObsrv_Fl6(&sAnglePos, &sAto);
}

```

2.5 AMCLIB_CtrlFluxWkng

The [AMCLIB_CtrlFluxWkng](#) function controls the motor magnetizing flux for a speed exceeding above the nominal speed of the motor. Where a higher maximum motor speed is required, the flux (field) weakening technique must be used. The basic task of the function is to maintain the motor magnetizing flux below the nominal level which does not require a higher supply voltage when the motor rotates above the nominal motor speed. The lower magnetizing flux is provided by maintaining the flux-producing current component i_D in the flux-weakening region, as shown in [Figure 1](#)).

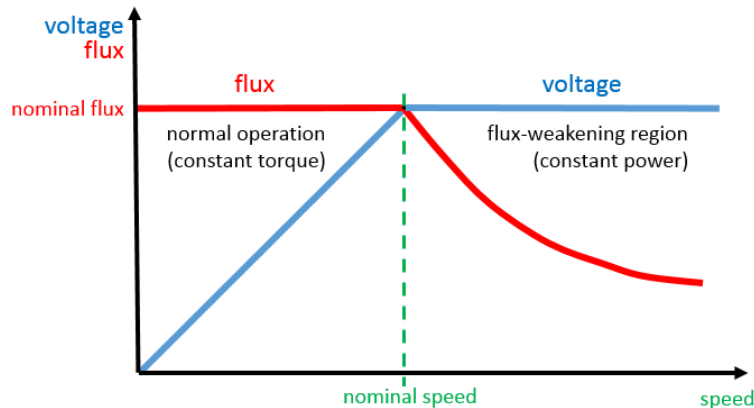


Figure 29. Flux weakening operating range

The [AMCLIB_CtrlFluxWkng](#) function processes the magnetizing flux by the PI controller function with the anti-windup functionality and output limitation. The controller integration can be stopped if the system is saturated by the input flag pointer in the flux-weakening controller structure. The flux-weakening controller algorithm is executed in the following steps:

1. The voltage error calculation from the voltage limit and the required voltage.

$$u_{err} = (u_{QLim} - |u_{Qreq}|) \cdot \frac{I_{gain}}{U_{gain}}$$

Figure 30.

where:

- u_{err} is the voltage error

- u_{QLim} is the Q voltage limit component
 - u_{Qreq} is the Q required voltage component
 - I_{gain} is the voltage scale - max. value (for fraction gain = 1)
 - U_{gain} is the current scale - max. value (for fraction gain = 1)
2. The input Q current error component must be positive and filtered by the infinite impulse response first-order filter.

$$i_{QerrIIR} = IIR1(|i_{Qerr}|)$$

Figure 31.

where:

- $i_{QerrIIR}$ is the Q current error component filtered by the first-order IIR
 - i_{Qerr} is the input Q current error component (calculated before calling the [AMCLIB_CtrFluxWkng](#) function from the measured and limited required Q current component value).
3. The flux error is obtained from the previously calculated voltage and current errors as follows:

$$i_{err} = i_{QerrIIR} \cdot u_{err}$$

Figure 32.

where:

- i_{err} is the Q current error component for the flux PI controller
 - $i_{QerrIIR}$ is the current error component filtered by the first-order IIR
 - u_{err} is the voltage error for the flux PI controller
4. Finally, the flux error (corresponding the I_D) is processed by the flux PI controller:

$$i_{Dreq} = CtrlPIpAW(i_{err})$$

Figure 33.

where:

- i_{Dreq} is the required D current component for the current control
- i_{err} is the flux error (corresponding the D current component) for the flux PI controller

The controller output should be used as the required D current component in the fast control loop and concurrently used as an input for the `GFLIB_VectorLimit1` function which limits the I_Q controller as follows:

$$i_{Qreq} \leq \sqrt{i_{max}^2 - i_{Dreq}^2}$$

Figure 34.

where:

- i_{Qreq} is the required Q current component for the current control
- i_{max} is application current limit
- i_{Dreq} is the required D current component for the current control

The following figure shows an example of applying the flux-weakening controller function in the control structure. The flux controller starts to operate when the I_Q controller is not able to compensate the I_{Qerr} and creates a deviation between its input and output. The flux controller processes the deviation and decreases the flux excitation (for ACIM, or starts to create the flux excitation against a permanent magnet flux in case of PMSM). A lower BEMF causes a higher I_Q and the motor speed increases. The speed controller with I_{Qreg} on the output should be limited by the vector limit1 function because a part of the current is used for flux excitation.

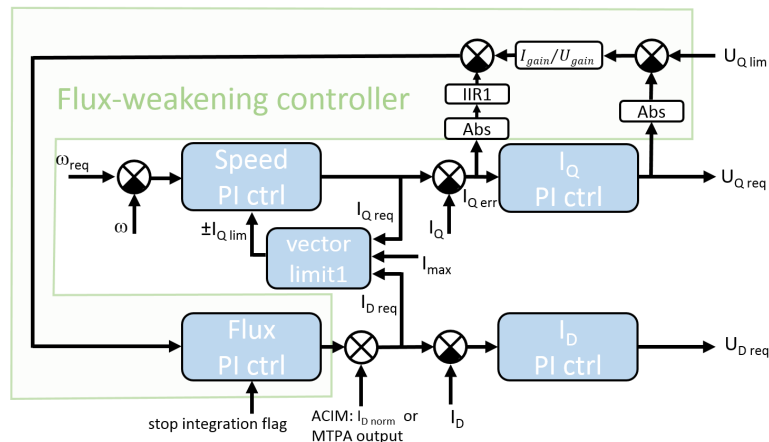


Figure 35. Flux weakening function in control block structure

2.5.1 Available versions

This function is available in the following versions:

- Fractional output - the output is the fractional portion of the result; the result is within the range $-1 ; 1$ in case of no limitation. The parameters are of fractional or accumulator types.
- Floating-point output - the output is the floating-point result within the type's full range in case of no limitation. The parameters are of a floating-point type as well.

The available versions of the AMCLIB_CtrlFluxWknglInit function are shown in the following table:

Table 10. Init function versions

Function name	Input type	Parameters	Result type
AMCLIB_CtrlFluxWkngInit_F16	frac16_t	AMCLIB_CtrlFluxWkngInit_A32*	void
	The inputs are a 16-bit fractional initial value for the flux PI controller integrating the part state and a pointer to the flux-weakening controller's parameters structure. The function initializes the flux PI controller and the IIR1 filter.		
AMCLIB_CtrlFluxWkngInit_FLT	float_t	AMCLIB_CtrlFluxWkngInit_FLT*	void
	The inputs are a 32-bit single precision floating-point initial value for the flux PI controller integrating the part state and a pointer to the flux-weakening controller's parameters structure. The function initializes the flux PI controller and the IIR1 filter.		

The available versions of the [AMCLIB_CtrlFluxWkng](#) function are shown in the following table:

Table 11. Function versions

Function name	Input type			Parameters	Result type
	Q current error	Q required voltage	Q voltage limit		
AMCLIB_CtrlFluxWkn g_F16	frac16_t	frac16_t	frac16_t	AMCLIB_CTRL_FLUX_WKNG_T_A32*	frac16_t
	The Q current error component value input (I _Q controller input) and the Q required voltage value input (I _Q controller output) are 16-bit fractional values within the range <-1 ; 1). The Q voltage limit value input				

Table continues on the next page...

Table 11. Function versions (continued)

Function name	Input type			Parameters	Result type
	Q current error	Q required voltage	Q voltage limit		
	(constant value) is a 16-bit fractional value within the range (0 ; 1). The parameters are pointed to by an input pointer. The function returns a 16-bit fractional value in the range <f16LowerLim ; f16UpperLim>.				
AMCLIB_CtrlFluxWkn g_FLT	float_t	float_t	float_t	AMCLIB_CTRL_FLUX_WKNG_T_FLT*	float_t
	The Q current error component value input (I _Q controller input) is a 32-bit single precision floating-point value within the full type's range. The Q required voltage value input (I _Q controller output) is a 32-bit single precision floating-point value within the full type's range.The Q voltage limit value (constant value) is a 32-bit single precision floating-point positive value. The parameters are pointed to by an input pointer. The function returns a 32-bit single precision floating-point value in the range <fltLowerLim ; fltUpperLim>.				

2.5.2 AMCLIB_CTRL_FLUX_WKNG_T_A32

Variable name	Input type	Description
sFWPiParam	GFLIB_CTRL_PI_P_AW_T_A32	The input pointer for the flux PI controller parameter structure. The flux controller output should be negative. Therefore, set at least the following parameters: • a32PGain - proportional gain, the range is <0 ; 65536.0>. • a32IGain - integral gain, the range is <0 ; 65536.0>. • f16UpperLim - upper limit, the zero value should be set. • f16LowerLim - the lower limit, the range is <-1; 0>.
slqErrIIR1Param	GDFLIB_FILTER_IIR1_T_F32	The input pointer for the IIR1 filter parameter structure. The IIR1 filters the absolute value of the Q current error component for the flux controller. Set at least the following parameters: • sFltCoeff.f32B0 - B0 coefficient, must be divided by 2. • sFltCoeff.f32B1 - B1 coefficient, must be divided by 2. • sFltCoeff.f32A1 - A1 (sign-inverted) coefficient, must be divided by -2 (negative two).
f16lqErrIIR1	frac32_t	The I_Q current error component, filtered by the IIR1 filter for the flux PI controller, as shown in Equation 2. The output value calculated by the algorithm.
f16UFWErr	frac16_t	The voltage error, as shown in Equation 1. The output value calculated by the algorithm.
f16FWErr	frac16_t	The flux-weakening error, as shown in Equation 3. The output value calculated by the algorithm.
*bStopIntegFlag	frac16_t	The integration of the PI controller is suspended if the stop flag is set. When it is cleared, the integration continues. The pointer is set by the user and controlled by the application.

2.5.3 AMCLIB_CTRL_FLUX_WKNG_T_FLT

Variable name	Input type	Description
sFWPiParam	GFLIB_CTRL_PI_P_AW_T_FLT	The input pointer for the flux PI controller parameter structure. The flux controller output should be negative. Therefore, set at least the following parameters: - fltPGain - the proportional gain, the parameter is a 32-bit single precision floating-point type non-negative value. - fltIGain - the integral gain, the parameter is a 32-bit single precision floating-point type non-negative value. - fltUpperLim - the upper limit, the zero value should be set. - fltLowerLim - the lower limit, the parameter is a 32-bit single precision floating-point type positive value.
slqErrIIR1Param	GDFLIB_FILTER_IIR1_T_FLT	The input pointer for the IIR1 filter parameter structure. The IIR1 filters the absolute value of the Q current error component for the flux controller. Set at least the following parameters: - sFltCoeff.fltB0 - B0 coefficient. - sFltCoeff.fltB1 - B1 coefficient. - sFltCoeff.fltA1 - A1 coefficient.
fltIqErrIIR1	float_t	The I_Q current error, filtered by the IIR1 filter for the flux PI controller, as shown in Equation 2. The output value calculated by the algorithm.
fltUFWErr	float_t	The voltage error, as shown in Equation 1. The output value calculated by the algorithm.
fltFWErr	float_t	The flux-weakening error, as shown in Equation 3. The output value calculated by the algorithm.
fltIGainUgain	float_t	The current/voltage scale, calculated according to: $fltIGainUgain = \frac{I_{gain}}{U_{gain}}$ Set by the user.
*bStopIntegFlag	float_t	The integration of the flux PI controller is suspended if the input stop flag is set. When it is cleared, the integration continues. The pointer is set by the user and controlled by the application.

2.5.4 Declaration

The available AMCLIB_CtrlFluxWkngInit functions have the following declarations:

```
void AMCLIB_CtrlFluxWkngInit_F16(frac16_t f16InitVal, AMCLIB_CTRL_FLUX_WKNG_T_A32 *psParam)
```

```
void AMCLIB_CtrlFluxWkngInit_FLT(float_t fltInitVal, AMCLIB_CTRL_FLUX_WKNG_T_FLT *psParam)
```

The available [AMCLIB_CtrlFluxWkng](#) functions have the following declarations:

```
frac16_t AMCLIB_CtrlFluxWkng_F16(frac16_t f16IQErr, frac16_t f16UQReq, frac16_t f16UQLim,
AMCLIB_CTRL_FLUX_WKNG_T_A32 *psParam)
```

```
float_t AMCLIB_CtrlFluxWkng_FLT(float_t fltIQErr, float_t fltUQReq, float_t fltUQLim,
AMCLIB_CTRL_FLUX_WKNG_T_FLT *psParam)
```

2.5.5 Function use

The use of the [AMCLIB_CtrlFluxWkngInit](#) and [AMCLIB_CtrlFluxWkng](#) functions is shown in the following examples:

Fixed-point version:

```
#include "amclib.h"

static AMCLIB_CTRL_FLUX_WKNG_T_A32 sCtrl;
static frac16_t f16IQErr, f16UQReq, f16UQLim;
static frac16_t f16IdReq, f16InitVal;
static bool_t bStopIntegFlag;

void Isr(void);

void main(void)
{
    /* Associate input stop integration flag */
    bStopIntegFlag = FALSE;
    sCtrl.bStopIntegFlag = &bStopIntegFlag;

    /* Set PI controller and IIR1 parameters */
    sCtrl.sFWPiParam.a32PGain = ACC32(0.1);
    sCtrl.sFWPiParam.a32IGain = ACC32(0.2);
    sCtrl.sFWPiParam.f16UpperLim = FRAC16(0.);
    sCtrl.sFWPiParam.f16LowerLim = FRAC16(-0.9);
    sCtrl.sIQErrIIRParam.sFltCoeff.f32B0 = FRAC32(0.245237275252786 / 2.0);
    sCtrl.sIQErrIIRParam.sFltCoeff.f32B1 = FRAC32(0.245237275252786 / 2.0);
    sCtrl.sIQErrIIRParam.sFltCoeff.f32A1 = FRAC32(-0.509525449494429 / -2.0);

    /* Flux weakening controller initialization */
    f16InitVal = FRAC16(0.0);
    AMCLIB_CtrlFluxWkngInit_F16(f16InitVal, &sCtrl);

    /* Assign input variable */
    f16IQErr = FRAC16(-0.1);
    f16UQReq = FRAC16(-0.2);
    f16UQLim = FRAC16(0.8);
}

/* Periodical function or interrupt */
void Isr()
{
    /* Flux weakening controller calculation */
    f16Result = AMCLIB_CtrlFluxWkng_F16(f16IQErr, f16UQReq, f16UQLim, &sCtrl);
}
```

Floating-point version:

```
#include "amclib.h"

static AMCLIB_CTRL_FLUX_WKNG_T_FLT sCtrl;
static float_t fltIQErr, fltUQReq, fltUQLim;
static float_t fltIdReq, fltInitVal;
static bool_t bStopIntegFlag;

void Isr(void);

void main(void)
{
    /* Associate input stop integration flag */
    bStopIntegFlag = FALSE;
    sCtrl.bStopIntegFlag = &bStopIntegFlag;

    /* Set PI controller and IIR1 parameters */
    sCtrl.sFWPiParam.fltPGain = 0.1F;
    sCtrl.sFWPiParam.fltIGain = 0.2F;
    sCtrl.sFWPiParam.fltUpperLim = 0.0F;
    sCtrl.sFWPiParam.fltLowerLim = -0.9F;
    sCtrl.sIQErrIIR1Param.sFltCoeff.fltB0 = 0.245237275252786f;
    sCtrl.sIQErrIIR1Param.sFltCoeff.fltB1 = 0.245237275252786f;
    sCtrl.sIQErrIIR1Param.sFltCoeff.fltA1 = -0.509525449494429f;

    /* Flux weakening controller initialization */
    fltInitVal = 0.0F;
    AMCLIB_CtrlFluxWkngInit_FLT(fltInitVal, &sCtrl);

    /* Assign input variable */
    fltIQErr = -0.1F;
    fltUQReq = -0.2F;
    fltUQLim = 0.8F;
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Flux weakening controller calculation */
    fltIdReq = AMCLIB_CtrlFluxWkng_FLT(fltIQErr, fltUQReq, fltUQLim, &sCtrl);
}
```

2.6 AMCLIB_PMSMBemfObsrvAB

The [AMCLIB_PMSMBemfObsrvAB](#) function calculates the algorithm of the back-electro-motive force (back-EMF) observer in a stationary reference frame. The estimation method for the rotor position and the angular speed is based on the mathematical model of an interior PMSM motor with an extended electro-motive force function, which is realized in the alpha/beta stationary reference frame.

The back-EMF observer detects the generated motor voltages, induced by the permanent magnets. The angle-tracking observer uses the back-EMF signals to calculate the position and speed of the rotor. The transformed model is then derived as:

$$\begin{bmatrix} u_\alpha \\ u_\beta \end{bmatrix} = \begin{bmatrix} R_S + sL_D & \omega_r \Delta L \\ -\omega_r \Delta L & R_S + sL_D \end{bmatrix} \cdot \begin{bmatrix} i_\alpha \\ i_\beta \end{bmatrix} + [\Delta L \cdot (\omega_r i_D - s i_Q) + \Psi_m \omega_r] \cdot \begin{bmatrix} -\sin(\theta_r) \\ \cos(\theta_r) \end{bmatrix}$$

Where:

- R_S is the stator resistance
- L_D and L_Q are the D-axis and Q-axis inductances
- $\Delta L = L_D - L_Q$ is the motor saliency
- Ψ_m is the back-EMF constant
- ω_r is the angular electrical rotor speed
- u_α and u_β are the estimated stator voltages
- i_α and i_β are the estimated stator currents
- θ_r is the estimated rotor electrical position
- s is the operator of the derivative

This extended back-EMF model includes both the position information from the conventionally defined back-EMF and the stator inductance as well. This enables extracting the rotor position and velocity information by estimating the extended back-EMF only.

Both the alpha and beta axes consist of the stator current observer based on the RL motor circuit which requires the motor parameters.

The current observer input is the sum of the actual applied motor voltage and the cross-coupled rotational term, which corresponds to the motor saliency ($L_D - L_Q$) and the compensator corrective output. The observer provides the back-EMF signals as a disturbance because the back-EMF is not included in the observer model.

The block diagram of the observer in the estimated reference frame is shown in [Figure 1](#). The observer compensator is substituted by a standard PI controller with following equation in the fractional arithmetic.

$$i_{sc}(k) \cdot i_{max} = K_P \cdot e_{sc}(k) \cdot e_{max} + T_s \cdot K_I \cdot e_{sc}(k) \cdot e_{max} + i_{sc}(k-1) \cdot i_{max}$$

where:

- K_P is the observer proportional gain [-]
- K_I is the observer integral gain [-]
- $i_{sc}(k) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the actual step
- $i_{sc}(k-1) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the previous step
- $e_{sc}(k) = [e_\gamma, e_\delta]$ is the scaled stator back-EMF voltage vector in the actual step
- i_{max} is the maximum current [A]
- e_{max} is the maximum back-EMF voltage [V]
- T_s is the sampling time [s]

As shown in [Figure 1](#), the observer model and hence also the PI controller gains in both axes are identical to each other.

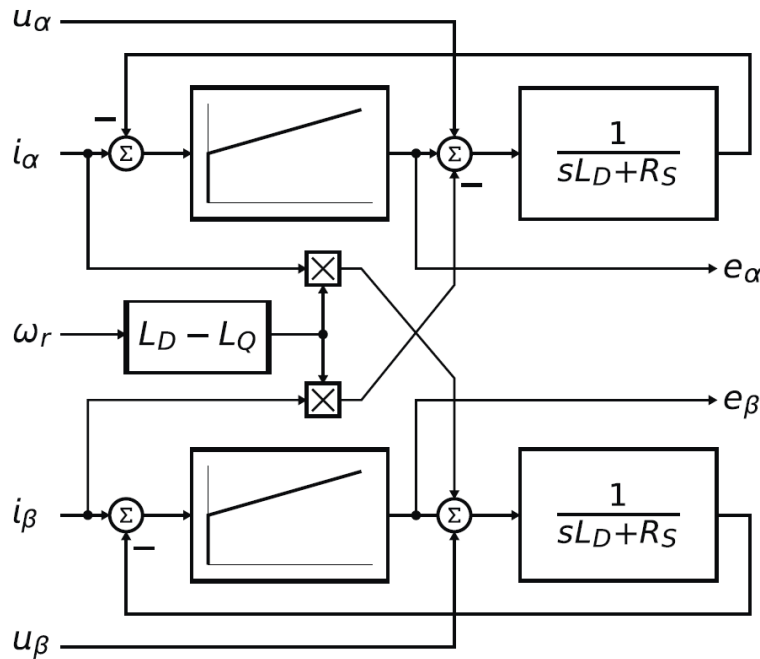


Figure 36. Block diagram of back-EMF observer

It is obvious that the accuracy of the back-EMF estimates is determined by the correctness of the motor parameters used (R , L), the fidelity of the reference stator voltage, and the quality of the compensator, such as the bandwidth, phase lag, and so on.

The appropriate dynamic behavior of the back-EMF observer is achieved by the placement of the poles of the stator current observer characteristic polynomial. This general method is based on matching the coefficients of the characteristic polynomial to the coefficients of the general second-order system.

$$\hat{E}_{\alpha\beta}(s) = -E_{\alpha\beta}(s) \cdot \frac{F_C(s)}{sL_D + R_S + F_C(s)}$$

The back-EMF observer is a Luenberger-type observer with a motor model, which is implemented using the backward Euler transformation as:

$$i(k) = \frac{T_s}{L_D + T_s R_S} \cdot u(k) + \frac{T_s}{L_D + T_s R_S} \cdot e(k) - \frac{\Delta L T_s}{L_D + T_s R_S} \cdot \omega_e(k) \cdot i'(k) + \frac{L_D}{L_D + T_s R_S} \cdot i(k-1)$$

Where:

- $i(k) = [i_\alpha, i_\beta]$ is the stator current vector in the actual step
- $i(k-1) = [i_\alpha, i_\beta]$ is the stator current vector in the previous step
- $u(k) = [u_\alpha, u_\beta]$ is the stator voltage vector in the actual step
- $e(k) = [e_\alpha, e_\beta]$ is the stator back-EMF voltage vector in the actual step
- $i'(k) = [i_\alpha, -i_\beta]$ is the complementary stator current vector in the actual step
- $\omega_e(k)$ is the electrical angular speed in the actual step
- T_s is the sampling time [s]

This equation is transformed into the fractional arithmetic as:

$$i_{sc}(k) \cdot i_{max} = \frac{T_s}{L_D + T_s R_S} \cdot u_{sc}(k) \cdot u_{max} + \frac{T_s}{L_D + T_s R_S} \cdot e_{sc}(k) \cdot e_{max} - \frac{\Delta L T_s}{L_D + T_s R_S} \cdot \omega_{esc}(k) \cdot \omega_{max} \cdot i'_{sc}(k) \cdot i_{max} + \frac{L_D}{L_D + T_s R_S} \cdot i_{sc}(k-1) \cdot i_{max}$$

Where:

- $i_{sc}(k) = [i_v, i_\delta]$ is the scaled stator current vector in the actual step
- $i_{sc}(k-1) = [i_v, i_\delta]$ is the scaled stator current vector in the previous step
- $u_{sc}(k) = [u_v, u_\delta]$ is the scaled stator voltage vector in the actual step
- $e_{sc}(k) = [e_v, e_\delta]$ is the scaled stator back-EMF voltage vector in the actual step
- $i'_{sc}(k) = [i_v, -i_\delta]$ is the scaled complementary stator current vector in the actual step
- $\omega_{esc}(k)$ is the scaled electrical angular speed in the actual step
- i_{max} is the maximum current [A]
- e_{max} is the maximum back-EMF voltage [V]
- u_{max} is the maximum stator voltage [V]
- ω_{max} is the maximum electrical angular speed in [rad / s]

If the Luenberger-type stator current observer is properly designed in the stationary reference frame, the back-EMF can be estimated as a disturbance produced by the observer controller. However, this is only valid when the back-EMF term is not included in the observer model. The observer is a closed-loop current observer, therefore, it acts as a state filter for the back-EMF term.

The estimate of the extended EMF term can be derived from [AMCLIB_PMSMBemfObsrvAB_Eq1](#) as:

$$-\frac{\hat{E}_{\gamma\phi}(s)}{E_{\gamma\phi}(s)} = \frac{sK_P + K_I}{s^2 L_D + sR_S + sK_P + K_I}$$

The observer controller can be designed by comparing the closed-loop characteristic polynomial to that of a standard second-order system as:

$$s^2 + \frac{K_P + R_S}{L_D} \cdot s + \frac{K_I}{L_D} = s^2 + 2\xi\omega_0 s + \omega_0^2$$

where:

- ω_0 is the natural frequency of the closed-loop system (loop bandwidth)
- ξ is the loop attenuation
- K_P is the proportional gain
- K_I is the integral gain

2.6.1 Available versions

This function is available in the following versions:

- Fractional output - the output is the fractional portion of the result; the result is within the range <-1 ; 1). The parameters use the accumulator types.
- Floating-point output - the output is the floating-point result within the type's full range.

The available versions of the [AMCLIB_PMSMBemfObsrvAB](#) function are shown in the following table:

Table 12. Init versions

Function name	Parameters	Result type
AMCLIB_PMSMBemfObsrvABInit_F16	AMCLIB_BEMF_OBSRV_AB_T_A32 *	void
	The initialization does not have an input.	
AMCLIB_PMSMBemfObsrvABInit_A32fff	AMCLIB_BEMF_OBSRV_AB_T_FLT *	void
	The initialization does not have an input.	

The available versions of the [AMCLIB_PMSMBemfObsrvAB](#) function are shown in the following table:

Table 13. Function versions

Function name	Input/output type		Result type
AMCLIB_PMSMBemfObsrvAB_F16	Input	GMCLIB_2COOR_ALBE_T_F16 *	void
		GMCLIB_2COOR_ALBE_T_F16 *	
		frac16_t	
	Parameters	AMCLIB_BEMF_OBSRV_AB_T_A32 *	
	The back-EMF observer with a 16-bit fractional input Alpha/Beta current and voltage, and a 16-bit electrical speed. All are within the range <-1 ; 1).		
AMCLIB_PMSMBemfObsrvAB_FLT	Input	GMCLIB_2COOR_ALBE_T_FLT *	void
		GMCLIB_2COOR_ALBE_T_FLT *	
		float_t	
	Parameters	AMCLIB_BEMF_OBSRV_AB_T_FLT *	
	The back-EMF observer with a 32-bit single precision floating-point input Alpha/Beta current and voltage, and a 32-bit single precision floating-point electrical speed. All are within the full range.		

2.6.2 AMCLIB_BEMF_OBSRV_AB_T_A32 type description

Variable name		Data type	Description
sEObsrv		GMCLIB_2COOR_ALBE_T_F32	The estimated back-EMF voltage structure.
sIObsrv		GMCLIB_2COOR_ALBE_T_F32	The estimated current structure.
sCtrl	f32IAlpha_1	frac32_t	The state variable in the alpha part of the observer, integral part at step k-1. The variable is within the range <-1 ; 1).
	f32IBeta_1	frac32_t	The state variable in the beta part of the observer, integral part at step k-1. The variable is within the range <-1 ; 1).
	a32PGain	acc32_t	The observer proportional gain is set up according to Equation 7 as: $(2\xi\omega_0 L_D - R_S) \frac{i_{max}}{e_{max}}$

Table continues on the next page...

Table continued from the previous page...

Variable name		Data type	Description
			The parameter is within the range <0 ; 65536.0). Set by the user.
	a32IGain	acc32_t	The observer integral gain is set up according to Equation 7 as: $\omega_0^2 L_D T_s \frac{i_{max}}{e_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32IGain		acc32_t	The current coefficient gain is set up according to Equation 5 as: $\frac{L_D}{L_D + T_s R_S}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32UGain		acc32_t	The voltage coefficient gain is set up according to Equation 5 as: $\frac{T_s}{L_D + T_s R_S} \cdot \frac{u_{max}}{i_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32WIGain		acc32_t	The angular speed coefficient gain is set up according to Equation 5 as: $\frac{\Delta L T_s}{L_D + T_s R_S} \cdot \omega_{max}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32EGain		acc32_t	The back-EMF coefficient gain is set up according to Equation 5 as: $\frac{T_s}{L_D + T_s R_S} \cdot \frac{e_{max}}{i_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
sUnityVctr		GMCLIB_2COOR_SINC OS_T_F16	The output - estimated angle as the sin/cos vector.

2.6.3 AMCLIB_BEMF_OBSRV_AB_T_FLT type description

Variable name	Data type	Description
sEObsrv	GMCLIB_2COOR_ALBE _T_FLT	The estimated back-EMF voltage structure.

Table continues on the next page...

Table continued from the previous page...

Variable name		Data type	Description
sIObsrv		GMCLIB_2COOR_ALBE_T_FLT	The estimated current structure.
sCtrl	fltAlpha_1	float_t	The state variable in the alpha part of the observer, integral part at step k-1. The variable is within the range <-1 ; 1).
	fltBeta_1	float_t	The state variable in the beta part of the observer, integral part at step k-1. The variable is within the range <-1 ; 1).
	fltPGain	float_t	The observer proportional gain is set up according to Equation 7 as: $2\xi\omega_0L_D - R_S$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
	fltIGain	float_t	The observer integral gain is set up according to Equation 7 as: $\omega_0^2L_D - R_S$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltIGain		float_t	The current coefficient gain is set up according to Equation 4 as: $\frac{L_D}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltUGain		float_t	The voltage coefficient gain is set up according to Equation 4 as: $\frac{T_s}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltWIGain		float_t	The angular speed coefficient gain is set up according to Equation 4 as: $\frac{\Delta L T_s}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltEGain		float_t	The back-EMF coefficient gain is set up according to Equation 4 as: $\frac{T_s}{L_D + T_s R_S}$

Table continues on the next page...

Table continued from the previous page...

Variable name	Data type	Description
		The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
sUnityVctr	GMCLIB_2COOR_SINC OS_T_FLT	The output - estimated angle as the sin/cos vector.

2.6.4 Declaration

The available AMCLIB_PMSMBemfObsrvABInit functions have the following declarations:

```
void AMCLIB_PMSMBemfObsrvABInit_F16(AMCLIB_BEMF_OBSRV_AB_T_A32 *psCtrl)
void AMCLIB_PMSMBemfObsrvABInit_FLT(AMCLIB_BEMF_OBSRV_AB_T_FLT *psCtrl)
```

The available AMCLIB_PMSMBemfObsrvAB functions have the following declarations:

```
void AMCLIB_PMSMBemfObsrvAB_F16(const GMCLIB_2COOR_ALBE_T_F16 *psIAIbe, const GMCLIB_2COOR_ALBE_T_F16
*psUAIbe, frac16_t f16Speed, AMCLIB_BEMF_OBSRV_AB_T_A32 *psCtrl)

void AMCLIB_PMSMBemfObsrvAB_FLT(const GMCLIB_2COOR_ALBE_T_FLT *psIAIbe, const GMCLIB_2COOR_ALBE_T_FLT
*psUAIbe, float_t fltSpeed, AMCLIB_BEMF_OBSRV_AB_T_FLT *psCtrl)
```

2.6.5 Function use

The use of the AMCLIB_PMSMBemfObsrvAB function is shown in the following examples:

Fixed-point version:

```
#include "amclib.h"

static GMCLIB_2COOR_ALBE_T_F16 sIAIbe, sUAIbe;
static AMCLIB_BEMF_OBSRV_AB_T_A32 sBemfObsrv;
static frac16_t f16Speed;

void Isr(void);

void main (void)
{
    sBemfObsrv.sCtrl.a32PGain= ACC32(1.697);
    sBemfObsrv.sCtrl.a32IGain= ACC32(0.134);
    sBemfObsrv.a32IGain = ACC32(0.986);
    sBemfObsrv.a32UGain = ACC32(0.170);
    sBemfObsrv.a32WIGain= ACC32(0.110);
    sBemfObsrv.a32EGain = ACC32(0.116);

    /* Initialization of the observer's structure */
    AMCLIB_PMSMBemfObsrvABInit_F16(&sBemfObsrv);

    sIAIbe.f16Alpha = FRAC16(0.05);
    sIAIbe.f16Beta  = FRAC16(0.1);
    sUAIbe.f16Alpha = FRAC16(0.2);
    sUAIbe.f16Beta  = FRAC16(-0.1);
```

```

}

/* Periodical function or interrupt */
void Isr(void)
{
    /* BEMF Observer calculation */
    AMCLIB_PMSMBemfObsrvAB_F16(&sIAlBe, &sUAlBe, f16Speed, &sBemfObsrv);
}

```

Floating-point version:

```

#include "amclib.h"

static GMCLIB_2COOR_ALBE_T_FLT sIAlBe, sUAlBe;
static AMCLIB_BEMF_OBSRV_AB_T_FLT sBemfObsrv;
static float_t fltSpeed;

void Isr(void);

void main (void)
{
    sBemfObsrv.sCtrl.fltIAlpha_1 = 0.0F;
    sBemfObsrv.sCtrl.fltIBeta_1 = 0.0F;
    sBemfObsrv.sCtrl.fltPGain = 1.697F;
    sBemfObsrv.sCtrl.fltIGain = 0.134F;
    sBemfObsrv.fltIGain = 0.986F;
    sBemfObsrv.fltUGain = 0.170F;
    sBemfObsrv.fltWIGain = 0.110F;
    sBemfObsrv.fltEGain = 0.116F;

    sIAlBe.fltAlpha = 0.05F;
    sIAlBe.fltBeta = 0.1F;
    sUAlBe.fltAlpha = 0.2F;
    sUAlBe.fltBeta = -0.1F;
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* BEMF Observer calculation */
    AMCLIB_PMSMBemfObsrvAB_FLT(&sIAlBe, &sUAlBe, fltSpeed, &sBemfObsrv);
}

```

2.7 AMCLIB_PMSMBemfObsrvDQ

The [AMCLIB_PMSMBemfObsrvDQ](#) function calculates the algorithm of back-electro-motive force observer in a rotating reference frame. The method for estimating the rotor position and angular speed is based on the mathematical model of an interior PMSM motor with an extended electro-motive force function, which is realized in an estimated quasi-synchronous reference frame $\gamma\delta$ as shown in [Figure 1](#).

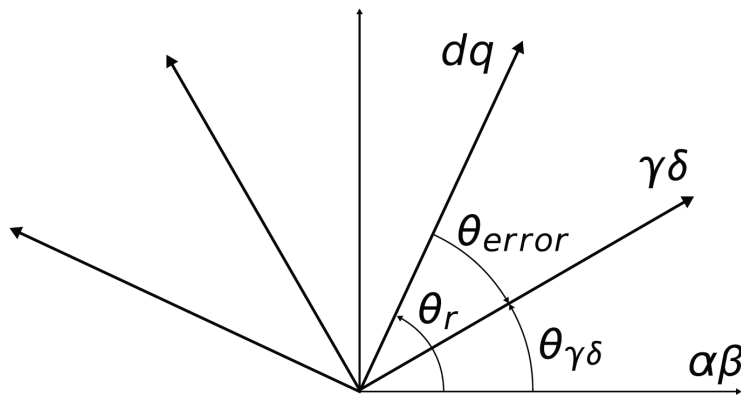


Figure 37. The estimated and real rotor dq synchronous reference frames

The back-EMF observer detects the generated motor voltages induced by the permanent magnets. A tracking observer uses the back-EMF signals to calculate the position and speed of the rotor. The transformed model is then derived as follows:

$$\begin{bmatrix} u_\gamma \\ u_\delta \end{bmatrix} = \begin{bmatrix} R_S + sL_D & -\omega_r L_Q \\ \omega_r L_Q & R_S + sL_D \end{bmatrix} \cdot \begin{bmatrix} i_\gamma \\ i_\delta \end{bmatrix} + (\Delta L \cdot (\omega_r i_D - s i_Q) + \Psi_m \omega_r) \cdot \begin{bmatrix} -\sin(\theta_{error}) \\ \cos(\theta_{error}) \end{bmatrix}$$

where:

- R_S is the stator resistance
- L_D and L_Q are the D-axis and Q-axis inductances
- Ψ_m is the back-EMF constant
- ω_r is the angular electrical rotor speed
- u_γ and u_δ are the estimated stator voltages
- i_γ and i_δ are the estimated stator currents
- θ_{error} is the error between the actual D-Q frame and the estimated frame position
- s is the operator of the derivative

The block diagram of the observer in the estimated reference frame is shown in [Figure 1](#). The observer compensator is substituted by a standard PI controller with following equation in the fractional arithmetic.

$$i_{sc}(k) \cdot i_{max} = K_P \cdot e_{sc}(k) \cdot e_{max} + T_s \cdot K_I \cdot e_{sc}(k) \cdot e_{max} + i_{sc}(k-1) \cdot i_{max}$$

where:

- K_P is the observer proportional gain [-]
- K_I is the observer integral gain [-]
- $i_{sc}(k) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the actual step
- $i_{sc}(k-1) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the previous step
- $e_{sc}(k) = [e_\gamma, e_\delta]$ is the scaled stator back-EMF voltage vector in the actual step
- i_{max} is the maximum current [A]
- e_{max} is the maximum back-EMF voltage [V]
- T_s is the sampling time [s]

As shown in [Figure 1](#), the observer model and hence also the PI controller gains in both axes are identical to each other.

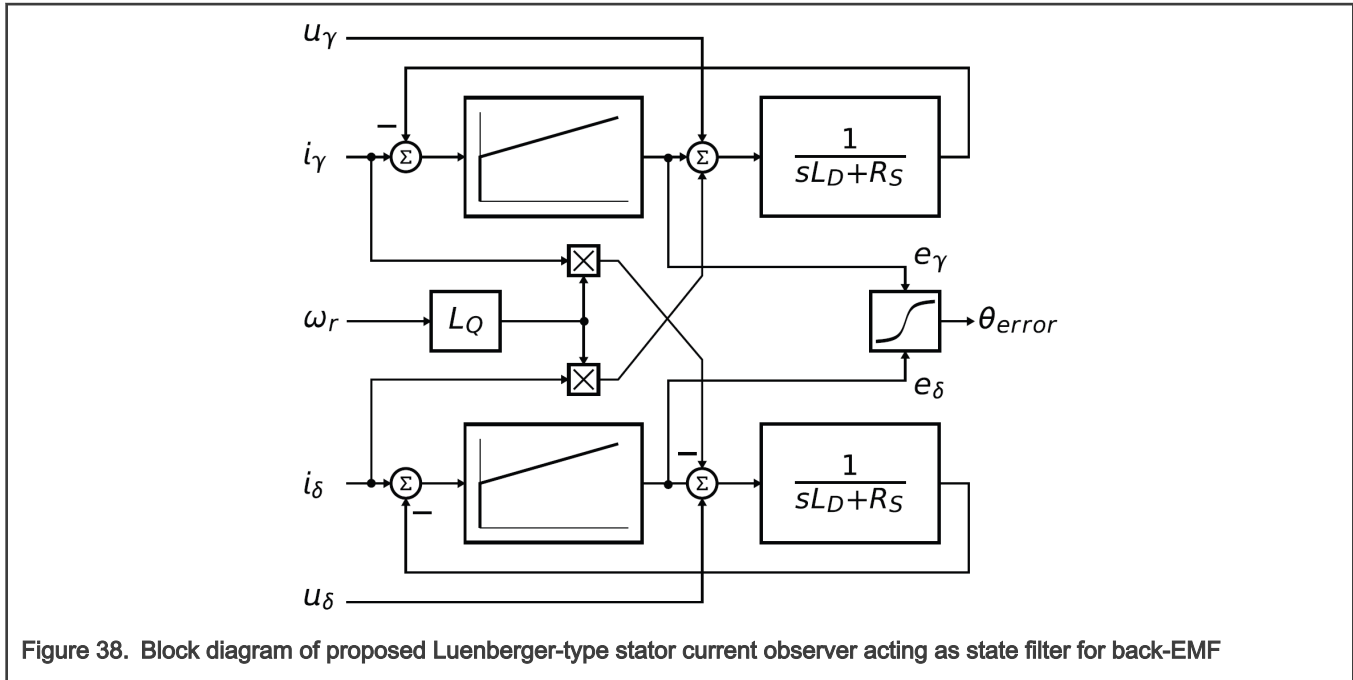


Figure 38. Block diagram of proposed Luenberger-type stator current observer acting as state filter for back-EMF

The position estimation can now be performed by extracting the θ_{error} term from the model, and adjusting the position of the estimated reference frame to achieve $\theta_{error} = 0$. Because the θ_{error} term is only included in the saliency-based EMF component of both u_γ and u_δ axis voltage equations, the Luenberger-based disturbance observer is designed to observe the u_γ and u_δ voltage components. The position displacement information θ_{error} is then obtained from the estimated back-EMFs as follows:

$$\theta_{error} = \text{atan}\left(\frac{-e_\gamma}{e_\delta}\right)$$

The estimated position

$$\hat{\theta}_e$$

can be obtained by driving the position of the estimated reference frame to achieve zero displacement $\theta_{error} = 0$. The phase-locked-loop mechanism can be adopted, where the loop compensator ensures correct tracking of the actual rotor flux position by keeping the error signal θ_{error} zeroed, $\theta_{error} = 0$.

A perfect match between the actual and estimated motor model parameters is assumed, and then the back-EMF transfer function can be simplified as follows:

$$\hat{E}_{\alpha\beta}(s) = -E_{\alpha\beta}(s) \cdot \frac{F_C(s)}{sL_D + R_S + F_C(s)}$$

The appropriate dynamic behavior of the back-EMF observer is achieved by the placement of the poles of the stator current observer characteristic polynomial. This general method is based on matching the coefficients of the characteristic polynomial with the coefficients of the general second-order system.

The back-EMF observer is a Luenberger-type observer with a motor model, which is implemented using the backward Euler transformation as follows:

$$i(k) = \frac{T_s}{L_D + T_s R_S} \cdot u(k) + \frac{T_s}{L_D + T_s R_S} \cdot e(k) + \frac{L_Q T_s}{L_D + T_s R_S} \cdot \omega_d(k) \cdot i'(k) + \frac{L_D}{L_D + T_s R_S} \cdot i(k-1)$$

where:

- $i(k) = [i_\gamma, i_\delta]$ is the stator current vector in the actual step
- $i(k-1) = [i_\gamma, i_\delta]$ is the stator current vector in the previous step

- $u(k) = [u_\gamma, u_\delta]$ is the stator voltage vector in the actual step
- $e(k) = [e_\gamma, e_\delta]$ is the stator back-EMF voltage vector in the actual step
- $i'(k) = [i_\gamma, -i_\delta]$ is the complementary stator current vector in the actual step
- $\omega_e(k)$ is the electrical angular speed in the actual step
- T_s is the sampling time [s]

This equation is transformed into the fractional arithmetic as follows:

$$i_{sc}(k) \cdot i_{max} = \frac{T_s}{L_D + T_s R_S} \cdot u_{sc}(k) \cdot u_{max} + \frac{T_s}{L_D + T_s R_S} \cdot e_{sc}(k) \cdot e_{max} + \frac{L_D T_s}{L_D + T_s R_S} \cdot \omega_{esc}(k) \cdot \omega_{max} \cdot i'_{sc}(k) \cdot i_{max} + \frac{L_D}{L_D + T_s R_S} \cdot i_{sc}(k-1) \cdot i_{max}$$

where:

- $i_{sc}(k) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the actual step
- $i_{sc}(k-1) = [i_\gamma, i_\delta]$ is the scaled stator current vector in the previous step
- $u_{sc}(k) = [u_\gamma, u_\delta]$ is the scaled stator voltage vector in the actual step
- $e_{sc}(k) = [e_\gamma, e_\delta]$ is the scaled stator back-EMF voltage vector in the actual step
- $i'_{sc}(k) = [i_\gamma, -i_\delta]$ is the scaled complementary stator current vector in the actual step
- $\omega_{esc}(k)$ is the scaled electrical angular speed in the actual step
- i_{max} is the maximum current [A]
- e_{max} is the maximum back-EMF voltage [V]
- u_{max} is the maximum stator voltage [V]
- ω_{max} is the maximum electrical angular speed in [rad / s]

If the Luenberger-type stator current observer is properly designed in the stationary reference frame, the back-EMF can be estimated as a disturbance produced by the observer controller. However, this is only valid when the back-EMF term is not included in the observer model. The observer is a closed-loop current observer, therefore it acts as a state filter for the back-EMF term.

The estimate of the extended EMF term can be derived from [AMCLIB_PMSMBemfObsrvDQ_Eq3](#) as follows:

$$-\frac{\hat{E}_{\gamma\delta}(s)}{E_{\gamma\delta}(s)} = \frac{sK_P + K_I}{s^2 L_D + sR_S + sK_P + K_I}$$

The observer controller can be designed by comparing the closed-loop characteristic polynomial with that of a standard second-order system as follows:

$$s^2 + \frac{K_P + R_S}{L_D} \cdot s + \frac{K_I}{L_D} = s^2 + 2\xi\omega_0 s + \omega_0^2$$

where:

- ω_0 is the natural frequency of the closed-loop system (loop bandwidth)
- ξ is the loop attenuation
- K_P is the proportional gain
- K_I is the integral gain

2.7.1 Available versions

This function is available in the following versions:

- Fractional output - the output is the fractional portion of the result; the result is within the range $[-1; 1)$. The parameters use the accumulator types.
- Accumulator output with floating-point inputs - the output is the accumulator result; the result is within the range $[-1; 1)$. The inputs are 32-bit single precision floating-point values.

The available versions of the [AMCLIB_PMSMBemfObsrvDQ](#) function are shown in the following table:

Table 14. Init versions

Function name	Parameters	Result type
AMCLIB_PMSMBemfObsrvDQInit_F16	AMCLIB_BEMF_OBSRV_DQ_T_A32 *	void
	Initialization does not have any input.	
AMCLIB_PMSMBemfObsrvDQInit_A32fff	AMCLIB_BEMF_OBSRV_DQ_T_FLT *	void
	Initialization does not have any input.	

Table 15. Function versions

Function name	Input/output type		Result type
AMCLIB_PMSMBemfObsrvDQ_F16	Input	GMCLIB_2COOR_DQ_T_F16 *	frac16_t
		GMCLIB_2COOR_DQ_T_F16 *	
		frac16_t	
	Parameters	AMCLIB_BEMF_OBSRV_DQ_T_A32 *	
	Back-EMF observer with a 16-bit fractional input D-Q current and voltage, and a 16-bit electrical speed. All are within the range <-1 ; 1).		
AMCLIB_PMSMBemfObsrvDQ_A32fff	Input	GMCLIB_2COOR_DQ_T_FLT *	acc32_t
		GMCLIB_2COOR_DQ_T_FLT *	
		float_t	
	Parameters	AMCLIB_BEMF_OBSRV_DQ_T_FLT *	
	Back-EMF observer with a 32-bit single precision floating-point input D-Q current and voltage, and a 32-bit single precision floating-point electrical speed. All are within the full range. The output is a 32-bit accumulator angle error normalized to the range <-1 ; 1) that represents an angle (in radians) within the range <-π ; π).		

2.7.2 AMCLIB_BEMF_OBSRV_DQ_T_A32 type description

Variable name		Data type	Description
sEObsrv		GMCLIB_2COOR_DQ_T_F32	Estimated back-EMF voltage structure.
sIObsrv		GMCLIB_2COOR_DQ_T_F32	Estimated current structure.
sCtrl	f32ID_1	frac32_t	State variable in the alpha part of the observer, integral part at step k - 1. The variable is within the range $[-1; 1)$.

Table continues on the next page...

Table continued from the previous page...

Variable name		Data type	Description
	f32IQ_1	frac32_t	State variable in the beta part of the observer, integral part at step k - 1. The variable is within the range <-1 ; 1).
	a32PGain	acc32_t	The observer proportional gain is set up according to Equation 7 as: $(2\xi\omega_0L_D - R_S) \frac{i_{max}}{e_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
	a32IGain	acc32_t	The observer integral gain is set up according to Equation 7 as: $\omega_0^2 L_D T_s \frac{i_{max}}{e_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32IGain		acc32_t	The current coefficient gain is set up according to Equation 5 as: $\frac{L_D}{L_D + T_s R_S}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32UGain		acc32_t	The voltage coefficient gain is set up according to Equation 5 as: $\frac{T_s}{L_D + T_s R_S} \cdot \frac{u_{max}}{i_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32WIGain		acc32_t	The angular speed coefficient gain is set up according to Equation 5 as: $\frac{L_Q T_s}{L_D + T_s R_S} \cdot \omega_{max}$ The parameter is within the range <0 ; 65536.0). Set by the user.
a32EGain		acc32_t	The back-EMF coefficient gain is set up according to Equation 5 as: $\frac{T_s}{L_D + T_s R_S} \cdot \frac{e_{max}}{i_{max}}$ The parameter is within the range <0 ; 65536.0). Set by the user.
f16Error		frac16_t	Output - estimated phase error between a real D / Q frame system and an estimated D / Q reference system. The error is within the range <-1 ; 1).

2.7.3 AMCLIB_BEMF_OBSRV_DQ_T_FLT type description

Variable name		Data type	Description
sEObsrv		GMCLIB_2COOR_DQ_T_FLT	Estimated back-EMF voltage structure.
sIObsrv		GMCLIB_2COOR_DQ_T_FLT	Estimated current structure.
sCtrl	fltID_1	float_t	State variable in the alpha part of the observer; integral part at step k - 1. The variable is within the range <-1 ; 1).
	fltIQ_1	float_t	State variable in the beta part of the observer; integral part at step k - 1. The variable is within the range <-1 ; 1).
	fltPGain	float_t	Observer proportional gain is set up according to Equation 7 as: $2\xi\omega_0 L_D - R_S$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
	fltIGain	float_t	The observer integral gain is set up according to Equation 7 as: $\omega_0^2 L_D - R_S$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltIGain		float_t	The current coefficient gain is set up according to Equation 4 as: $\frac{L_D}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltUGain		float_t	The voltage coefficient gain is set up according to Equation 4 as: $\frac{T_s}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltWIGain		float_t	The angular speed coefficient gain is set up according to Equation 4 as: $\frac{L_Q T_s}{L_D + T_s R_S}$ The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
fltEGain		float_t	The back-EMF coefficient gain is set up according to Equation 4 as: $\frac{T_s}{L_D + T_s R_S}$

Table continues on the next page...

Table continued from the previous page...

Variable name	Data type	Description
		The parameter is a 32-bit single precision floating-point type non-negative value. Set by the user.
a32Error	acc32_t	Output - estimated phase error between a real D / Q frame system and an estimated D / Q reference system. The error is within the range $-1 ; 1$).

2.7.4 Declaration

The available [AMCLIB_PMSMBemfObsrvDQInit](#) functions have the following declarations:

```
void AMCLIB_PMSMBemfObsrvDQInit_F16(AMCLIB_BEMF_OBSRV_DQ_T_A32 *psCtrl)
void AMCLIB_PMSMBemfObsrvDQInit_A32fff(AMCLIB_BEMF_OBSRV_DQ_T_FLT *psCtrl)
```

The available [AMCLIB_PMSMBemfObsrvDQ](#) functions have the following declarations:

```
frac16_t AMCLIB_PMSMBemfObsrvDQ_F16(const GMCLIB_2COOR_DQ_T_F16 *psIDQ, const GMCLIB_2COOR_DQ_T_F16
*psUDQ, frac16_t f16Speed, AMCLIB_BEMF_OBSRV_DQ_T_A32 *psCtrl)

acc32_t AMCLIB_PMSMBemfObsrvDQ_A32fff(const GMCLIB_2COOR_DQ_T_FLT *psIDQ, const GMCLIB_2COOR_DQ_T_FLT
*psUDQ, float_t fltSpeed, AMCLIB_BEMF_OBSRV_DQ_T_FLT *psCtrl)
```

2.7.5 Function use

The use of the [AMCLIB_PMSMBemfObsrvDQ](#) function is shown in the following example:

```
#include "amclib.h"

static GMCLIB_2COOR_DQ_T_F16      sIdq, sUdq;
static AMCLIB_BEMF_OBSRV_DQ_T_A32 sBemfObsrv;
static frac16_t f16Speed, f16Error;

void Isr(void);

void main (void)
{
    sBemfObsrv.sCtrl.a32PGain= ACC32(1.697);
    sBemfObsrv.sCtrl.a32IGain= ACC32(0.134);
    sBemfObsrv.a32IGain = ACC32(0.986);
    sBemfObsrv.a32UGain = ACC32(0.170);
    sBemfObsrv.a32WIGain= ACC32(0.110);
    sBemfObsrv.a32EGain = ACC32(0.116);

    /* Initialization of the observer's structure */
    AMCLIB_PMSMBemfObsrvDQInit_F16(&sBemfObsrv);

    sIdq.f16D = FRAC16(0.05);
    sIdq.f16Q = FRAC16(0.1);
    sUdq.f16D = FRAC16(0.2);
    sUdq.f16Q = FRAC16(-0.1);
}
```

```

/* Periodical function or interrupt */
void Isr(void)
{
    /* BEMF Observer calculation */
    fl6Error = AMCLIB_PMSMBemfObsrvDQ_Fl6(&sIdq, &sUdq, fl6Speed, &sBemfObsrv);
}

```

2.8 AMCLIB_TrackObsrv

The [AMCLIB_TrackObsrv](#) function calculates a tracking observer for the determination of angular speed and position of the input error functional signal. The tracking-observer algorithm uses the phase-locked-loop mechanism. It is recommended to call this function at every sampling period. It requires a single input argument as a phase error. A phase-tracking observer with a standard PI controller used as the loop compensator is shown in [Figure 1](#).

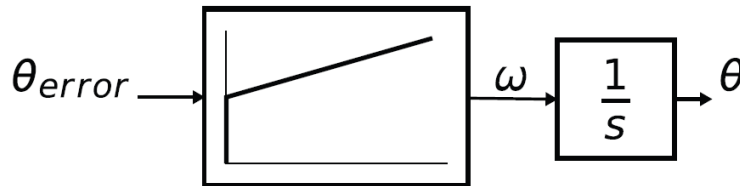


Figure 39. Block diagram of proposed PLL scheme for position estimation

The depicted tracking observer structure has the following transfer function:

$$\frac{\hat{\theta}(s)}{\theta(s)} = \frac{sK_P + K_I}{s^2 + sK_P + K_I}$$

The controller gains K_P and K_I are calculated by comparing the characteristic polynomial of the resulting transfer function to a standard second-order system polynomial.

The essential equations for implementation of the tracking observer according to the block scheme in [Figure 1](#) are as follows:

$$\omega(k) = K_P \cdot e(k) + T_s \cdot K_I \cdot e(k) + \omega(k-1)$$

$$\theta(k) = T_s \cdot \omega(k) + \theta(k-1)$$

where:

- K_P is the proportional gain
- K_I is the integral gain
- T_s is the sampling period [s]
- $e(k)$ is the position error in step k
- $\omega(k)$ is the rotor speed [rad / s] in step k
- $\omega(k-1)$ is the rotor speed [rad / s] in step k - 1
- $\theta(k)$ is the rotor angle [rad] in step k
- $\theta(k-1)$ is the rotor angle [rad] in step k - 1

In the fractional arithmetic, [AMCLIB_TrackObsrv_Eq1](#) and [AMCLIB_TrackObsrv_Eq2](#) are as follows:

$$\omega_{sc}(k) \cdot \omega_{max} = K_P \cdot e_{sc}(k) + T_s \cdot K_I \cdot e_{sc}(k) + \omega_{sc}(k-1) \cdot \omega_{max}$$

$$\theta_{sc}(k) \cdot \theta_{max} = T_s \cdot \omega_{sc}(k) \cdot \omega_{max} + \theta_{sc}(k-1) \cdot \theta_{max}$$

where:

- $e_{sc}(k)$ is the scaled position error in step k
- $\omega_{sc}(k)$ is the scaled rotor speed [rad / s] in step k
- $\omega_{sc}(k-1)$ is the scaled rotor speed [rad / s] in step k - 1
- $\theta_{sc}(k)$ is the scaled rotor angle [rad] in step k
- $\theta_{sc}(k-1)$ is the scaled rotor angle [rad] in step k - 1
- ω_{max} is the maximum speed
- θ_{max} is the maximum rotor angle (typically)

2.8.1 Available versions

The function is available in the following versions:

- Fractional output - the output is the fractional portion of the result; the result is within the range $<-1 ; 1$).
- Accumulator output with floating point structure - the output is the accumulator result; the result is within the range $<-1 ; 1$). The structure of the parameters contains the 32-bit single precision floating-point values.

The available versions of the [AMCLIB_TrackObsrv](#) function are shown in the following table:

Table 16. Init versions

Function name	Init angle	Parameters	Result type
AMCLIB_TrackObsrvInit_F16	frac16_t	AMCLIB_TRACK_OBSRV_T_F32 *	void
	The input is a 16-bit fractional value of the angle normalized to the range $<-1 ; 1$) that represents an angle (in radians) within the range $<-\pi ; \pi$).		
AMCLIB_TrackObsrvInit_A32af	acc32_t	AMCLIB_TRACK_OBSRV_T_FLT *	void
	Input is the 32-bit accumulator value of the angle normalized to the range $<-1 ; 1$) that represents an angle in radians within the range $<-\pi ; \pi$). The parameters are 32-bit single precision values.		

Table 17. Function versions

Function name	Input type	Parameters	Result type
AMCLIB_TrackObsrv_F16	frac16_t	AMCLIB_TRACK_OBSRV_T_F32 *	frac16_t
	Tracking observer with a 16-bit fractional position error input divided by π . The output from the observer is a 16-bit fractional position normalized to the range $<-1 ; 1$) that represents an angle (in radians) within the range $<-\pi ; \pi$).		
AMCLIB_TrackObsrv_A32af	acc32_t	AMCLIB_TRACK_OBSRV_T_FLT *	acc32_t
	Tracking observer with a 32-bit accumulator position divided by π . The output from the observer is a 32-bit accumulator position normalized to the range $<-1 ; 1$) that represents an angle (in radians) within the range $<-\pi ; \pi$). The parameters are 32-bit single precision values.		

2.8.2 AMCLIB_TRACK_OBSRV_T_F32

Variable name	Input type	Description
f32Theta	frac32_t	Estimated position as the output of the second numerical integrator. The parameter is within the range <-1 ; 1). Controlled by the algorithm.
f32Speed	frac32_t	Estimated speed as the output of the first numerical integrator. The parameter is within the range <-1 ; 1). Controlled by the algorithm.
f32I_1	frac32_t	State variable in the controller part of the observer; integral part at step k - 1. The parameter is within the range <-1 ; 1). Controlled by the algorithm.
f16IGain	frac16_t	The observer integral gain is set up according to Equation 4 as: $T_s \cdot K_I \cdot \frac{1}{\omega_{max}} \cdot 2^{-Ish}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16IGainSh	int16_t	The observer integral gain shift takes care of keeping the f16IGain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(T_s \cdot K_I \cdot \frac{1}{\omega_{max}}) - \log_2 1 < Ish \leq \log_2(T_s \cdot K_I \cdot \frac{1}{\omega_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.
f16PGain	frac16_t	The observer proportional gain is set up according to Equation 4 as: $K_P \cdot \frac{1}{\omega_{max}} \cdot 2^{-Psh}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16PGainSh	int16_t	The observer proportional gain shift takes care of keeping the f16PGain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(K_P \cdot \frac{1}{\omega_{max}}) - \log_2 1 < Psh \leq \log_2(K_P \cdot \frac{1}{\omega_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.
f16ThGain	frac16_t	The observer gain for the output position integrator is set up according to Equation 5 as: $T_s \cdot \frac{\omega_{max}}{\theta_{max}} \cdot 2^{-Thsh}$ The parameter is a 16-bit fractional type within the range <0 ; 1). Set by the user.
i16ThGainSh	int16_t	The observer gain shift for the position integrator takes care of keeping the f16ThGain variable within the fractional range <-1 ; 1). The shift is determined as: $\log_2(T_s \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 1 < THsh \leq \log_2(T_s \cdot \frac{\omega_{max}}{\theta_{max}}) - \log_2 0.5$ The parameter is a 16-bit integer type within the range <-15 ; 15>. Set by the user.

2.8.3 AMCLIB_TRACK_OBSRV_T_FLT

Variable name	Input type	Description
f32Theta	frac32_t	Estimated position as the output of the second numerical integrator. The parameter is within the range $[-1; 1]$. Controlled by the algorithm.
fltSpeed	float_t	Estimated speed as the output of the first numerical integrator. The parameter is within the full range. Controlled by the algorithm.
fltI_1	float_t	State variable in the controller part of the observer; integral part at the step $k - 1$. The parameter is within the full range. Controlled by the algorithm.
fltIGain	float_t	The observer integral gain is set up according to Equation 2 as: $K_I T_s$ The parameter is a 32-bit single precision floating-point value in range (0; 16383.99999). Set by the user.
fltPGain	float_t	The observer proportional gain is set up according to Equation 2 as: K_P The parameter is a 32-bit single precision floating-point value in range (0; 32767.99998). Set by the user.
fltThGain	float_t	The observer gain for the output position integrator is set up according to Equation 3 as: T_s The parameter is a 32-bit single precision floating-point value in range (0; 1). Set by the user.

2.8.4 Declaration

The available AMCLIB_TrackObsrvInit functions have the following declarations:

```
void AMCLIB_TrackObsrvInit_F16(frac16\_t f16ThetaInit, AMCLIB\_TRACK\_OBSRV\_T\_F32 *psCtrl)
void AMCLIB_TrackObsrvInit_A32af(acc32\_t a32ThetaInit, AMCLIB\_TRACK\_OBSRV\_T\_FLT *psCtrl)
```

The available [AMCLIB_TrackObsrv](#) functions have the following declarations:

```
frac16\_t AMCLIB_TrackObsrv_F16(frac16\_t f16Error, AMCLIB\_TRACK\_OBSRV\_T\_F32 *psCtrl)
acc32\_t AMCLIB_TrackObsrv_A32af(acc32\_t a32Error, AMCLIB\_TRACK\_OBSRV\_T\_FLT *psCtrl)
```

2.8.5 Function use

The use of the [AMCLIB_TrackObsrv](#) function is shown in the following example:

```
#include "amclib.h"

static AMCLIB\_TRACK\_OBSRV\_T\_F32 sTo;
static frac16\_t f16ThetaError;
static frac16\_t f16PositionEstim;

void Isr(void);

void main(void)
{
    sTo.f16IGain = FRAC16(0.6434);
    sTo.i16IGainSh = -9;
    sTo.f16PGain = FRAC16(0.6801);
    sTo.i16PGainSh = -2;
```

```
sTo.fl6ThGain    = FRAC16(0.6400);
sTo.il6ThGainSh  = -4;

AMCLIB_TrackObsrvInit_Fl6(FRAC16(0.0), &sTo);

fl6ThetaError    = FRAC16(0.5);
}

/* Periodical function or interrupt */
void Isr(void)
{
    /* Tracking observer calculation */
    fl6PositionEstim = AMCLIB_TrackObsrv_Fl6(fl6ThetaError, &sTo);
}
```

Appendix A

Library types

A.1 bool_t

The `bool_t` type is a logical 16-bit type. It is able to store the boolean variables with two states: TRUE (1) or FALSE (0). Its definition is as follows:

```
typedef unsigned short bool_t;
```

The following figure shows the way in which the data is stored by this type:

Table 18. Data storage

	15	14	13	12	11	10	9	8	7	6	5	4	3	2	1	0
Value	Unused															Logical
TRUE	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1
FALSE	0				0				0				1			
	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
	0				0				0				0			

To store a logical value as `bool_t`, use the `FALSE` or `TRUE` macros.

A.2 uint8_t

The `uint8_t` type is an unsigned 8-bit integer type. It is able to store the variables within the range <0 ; 255>. Its definition is as follows:

```
typedef unsigned char uint8_t;
```

The following figure shows the way in which the data is stored by this type:

Table 19. Data storage

	7	6	5	4	3	2	1	0
Value	Integer							
255	1	1	1	1	1	1	1	1
	F				F			

Table continues on the next page...

Table 19. Data storage (continued)

11	0	0	0	0	1	0	1	1
	0				B			
124	0	1	1	1	1	1	0	0
	7				C			
159	1	0	0	1	1	1	1	1
	9				F			

A.3 uint16_t

The `uint16_t` type is an unsigned 16-bit integer type. It is able to store the variables within the range $<0 ; 65535>$. Its definition is as follows:

```
typedef unsigned short uint16_t;
```

The following figure shows the way in which the data is stored by this type:

Table 20. Data storage

	15	14	13	12	11	10	9	8	7	6	5	4	3	2	1	0
Value	Integer															
65535	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
	F				F				F				F			
5	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	1
	0				0				0				5			
15518	0	0	1	1	1	1	0	0	1	0	0	1	1	1	1	0
	3				C				9				E			
40768	1	0	0	1	1	1	1	1	0	1	0	0	0	0	0	0
	9				F				4				0			

A.4 uint32_t

The `uint32_t` type is an unsigned 32-bit integer type. It is able to store the variables within the range $\langle 0 ; 4294967295 \rangle$. Its definition is as follows:

```
typedef unsigned long uint32_t;
```

The following figure shows the way in which the data is stored by this type:

Table 21. Data storage

	31	24	23	16	15	8	7	0
Value	Integer							
4294967295	F	F	F	F	F	F	F	F
2147483648	8	0	0	0	0	0	0	0
55977296	0	3	5	6	2	5	5	0
3451051828	C	D	B	2	D	F	3	4

A.5 int8_t

The `int8_t` type is a signed 8-bit integer type. It is able to store the variables within the range $\langle -128 ; 127 \rangle$. Its definition is as follows:

```
typedef char int8_t;
```

The following figure shows the way in which the data is stored by this type:

Table 22. Data storage

	7	6	5	4	3	2	1	0
Value	Sign	Integer						
127	0	1	1	1	1	1	1	1
	7				F			
-128	1	0	0	0	0	0	0	0
	8				0			
60	0	0	1	1	1	1	0	0
	3				C			

Table continues on the next page...

Table 22. Data storage (continued)

-97	1	0	0	1	1	1	1	1	
	9				F				

A.6 int16_t

The `int16_t` type is a signed 16-bit integer type. It is able to store the variables within the range $<-32768 ; 32767>$. Its definition is as follows:

```
typedef short int16_t;
```

The following figure shows the way in which the data is stored by this type:

Table 23. Data storage

	15	14	13	12	11	10	9	8	7	6	5	4	3	2	1	0
Value	Sign	Integer														
32767	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
	7				F				F				F			
-32768	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
	8				0				0				0			
15518	0	0	1	1	1	1	0	0	1	0	0	1	1	1	1	0
	3				C				9				E			
-24768	1	0	0	1	1	1	1	1	0	1	0	0	0	0	0	0
	9				F				4				0			

A.7 int32_t

The `int32_t` type is a signed 32-bit integer type. It is able to store the variables within the range $<-2147483648 ; 2147483647>$. Its definition is as follows:

```
typedef long int32_t;
```

The following figure shows the way in which the data is stored by this type:

Table 24. Data storage

--

Table continues on the next page...

Table 24. Data storage (continued)

	31	24	23	16	15	8	7	0
Value	S	Integer						
2147483647	7	F	F	F	F	F	F	F
-2147483648	8	0	0	0	0	0	0	0
55977296	0	3	5	6	2	5	5	0
-843915468	C	D	B	2	D	F	3	4

A.8 frac8_t

The `frac8_t` type is a signed 8-bit fractional type. It is able to store the variables within the range $<-1 ; 1$). Its definition is as follows:

```
typedef char frac8_t;
```

The following figure shows the way in which the data is stored by this type:

Table 25. Data storage

	7	6	5	4	3	2	1	0
Value	Sign	Fractional						
0.99219	0	1	1	1	1	1	1	1
	7				F			
-1.0	1	0	0	0	0	0	0	0
	8				0			
0.46875	0	0	1	1	1	1	0	0
	3				C			
-0.75781	1	0	0	1	1	1	1	1
	9				F			

To store a real number as `frac8_t`, use the `FRAC8` macro.

A.9 frac16_t

The `frac16_t` type is a signed 16-bit fractional type. It is able to store the variables within the range $<-1 ; 1$). Its definition is as follows:

```
typedef short frac16_t;
```

The following figure shows the way in which the data is stored by this type:

Table 26. Data storage

	15	14	13	12	11	10	9	8	7	6	5	4	3	2	1	0
Value	Sign	Fractional														
0.99997	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
-1.0	7				F				F				F			
	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
0.47357	8				0				0				0			
	0	0	1	1	1	1	0	0	1	0	0	1	1	1	1	0
-0.75586	3				C				9				E			
	1	0	0	1	1	1	1	1	0	1	0	0	0	0	0	0
	9				F				4				0			

To store a real number as `frac16_t`, use the `FRAC16` macro.

A.10 frac32_t

The `frac32_t` type is a signed 32-bit fractional type. It is able to store the variables within the range $<-1 ; 1$). Its definition is as follows:

```
typedef long frac32_t;
```

The following figure shows the way in which the data is stored by this type:

Table 27. Data storage

Table continues on the next page...

Table 27. Data storage (continued)

-1.0	8	0	0	0	0	0	0
0.02606645970	0	3	5	6	2	5	5
-0.3929787632	C	D	B	2	D	F	3

To store a real number as `frac32_t`, use the `FRAC32` macro.

A.11 `acc16_t`

The `acc16_t` type is a signed 16-bit fractional type. It is able to store the variables within the range $<-256 ; 256$). Its definition is as follows:

```
typedef short acc16_t;
```

The following figure shows the way in which the data is stored by this type:

Table 28. Data storage

	15	14	13	12	11	10	9	8	7	6	5	4	3	2	1	0	
Value	Sign	Integer								Fractional							
255.9921875	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	
	7				F				F				F				
-256.0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	
	8				0				0				0				
1.0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	
	0				0				8				0				
-1.0	1	1	1	1	1	1	1	1	1	0	0	0	0	0	0	0	
	F				F				8				0				
13.7890625	0	0	0	0	0	1	1	0	1	1	1	0	0	1	0	1	
	0				6				E				5				
-89.71875	1	1	0	1	0	0	1	1	0	0	1	0	0	1	0	0	
	D				3				2				4				

To store a real number as `acc16_t`, use the `ACC16` macro.

A.12 `acc32_t`

The `acc32_t` type is a signed 32-bit accumulator type. It is able to store the variables within the range $<-65536 ; 65536$). Its definition is as follows:

```
typedef long acc32_t;
```

The following figure shows the way in which the data is stored by this type:

Table 29. Data storage

	3124231615870								
Value	S	Integer				Fractional			
65535.999969	7	F	F	F	F	F	F	F	
-65536.0	8	0	0	0	0	0	0	0	
1.0	0	0	0	0	8	0	0	0	
-1.0	F	F	F	F	8	0	0	0	
23.789734	0	0	0	B	E	5	1	6	
-1171.306793	F	D	B	6	5	8	B	C	

To store a real number as `acc32_t`, use the `ACC32` macro.

A.13 `float_t`

The `float_t` type is a signed 32-bit single precision floating-point type, defined by IEEE 754. It is able to store the full precision (normalized) finite variables within the range $<-3.40282 \cdot 10^{38} ; 3.40282 \cdot 10^{38}$) with the minimum resolution of 2^{-23} . The smallest normalized number is $\pm 1.17549 \cdot 10^{-38}$. Nevertheless, the denormalized numbers (with reduced precision) reach yet lower values, from $\pm 1.40130 \cdot 10^{-45}$ to $\pm 1.17549 \cdot 10^{-38}$. The standard also defines the additional values:

- Negative zero
- Infinity
- Negative infinity
- Not a number

The 32-bit type is composed of:

- Sign (bit 31)
- Exponent (bits 23 to 30)
- Mantissa (bits 0 to 22)

The conversion of the number is straightforward. The sign of the number is stored in bit 31. The binary exponent is decoded as an integer from bits 23 to 30 by subtracting 127. The mantissa (fraction) is stored in bits 0 to 22. An invisible leading bit (it is not

actually stored) with value 1.0 is placed in front; therefore, bit 23 has a value of 0.5, bit 22 has a value 0.25, and so on. As a result, the mantissa has a value between 1.0 and 2. If the exponent reaches -127 (binary 00000000), the leading 1.0 is no longer used to enable the gradual underflow.

The `float_t` type definition is as follows:

```
typedef float float_t;
```

The following figure shows the way in which the data is stored by this type:

Table 30. Data storage - normalized values

	31	24	23	16	15	8	7	0																							
Value	S	Exponent							Mantissa																						
$(2.0 - 2^{-23}) \cdot 2^{127}$	0	1	1	1	1	1	1	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1			
$\approx 3.40282 \cdot 10^{38}$	7			F			7			F			F			F			F			F									
$-(2.0 - 2^{-23}) \cdot 2^{127}$	1	1	1	1	1	1	1	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1			
$\approx -3.40282 \cdot 10^{38}$	F			F			7			F			F			F			F												
2^{-126}	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
$\approx 1.17549 \cdot 10^{-38}$	0			0			8			0			0			0			0			0									
-2^{-126}	1	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
$\approx -1.17549 \cdot 10^{-38}$	8			0			8			0			0			0			0			0									
1.0	0	0	1	1	1	1	1	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
	3			F			8			0			0			0			0			0									
-1.0	1	0	1	1	1	1	1	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
	B			F			8			0			0			0			0			0									
π	0	1	0	0	0	0	0	0	0	1	0	0	1	0	0	1	0	0	0	1	1	1	1	1	1	0	1	1			
≈ 3.1415927	4			0			4			9			0			F			D			B									
-20810.086	1	1	0	0	0	1	1	0	1	0	1	0	0	1	0	1	0	0	0	0	1	0	1	1	0	0	0	0			
	C			6			A			2			9			4			2			C									

Table continues on the next page...

Table 30. Data storage - normalized values (continued)

--

Table 31. Data storage - denormalized values

	31	24	23	16	15	8	7	0																					
Value	S	Exponent							Mantissa																				
0.0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
		0		0		0		0		0		0		0		0		0		0		0		0		0		0	
-0.0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
		8		0		0		0		0		0		0		0		0		0		0		0		0		0	
$(1.0 - 2^{-23}) \cdot 2^{-126}$	0	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
$\approx 1.17549 \cdot 10^{-38}$		0		0		7		F		F		F		F		F		F		F		F		F		F		F	
$-(1.0 - 2^{-23}) \cdot 2^{-126}$	1	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
$\approx -1.17549 \cdot 10^{-38}$		8		0		7		F		F		F		F		F		F		F		F		F		F		F	
$2^{-1} \cdot 2^{-126}$	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
$\approx 5.87747 \cdot 10^{-39}$		0		0		4		0		0		0		0		0		0		0		0		0		0		0	
$-2^{-1} \cdot 2^{-126}$	1	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
$\approx -5.87747 \cdot 10^{-39}$		8		0		4		0		0		0		0		0		0		0		0		0		0		0	
$2^{-23} \cdot 2^{-126}$	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1
$\approx 1.40130 \cdot 10^{-45}$		0		0		0		0		0		0		0		0		0		0		0		0		0		1	
$-2^{-23} \cdot 2^{-126}$	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1
$\approx -1.40130 \cdot 10^{-45}$		8		0		0		0		0		0		0		0		0		0		0		0		0		1	

Table 32. Data storage - special values

	31	24	23	16	15	8	7	0																				
Value	S	Exponent							Mantissa																			
∞	0	1	1	1	1	1	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
	7 F							8	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	
$-\infty$	1	1	1	1	1	1	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	
	F F							8	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0		
Not a number	*	1	1	1	1	1	1	1	non zero																			
	7/F F							800001 to FFFFFFFF																				

A.14 GMCLIB_3COOR_T_F16

The `GMCLIB_3COOR_T_F16` structure type corresponds to the three-phase stationary coordinate system, based on the A, B, and C components. Each member is of the `frac16_t` data type. The structure definition is as follows:

```
typedef struct
{
    frac16_t f16A;
    frac16_t f16B;
    frac16_t f16C;
} GMCLIB_3COOR_T_F16;
```

The structure description is as follows:

Table 33. GMCLIB_3COOR_T_F16 members description

Type	Name	Description
frac16_t	f16A	A component; 16-bit fractional type
frac16_t	f16B	B component; 16-bit fractional type
frac16_t	f16C	C component; 16-bit fractional type

A.15 GMCLIB_3COOR_T_FLT

The `GMCLIB_3COOR_T_FLT` structure type corresponds to the three-phase stationary coordinate system, based on the A, B, and C components. Each member is of the `float t` data type. The structure definition is as follows:

```
typedef struct
{
    float_t fltA;
    float_t fltB;
```

```
float_t fltC;
} GMCLIB_3COOR_T_FLT;
```

The structure description is as follows:

Table 34. GMCLIB_3COOR_T_FLT members description

Type	Name	Description
float_t	fltA	A component; 32-bit single precision floating-point type
float_t	fltB	B component; 32-bit single precision floating-point type
float_t	fltC	C component; 32-bit single precision floating-point type

A.16 GMCLIB_2COOR_ALBE_T_F16

The [GMCLIB_2COOR_ALBE_T_F16](#) structure type corresponds to the two-phase stationary coordinate system, based on the Alpha and Beta orthogonal components. Each member is of the [frac16_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    frac16_t f16Alpha;
    frac16_t f16Beta;
} GMCLIB_2COOR_ALBE_T_F16;
```

The structure description is as follows:

Table 35. GMCLIB_2COOR_ALBE_T_F16 members description

Type	Name	Description
frac16_t	f16Apha	α -component; 16-bit fractional type
frac16_t	f16Beta	β -component; 16-bit fractional type

A.17 GMCLIB_2COOR_ALBE_T_FLT

The [GMCLIB_2COOR_ALBE_T_FLT](#) structure type corresponds to the two-phase stationary coordinate system based on the Alpha and Beta orthogonal components. Each member is of the [float_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    float_t fltAlpha;
    float_t fltBeta;
} GMCLIB_2COOR_ALBE_T_FLT;
```

The structure description is as follows:

Table 36. GMCLIB_2COOR_ALBE_T_FLT members description

Type	Name	Description
float_t	fltApha	α -component; 32-bit single precision floating-point type
float_t	fltBeta	β -component; 32-bit single precision floating-point type

A.18 GMCLIB_2COOR_DQ_T_F16

The [GMCLIB_2COOR_DQ_T_F16](#) structure type corresponds to the two-phase rotating coordinate system, based on the D and Q orthogonal components. Each member is of the [frac16_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    frac16_t f16D;
    frac16_t f16Q;
} GMCLIB_2COOR_DQ_T_F16;
```

The structure description is as follows:

Table 37. GMCLIB_2COOR_DQ_T_F16 members description

Type	Name	Description
frac16_t	f16D	D-component; 16-bit fractional type
frac16_t	f16Q	Q-component; 16-bit fractional type

A.19 GMCLIB_2COOR_DQ_T_F32

The [GMCLIB_2COOR_DQ_T_F32](#) structure type corresponds to the two-phase rotating coordinate system, based on the D and Q orthogonal components. Each member is of the [frac32_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    frac32_t f32D;
    frac32_t f32Q;
} GMCLIB_2COOR_DQ_T_F32;
```

The structure description is as follows:

Table 38. GMCLIB_2COOR_DQ_T_F32 members description

Type	Name	Description
frac32_t	f32D	D-component; 32-bit fractional type
frac32_t	f32Q	Q-component; 32-bit fractional type

A.20 GMCLIB_2COOR_DQ_T_FLT

The [GMCLIB_2COOR_DQ_T_FLT](#) structure type corresponds to the two-phase rotating coordinate system, based on the D and Q orthogonal components. Each member is of the [float_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    float_t fltD;
    float_t fltQ;
} GMCLIB_2COOR_DQ_T_FLT;
```

The structure description is as follows:

Table 39. GMCLIB_2COOR_DQ_T_FLT members description

Type	Name	Description
float_t	fltD	D-component; 32-bit single precision floating-point type
float_t	fltQ	Q-component; 32-bit single precision floating-point type

A.21 GMCLIB_2COOR_SINCOS_T_F16

The [GMCLIB_2COOR_SINCOS_T_F16](#) structure type corresponds to the two-phase coordinate system, based on the Sin and Cos components of a certain angle. Each member is of the [frac16_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    frac16_t f16Sin;
    frac16_t f16Cos;
} GMCLIB_2COOR_SINCOS_T_F16;
```

The structure description is as follows:

Table 40. GMCLIB_2COOR_SINCOS_T_F16 members description

Type	Name	Description
frac16_t	f16Sin	Sin component; 16-bit fractional type
frac16_t	f16Cos	Cos component; 16-bit fractional type

A.22 GMCLIB_2COOR_SINCOS_T_FLT

The [GMCLIB_2COOR_SINCOS_T_FLT](#) structure type corresponds to the two-phase coordinate system, based on the Sin and Cos components of a certain angle. Each member is of the [float_t](#) data type. The structure definition is as follows:

```
typedef struct
{
    float_t fltSin;
    float_t fltCos;
} GMCLIB_2COOR_SINCOS_T_FLT;
```

The structure description is as follows:

Table 41. GMCLIB_2COOR_SINCOS_T_FLT members description

Type	Name	Description
float_t	fltSin	Sin component; 32-bit single precision floating-point type
float_t	fltCos	Cos component; 32-bit single precision floating-point type

A.23 FALSE

The **FALSE** macro serves to write a correct value standing for the logical FALSE value of the **bool_t** type. Its definition is as follows:

```
#define FALSE ((bool_t)0)
```

```
#include "mlib.h"

static bool_t bVal;

void main(void)
{
    bVal = FALSE;           /* bVal = FALSE */
}
```

A.24 TRUE

The **TRUE** macro serves to write a correct value standing for the logical TRUE value of the **bool_t** type. Its definition is as follows:

```
#define TRUE ((bool_t)1)
```

```
#include "mlib.h"

static bool_t bVal;

void main(void)
{
    bVal = TRUE;           /* bVal = TRUE */
}
```

A.25 FRAC8

The **FRAC8** macro serves to convert a real number to the **frac8_t** type. Its definition is as follows:

```
#define FRAC8(x) ((frac8_t)((x) < 0.9921875 ? ((x) >= -1 ? (x)*0x80 : 0x80) : 0x7F))
```

The input is multiplied by 128 ($=2^7$). The output is limited to the range <0x80 ; 0x7F>, which corresponds to <-1.0 ; $1.0 \cdot 2^{-7}$ >.

```
#include "mlib.h"

static frac8_t f8Val;

void main(void)
{
    f8Val = FRAC8(0.187);   /* f8Val = 0.187 */
}
```

A.26 FRAC16

The **FRAC16** macro serves to convert a real number to the **frac16_t** type. Its definition is as follows:

```
#define FRAC16(x) ((frac16_t)((x) < 0.999969482421875 ? ((x) >= -1 ? (x)*0x8000 : 0x8000) : 0x7FFF))
```

The input is multiplied by 32768 ($=2^{15}$). The output is limited to the range <0x8000 ; 0x7FFF>, which corresponds to <-1.0 ; $1.0 \cdot 2^{-15}$ >.

```
#include "mlib.h"

static frac16_t f16Val;

void main(void)
{
    f16Val = FRAC16(0.736);          /* f16Val = 0.736 */
}
```

A.27 FRAC32

The **FRAC32** macro serves to convert a real number to the **frac32_t** type. Its definition is as follows:

```
#define FRAC32(x) ((frac32_t)((x) < 1 ? ((x) >= -1 ? (x)*0x80000000 : 0x80000000) : 0x7FFFFFFF))
```

The input is multiplied by 2147483648 ($=2^{31}$). The output is limited to the range <0x80000000 ; 0x7FFFFFFF>, which corresponds to <-1.0 ; $1.0 \cdot 2^{-31}$ >.

```
#include "mlib.h"

static frac32_t f32Val;

void main(void)
{
    f32Val = FRAC32(-0.1735667);    /* f32Val = -0.1735667 */
}
```

A.28 ACC16

The **ACC16** macro serves to convert a real number to the **acc16_t** type. Its definition is as follows:

```
#define ACC16(x) ((acc16_t)((x) < 255.9921875 ? ((x) >= -256 ? (x)*0x80 : 0x8000) : 0x7FFF))
```

The input is multiplied by 128 ($=2^7$). The output is limited to the range <0x8000 ; 0x7FFF> that corresponds to <-256.0 ; 255.9921875>.

```
#include "mlib.h"

static acc16_t a16Val;

void main(void)
{
```

```
a16Val = ACC16(19.45627);          /* a16Val = 19.45627 */
}
```

A.29 ACC32

The **ACC32** macro serves to convert a real number to the **acc32_t** type. Its definition is as follows:

```
#define ACC32(x) ((acc32_t)((x) < 65535.999969482421875 ? ((x) >= -65536 ? (x)*0x8000 : 0x80000000) : 0x7FFFFFFF))
```

The input is multiplied by 32768 ($=2^{15}$). The output is limited to the range $<0x80000000 ; 0x7FFFFFFF>$, which corresponds to $<-65536.0 ; 65536.0 \cdot 2^{-15}>$.

```
#include "mlib.h"

static acc32_t a32Val;

void main(void)
{
    a32Val = ACC32(-13.654437);      /* a32Val = -13.654437 */
}
```

How To Reach Us

Home Page:

nxp.com

Web Support:

nxp.com/support

Limited warranty and liability — Information in this document is provided solely to enable system and software implementers to use NXP products. There are no express or implied copyright licenses granted hereunder to design or fabricate any integrated circuits based on the information in this document. NXP reserves the right to make changes without further notice to any products herein.

NXP makes no warranty, representation, or guarantee regarding the suitability of its products for any particular purpose, nor does NXP assume any liability arising out of the application or use of any product or circuit, and specifically disclaims any and all liability, including without limitation consequential or incidental damages. "Typical" parameters that may be provided in NXP data sheets and/or specifications can and do vary in different applications, and actual performance may vary over time. All operating parameters, including "typicals," must be validated for each customer application by customer's technical experts. NXP does not convey any license under its patent rights nor the rights of others. NXP sells products pursuant to standard terms and conditions of sale, which can be found at the following address: nxp.com/SalesTermsandConditions.

Right to make changes - NXP Semiconductors reserves the right to make changes to information published in this document, including without limitation specifications and product descriptions, at any time and without notice. This document supersedes and replaces all information supplied prior to the publication hereof.

Security — Customer understands that all NXP products may be subject to unidentified or documented vulnerabilities. Customer is responsible for the design and operation of its applications and products throughout their lifecycles to reduce the effect of these vulnerabilities on customer's applications and products. Customer's responsibility also extends to other open and/or proprietary technologies supported by NXP products for use in customer's applications. NXP accepts no liability for any vulnerability. Customer should regularly check security updates from NXP and follow up appropriately. Customer shall select products with security features that best meet rules, regulations, and standards of the intended application and make the ultimate design decisions regarding its products and is solely responsible for compliance with all legal, regulatory, and security related requirements concerning its products, regardless of any information or support that may be provided by NXP. NXP has a Product Security Incident Response Team (PSIRT) (reachable at PSIRT@nxp.com) that manages the investigation, reporting, and solution release to security vulnerabilities of NXP products.

NXP, the NXP logo, NXP SECURE CONNECTIONS FOR A SMARTER WORLD, COOLFLUX, EMBRACE, GREENCHIP, HITAG, ICODE, JCOP, LIFE, VIBES, MIFARE, MIFARE CLASSIC, MIFARE DESFire, MIFARE PLUS, MIFARE FLEX, MANTIS, MIFARE ULTRALIGHT, MIFARE4MOBILE, MIGLO, NTAG, ROADLINK, SMARTLX, SMARTMX, STARPLUG, TOPFET, TRENCHMOS, UCODE, Freescale, the Freescale logo, Altivec, CodeWarrior, ColdFire, ColdFire+, the Energy Efficient Solutions logo, Kinetis, Layerscape, MagniV, mobileGT, PEG, PowerQUICC, Processor Expert, QorIQ, QorIQ Qonverge, SafeAssure, the SafeAssure logo, StarCore, Symphony, VortiQa, Vybrid, Airfast, BeeKit, BeeStack, CoreNet, Flexis, MXC, Platform in a Package, QUICC Engine, Tower, TurboLink, EdgeScale, EdgeLock, eIQ, and Immersive3D are trademarks of NXP B.V. All other product or service names are the property of their respective owners. AMBA, Arm, Arm7, Arm7TDMI, Arm9, Arm11, Artisan, big.LITTLE, Cordio, CoreLink, CoreSight, Cortex, DesignStart, DynamIQ, Jazelle, Keil, Mali, Mbed, Mbed Enabled, NEON, POP, RealView, SecurCore, Socrates, Thumb, TrustZone, ULINK, ULINK2, ULINK-ME, ULINK-PLUS, ULINKpro, uVision, Versatile are trademarks or registered trademarks of Arm Limited (or its subsidiaries) in the US and/or elsewhere. The related technology may be protected by any or all of patents, copyrights, designs and trade secrets. All rights reserved. Oracle and Java are registered trademarks of Oracle and/or its affiliates. The Power Architecture and Power.org word marks and the Power and Power.org logos and related marks are trademarks and service marks licensed by Power.org. M, M Mobileye and other Mobileye trademarks or logos appearing herein are trademarks of Mobileye Vision Technologies Ltd. in the United States, the EU and/or other jurisdictions.

© NXP B.V. 2021.

All rights reserved.

For more information, please visit: <http://www.nxp.com>

For sales office addresses, please send an email to: salesaddresses@nxp.com

Date of release: 01 November 2021
Document identifier: CM4FAMCLIBUG