
xCORE-200 DSP Library

This API reference manual describes the XMOS fixed-point digital signal processing software library. The library implements a suite of common signal processing functions for use on XMOS xCORE-200 multi-core microcontrollers.

Required tools

- xTIMEcomposer Tools Version 14.2.1 or later

Required hardware

Only XMOS xCORE-200 based multicore microcontrollers are supported with this library. The previous generation XS1 based multicore microcontrollers are not supported.

The xCORE-200 has a single cycle 32x32->64 bit multiply/accumulate unit, single cycle double-word load and store, dual issue instruction execution, and other instruction set enhancements. These features make xCORE-200 an efficient platform for executing digital signal processing algorithms.

Prerequisites

This document assumes familiarity with the XMOS xCORE architecture, the XMOS tool chain, the 'C' programming language, and digital signal processing concepts.

Software version and dependencies

The CHANGELOG contains information about the current and previous versions. For a list of direct dependencies, look for `DEPENDENT_MODULES` in `lib_dsp/module_build_info`.

Related application notes

The following application notes use this library:

- AN00209 - xCORE-200 DSP Library

1 Overview

1.1 Introduction

This API reference manual describes the Xmos xCORE-200 fixed-point digital signal processing firmware library. The library implements a suite of common signal processing functions for use on Xmos xCORE-200 multicore microcontrollers.

1.2 Library Organization

The library is divided into function collections with each collection covering a specific digital signal processing algorithm category. The API and implementation for each category are provided by a single 'C' header file and implementation file.

Category	Source Files	Functions
Fixed point	dsp_qformat	Q8 through Q31 formats, fixed and floating point conversions
Filters	dsp_filters	FIR, biquad, cascaded biquad, and convolution
Adaptive	dsp_adaptive	LMS and NLMS Adaptive filters
Scalar math	dsp_math	Multiply, divide, square root, exponential, natural logarithm trigonometric, hyperbolic
Integer math	dsp_math	Square root
Complex math	dsp_complex	Complex math functions plus real, imaginary combine/split
Vector math	dsp_vector	Scalar/vector add/subtract/multiply, dot product
Matrix math	dsp_matrix	Scalar/matrix add/subtract/multiply, inverse and transpose
Statistics	dsp_statistics	Vector mean, sum-of-squares, root-mean-square, variance
Design	dsp_design	Biquad coefficient generation for various filter types
FFT	dsp_fft	Forward and inverse Fast Fourier Transforms.
DCT	dsp_dct	Forward and inverse Discrete Cosine Transforms.

Table 1: DSP library organization

2 Fixed-Point Format

2.1 Q Format Introduction

The library functions support 32 bit input and output data, with internal 64 bit accumulator. The output data can be scaled to any of the supported Q Formats (Q8 through Q31), for all functions except for sample rate conversion, which uses a fixed Q31 format.

Further details about Q Format numbers is available [here](#)¹.

2.2 The 'q_format' Parameter

All XMOS DSP library functions that incorporate a multiply operation accept a parameter called q_format. This parameter can naively be used to specify the fixed point format for all operands and results (if applicable) where the formats are the same for all parameters. For example:

```
result_q28 = dsp_math_multiply( input1_q28, input2_q28, 28 );
```

The 'q_format' parameter, being used after one or more sequences of multiply and/or multiply-accumulate, is used to right-shift the 64-bit accumulator before truncating the value back to a 32-bit integer (i.e. the 32-bit fixed-point result). Therefore the 'q_format' parameter can be used to perform the proper fixed-point adjustment for any combination of input operand fixed-point format and desired fixed-point result format.

The output fixed-point fraction bit count is equal to the sum of the two input fraction bit counts minus the desired result fraction bit count:

```
q_format = input1 fraction bit count + input2 fraction bit count - result fraction bit count
```

For example:

```
// q_format_parameter = 31 = 30 + 29 - 28
result_q28 = dsp_math_multiply( input1_q30, input2_q29, 31 );

// q_format_parameter = 27 = 28 + 29 - 30
result_q30 = dsp_math_multiply( input1_q28, input2_q29, 27 );
```

¹[https://en.wikipedia.org/wiki/Q_\(number_format\)](https://en.wikipedia.org/wiki/Q_(number_format))

3 Filter Functions: Finite Impulse Response (FIR) Filter

Function	dsp_filters_fir
Description	This function implements a Finite Impulse Response (FIR) filter. The function operates on a single sample of input and output data (i.e. each call to the function processes one sample). The FIR filter algorithm is based upon a sequence of multiply-accumulate (MAC) operations. Each filter coefficient $b[i]$ is multiplied by a state variable which equals a previous input sample, or $y[n] = x[n]*b_0 + x[n-1]*b_1 + x[n-2]*b_2 + \dots + x[n-N+1]*b_{N-1}$. The parameter <code>filter_coeffs</code> points to a coefficient array of size $N = \text{num_taps}$. The filter coefficients are stored in forward order (e.g. $b_0, b_1, \dots, b_{N-1}$). The following example shows a five-tap (4th order) FIR filter with samples and coefficients represented in Q28 fixed-point format. <pre>int32_t filter_coeff[5] = { Q28(0.5), Q28(-0.5), Q28(0.0), Q28(-0.5), Q28(0.5) ↪ }; int32_t filter_state[4] = { 0, 0, 0, 0 }; int32_t result = dsp_filters_fir(sample, filter_coeff, filter_state, 5, ↪ 28);</pre> The FIR algorithm involves multiplication between 32-bit filter coefficients and 32-bit state data producing a 64-bit result for each coefficient and state data pair. Multiplication results are accumulated in a 64-bit accumulator. If overflow occurs in the final 64-bit result, it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. The saturation is only done after the last multiplication. To avoid 64-bit overflow in the intermediate results, the fixed point format must be chosen according to <code>num_taps</code>.
Type	<pre>int32_t dsp_filters_fir(int32_t input_sample, const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_taps, const int32_t q_format)</pre>
Parameters	<code>input_sample</code> The new sample to be processed. <code>filter_coeffs</code> Pointer to FIR coefficients array arranged as $[b_0, b_1, b_2, \dots, b_{N-1}]$. <code>state_data</code> Pointer to filter state data array of length $N-1$. Must be initialized at startup to all zeros. <code>num_taps</code> Number of filter taps ($N = \text{num_taps} = \text{filter order} + 1$). <code>q_format</code> Fixed point format (i.e. number of fractional bits).
Returns	The resulting filter output sample.

4 Filter Functions: Finite Impulse Response (FIR) Filter Add Sample

Function	dsp_filters_fir_add_sample
Description	This function pushes samples into an Finite Impulse Response (FIR) filter state array, without processing the filter. The function operates on a single sample of input and output data (i.e. each call to the function processes one sample). The FIR filter state array is updated with the new data. The parameter <code>filter_coeffs</code> points to a coefficient array of size $N = \text{num_taps}$. The filter coefficients are stored in forward order (e.g. $b_0, b_1, \dots, b_{N-1}$). The following example shows a five-tap (4th order) FIR filter. <pre>int32_t filter_state[4] = { 0, 0, 0, 0 }; int32_t result = dsp_filters_fir_add_sample(sample, filter_state, 5);</pre>
Type	<pre>void dsp_filters_fir_add_sample(int32_t input_sample, int32_t state_data[], const int32_t num_taps)</pre>
Parameters	<code>input_sample</code> The new sample to be processed. <code>state_data</code> Pointer to filter state data array of length $N-1$. Must be initialized at startup to all zeros. <code>num_taps</code> Number of filter taps ($N = \text{num_taps} = \text{filter order} + 1$).
Returns	Void.

5 Filter Functions: Interpolating FIR Filter

Function	<code>dsp_filters_interpolate</code>
Description	This function implements an interpolating FIR filter. The function operates on a single input sample and outputs a set of samples representing the interpolated data, whose sample count is equal to <code>interp_factor</code>. (i.e. and each call to the function processes one sample and results in <code>interp_factor</code> output samples). The FIR filter algorithm is based upon a sequence of multiply-accumulate (MAC) operations. Each filter coefficient <code>b[i]</code> is multiplied by a state variable which equals a previous input sample, or $y[n] = x[n]*b_0 + x[n-1]*b_1 + x[n-2]*b_2 + \dots + x[n-N+1]*b_{N-1}$. <code>filter_coeffs</code> points to a coefficient array of size $N = \text{num_taps}$. The filter coefficients are stored in forward order (e.g. <code>b0, b1, \dots, b_{N-1}</code>). Multiplication results are accumulated in a 64-bit accumulator. If overflow occurs in the final 64-bit result, it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. The saturation is only done after the last multiplication. To avoid 64-bit overflow in the intermediate results, the fixed point format must be chosen according to <code>num_taps</code>.
Type	<pre>void dsp_filters_interpolate(int32_t input_sample, const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_taps, const int32_t interp_factor, int32_t output_samples[], const int32_t q_format)</pre>

Continued on next page

Parameters	input_sample The new sample to be processed. filter_coeffs Pointer to FIR coefficients array arranged as: $b_M, b_{(1L+M)}, b_{(2L+M)}, b_{((N-1)L+M)}, b_1, b_{(1L+1)}, b_{(2L+1)}, b_{((N-1)L+1)}, b_0, b_{(1L+0)}, b_{(2L+0)}, b_{((N-1)L+0)}$, where $M = N-1$ state_data Pointer to filter state data array of length $N-1$. Must be initialized at startup to all zeros. num_taps Number of filter taps ($N = \text{num_taps} = \text{filter order} + 1$). interp_factor The interpolation factor/index (i.e. the up-sampling ratio). The interpolation factor/index can range from 2 to 16. output_samples The resulting interpolated samples. q_format Fixed point format (i.e. number of fractional bits).
------------	--

6 Filter Functions: Decimating FIR Filter

Function	dsp_filters_decimate
Description	This function implements an decimating FIR filter. The function operates on a single set of input samples whose count is equal to the decimation factor. (i.e. and each call to the function processes <code>decim_factor</code> samples and results in one sample). The FIR filter algorithm is based upon a sequence of multiply-accumulate (MAC) operations. Each filter coefficient <code>b[i]</code> is multiplied by a state variable which equals a previous input sample, or $y[n] = x[n]*b_0 + x[n-1]*b_1 + x[n-2]*b_2 + \dots + x[n-N+1]*b_{N-1}$ <code>filter_coeffs</code> points to a coefficient array of size $N = \text{num_taps}$. The filter coefficients are stored in forward order (e.g. <code>b0, b1, \dots, b_{N-1}</code>). The FIR algorithm involves multiplication between 32-bit filter coefficients and 32-bit state data producing a 64-bit result for each coefficient and state data pair. Multiplication results are accumulated in a 64-bit accumulator. If overflow occurs in the final 64-bit result, it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. The saturation is only done after the last multiplication. To avoid 64-bit overflow in the intermediate results, the fixed point format must be chosen according to <code>num_taps</code>.
Type	<pre>int32_t dsp_filters_decimate(int32_t input_samples[], const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_taps, const int32_t decim_factor, const int32_t q_format)</pre>
Parameters	input_samples The new samples to be decimated. filter_coeffs Pointer to FIR coefficients array arranged as <code>[b0, b1, b2, \dots, b_{N-1}]</code>. state_data Pointer to filter state data array of length $N-1$. Must be initialized at startup to all zeros. num_taps Number of filter taps ($N = \text{num_taps} = \text{filter order} + 1$). decim_factor The decimation factor/index (i.e. the down-sampling ratio). q_format Fixed point format (i.e. number of fractional bits).
Returns	The resulting decimated sample.

7 Filter Functions: Bi-Quadratic (BiQuad) IIR Filter

Function	dsp_filters_biquad
Description	This function implements a second order IIR filter (direct form I). The function operates on a single sample of input and output data (i.e. and each call to the function processes one sample). The IIR filter algorithm executes a difference equation on current and past input values x and past output values y: $y[n] = x[n]*b_0 + x[n-1]*b_1 + x[n-2]*b_2 + y[n-1]*-a_1 + y[n-2]*-a_2$ The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). Example showing a single Biquad filter with samples and coefficients represented in Q28 fixed-point format: <pre>int32_t filter_coeff[DSP_NUM_COEFFS_PER_BIQUAD] = { Q28(+0.5), Q28(-0.1), Q28(-0.5), ↪ Q28(-0.1), Q28(0.1) }; int32_t filter_state[DSP_NUM_STATES_PER_BIQUAD] = { 0, 0, 0, 0 }; int32_t result = dsp_filters_biquad(sample, filter_coeff, filter_state, 28);</pre> The IIR algorithm involves multiplication between 32-bit filter coefficients and 32-bit state data producing a 64-bit result for each coefficient and state data pair. Multiplication results are accumulated in a 64-bit accumulator. If overflow occurs in the final 64-bit result, it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits.
Type	<pre>int32_t dsp_filters_biquad(int32_t input_sample, const int32_t filter_coeffs[DSP_NUM_COEFFS_PER_BIQUAD], int32_t state_data[DSP_NUM_STATES_PER_BIQUAD], const int32_t q_format)</pre>
Parameters	input_sample The new sample to be processed. filter_coeffs Pointer to biquad coefficients array arranged as $[b_0, b_1, b_2, -a_1, -a_2]$. state_data Pointer to filter state data array (initialized at startup to zeros). The length of the state data array is 4. q_format Fixed point format (i.e. number of fractional bits).
Returns	The resulting filter output sample.

8 Filter Functions: Cascaded BiQuad Filter

Function	dsp_filters_biquads
Description	This function implements a cascaded direct form I BiQuad filter. The function operates on a single sample of input and output data (i.e. and each call to the function processes one sample). The IIR filter algorithm executes a difference equation on current and past input values x and past output values y: $y[n] = x[n]*b0 + x[n-1]*b1 + x[n-2]*b2 + y[n-1]*-a1 + y[n-2]*-a2$. The filter coefficients are stored in forward order (e.g. section1:b0,b1,b2,-a1,-a2,sectionN:b0,b1,b2,-a1,-a2). Example showing a 4x cascaded Biquad filter with samples and coefficients represented in Q28 fixed-point format: <pre>int32_t filter_coeff[4*DSP_NUM_COEFFS_PER_BIQUAD] = { Q28(+0.5), Q28(-0.1), Q28(-0.5), Q28(-0.1), Q28(0.1), Q28(+0.5), Q28(-0.1), Q28(-0.5), Q28(-0.1), Q28(0.1), Q28(+0.5), Q28(-0.1), Q28(-0.5), Q28(-0.1), Q28(0.1), Q28(+0.5), Q28(-0.1), Q28(-0.5), Q28(-0.1), Q28(0.1) }; int32_t filter_state[4*DSP_NUM_STATES_PER_BIQUAD] = { 0,0,0,0, 0,0,0,0, 0,0,0,0, ↪ 0,0,0,0 }; int32_t result = dsp_filters_biquads(sample, filter_coeff, filter_state, 4, 28);</pre> The IIR algorithm involves multiplication between 32-bit filter coefficients and 32-bit state data producing a 64-bit result for each coefficient and state data pair. Multiplication results are accumulated in a 64-bit accumulator. If overflow occurs in the final 64-bit result, it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits.
Type	<pre>int32_t dsp_filters_biquads(int32_t input_sample, const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_sections, const int32_t q_format)</pre>
Parameters	input_sample The new sample to be processed. filter_coeffs Pointer to biquad coefficients array for all BiQuad sections. Arranged as [section1:b0,b1,b2,-a1,-a2,...sectionN:b0,b1,b2,-a1,-a2]. state_data Pointer to filter state data array (initialized at startup to zeros). The length of the state data array is $num_sections * 4$. num_sections Number of BiQuad sections. q_format Fixed point format (i.e. number of fractional bits).
Returns	The resulting filter output sample.

9 Adaptive Filter Functions: LMS Adaptive Filter

Function	<code>dsp_adaptive_lms</code>
Description	This function implements a least-mean-squares adaptive FIR filter. LMS filters are a class of adaptive filters that adjust filter coefficients in order to create a transfer function that minimizes the error between the input and reference signals. FIR coefficients are adjusted on a per sample basis by an amount calculated from the given step size and the instantaneous error. The function operates on a single sample of input and output data (i.e. and each call to the function processes one sample and each call results in changes to the FIR coefficients). The general LMS algorithm, on a per sample basis, is to: 1) Apply the transfer function: <code>output = FIR(input)</code> 2) Compute the instantaneous error value: <code>error = reference - output</code> 3) Compute current coefficient adjustment delta: <code>delta = mu * error</code> 4) Adjust transfer function coefficients: <code>FIR_COEFFS[n] = FIR_COEFFS[n] + FIR_STATE[n] * delta</code> Example of a 100-tap LMS filter with samples and coefficients represented in Q28 fixed-point format: <pre>int32_t filter_coeff[100] = { ... not shown for brevity }; int32_t filter_state[100] = { 0, 0, 0, 0, ... not shown for brevity }; int32_t output_sample = dsp_adaptive_lms (input_sample, reference_sample, &error_sample, filter_coeff_array, filter_state_array, 100, Q28(0.01), 28);</pre> The LMS filter algorithm involves multiplication between two 32-bit values and 64-bit accumulation as a result of using a FIR filter as well as coefficient step size calculations. Multiplication results are accumulated in a 64-bit accumulator with the final result shifted to the required fixed-point format. Therefore overflow behavior of the 32-bit multiply operation and truncation behavior from final shifing of the accumulated multiplication results must be considered for both FIR operations as well as for coefficient step size calculation and FIR coefficient adjustment.
Type	<pre>int32_t dsp_adaptive_lms(int32_t input_sample, int32_t reference_sample, int32_t *error_sample, const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_taps, const int32_t mu, int32_t q_format)</pre>

Continued on next page

Parameters	<code>input_sample</code> The new sample to be processed. <code>reference_sample</code> Reference sample. <code>error_sample</code> Pointer to resulting error sample (error = reference - output) <code>filter_coeffs</code> Pointer to FIR coefficients arranged as [b0,b1,b2, ...,bN-1]. <code>state_data</code> Pointer to FIR filter state data array of length N. Must be initialized at startup to all zeros. <code>num_taps</code> Filter tap count where $N = \text{num_taps} = \text{filter order} + 1$. <code>mu</code> Coefficient adjustment step size, controls rate of convergence. <code>q_format</code> Fixed point format (i.e. number of fractional bits).
Returns	The resulting filter output sample.

10 Adaptive Filter Functions: Normalized LMS Filter

Function	dsp_adaptive_nlms
Description	This function implements a normalized LMS FIR filter. LMS filters are a class of adaptive filters that adjust filter coefficients in order to create the a transfer function that minimizes the error between the input and reference signals. FIR coefficients are adjusted on a per sample basis by an amount calculated from the given step size and the instantaneous error. The function operates on a single sample of input and output data (i.e. and each call to the function processes one sample and each call results in changes to the FIR coefficients). The general NLMS algorithm, on a per sample basis, is to: 1) Apply the transfer function: $\text{output} = \text{FIR}(\text{input})$ 2) Compute the instantaneous error value: $\text{error} = \text{reference} - \text{output}$ 3) Normalize the error using the instantaneous power computed by: $E = x[n]^2 + \dots + x[n-N+1]^2$ 4) Update error value: $\text{error} = (\text{reference} - \text{output}) / E$ 5) Compute current coefficient adjustment delta: $\text{delta} = \mu * \text{error}$ 6) Adjust transfer function coefficients: $\text{FIR_COEFFS}[n] = \text{FIR_COEFFS}[n] + \text{FIR_STATE}[n] * \text{delta}$ Example of a 100-tap NLMS filter with samples and coefficients represented in Q28 fixed-point format: <pre>int32_t filter_coeff[100] = { ... not shown for brevity }; int32_t filter_state[100] = { 0, 0, 0, 0, ... not shown for brevity }; int32_t output_sample = dsp_adaptive_nlms (input_sample, reference_sample, &error_sample, filter_coeff_array, filter_state_array, 100, Q28(0.01), 28);</pre> The LMS filter algorithm involves multiplication between two 32-bit values and 64-bit accumulation as a result of using a FIR filter as well as coefficient step size calculations. Multiplication results are accumulated in a 64-bit accumulator with the final result shifted to the required fixed-point format. Therefore overflow behavior of the 32-bit multiply operation and truncation behavior from final shifing of the accumulated multiplication results must be considered for both FIR operations as well as for coefficient step size calculation and FIR coefficient adjustment. Computing the coefficient adjustment involves taking the reciprocal of the instantaneous power computed by $E = x[n]^2 + x[n-1]^2 + \dots + x[n-N+1]^2$. The reciprocal is subject to overflow since the instantaneous power may be less than one.

Continued on next page

Type	<pre>int32_t dsp_adaptive_nlms(int32_t input_sample, int32_t reference_sample, int32_t *error_sample, const int32_t filter_coeffs[], int32_t state_data[], const int32_t num_taps, const int32_t mu, int32_t q_format)</pre>
Parameters	input_sample The new sample to be processed. reference_sample Reference sample. error_sample Pointer to resulting error sample (error = reference - output) filter_coeffs Pointer to FIR coefficients arranged as [b0,b1,b2, ...,bN-1]. state_data Pointer to FIR filter state data array of length N. Must be initialized at startup to all zeros. num_taps Filter tap count where $N = \text{num_taps} = \text{filter order} + 1$. mu Coefficient adjustment step size, controls rate of convergence. q_format Fixed point format (i.e. number of fractional bits).
Returns	The resulting filter output sample.

11 Scalar Math Functions: Multiply

Function	dsp_math_multiply
Description	Scalar multiplication. This function multiplies two scalar values and produces a result according to fixed-point format specified by the <code>q_format</code> parameter. The two operands are multiplied to produce a 64-bit result, and shifted right by <code>q_format</code> bits. Algorithm: <pre> 1) Y = X1 * X2 3) Y = Y >> q_format </pre> Example: <pre> int32_t result; result = dsp_math_multiply(Q28(-0.33), sample, 28); </pre>
Type	<pre> int32_t dsp_math_multiply(int32_t input1_value, int32_t input2_value, const int32_t q_format) </pre>
Parameters	<code>input1_value</code> Multiply operand #1. <code>input2_value</code> Multiply operand #2. <code>q_format</code> Fixed point format (i.e. number of fractional bits).
Returns	<code>input1_value * input2_value.</code>

12 Scalar Math Functions: Multiply Saturated

Function	dsp_math_multiply_sat
Description	Scalar saturated multiplication. This function multiplies two scalar values and produces a result according to fixed-point format specified by the q_format parameter. The two operands are multiplied to produce a 64-bit result, saturated at the minimum/maximum value given the fixed-point format if overflow occurs, and finally shifted right by q_format bits. Algorithm: <pre> 1) Y = X1 * X2 2) Y = min(max(Q_FORMAT_MIN, Y), Q_FORMAT_MAX, Y) 3) Y = Y >> q_format </pre> Example: <pre> int32_t result; result = dsp_math_multiply_sat(Q28(-0.33), sample, 28); </pre>
Type	<pre> int32_t dsp_math_multiply_sat(int32_t input1_value, int32_t input2_value, const int32_t q_format) </pre>
Parameters	<pre> input1_value Multiply operand #1. input2_value Multiply operand #2. q_format Fixed point format (i.e. number of fractional bits). </pre>
Returns	input1_value * input2_value.

13 Scalar Math Functions: Signed Division

Function	dsp_math_divide						
Description	Signed Division. This function divides two signed integer values and produces a result according to fixed-point format specified by the <code>q_format</code> parameter. It was optimised for performance using a dedicated instruction for unsigned long division. Example: <pre>uint32_t quotient; quotient = dsp_math_divide(divident, divisor, 24);</pre>						
Type	<pre>int32_t dsp_math_divide(int32_t dividend, int32_t divisor, uint32_t q_format)</pre>						
Parameters	<table> <tr> <td><code>dividend</code></td><td>Value to be divided</td></tr> <tr> <td><code>divisor</code></td><td>Dividing value</td></tr> <tr> <td><code>q_format</code></td><td>Fixed point32_t format (i.e. number of fractional bits).</td></tr> </table>	<code>dividend</code>	Value to be divided	<code>divisor</code>	Dividing value	<code>q_format</code>	Fixed point32_t format (i.e. number of fractional bits).
<code>dividend</code>	Value to be divided						
<code>divisor</code>	Dividing value						
<code>q_format</code>	Fixed point32_t format (i.e. number of fractional bits).						
Returns	Quotient of dividend/divisor						

14 Scalar Math Functions: Unsigned Division

Function	dsp_math_divide_unsigned						
Description	Unsigned Division. This function divides two unsigned integer values and produces a result according to fixed-point format specified by the <code>q_format</code> parameter. It was optimised for performance using a dedicated instruction for unsigned long division. Example: <pre>uint32_t quotient; quotient = dsp_math_divide_unsigned(divident, divisor, 24);</pre>						
Type	<pre>uint32_t dsp_math_divide_unsigned(uint32_t dividend, uint32_t divisor, uint32_t q_format)</pre>						
Parameters	<table> <tr> <td><code>dividend</code></td><td>Value to be divided</td></tr> <tr> <td><code>divisor</code></td><td>Dividing value</td></tr> <tr> <td><code>q_format</code></td><td>Fixed point32_t format (i.e. number of fractional bits).</td></tr> </table>	<code>dividend</code>	Value to be divided	<code>divisor</code>	Dividing value	<code>q_format</code>	Fixed point32_t format (i.e. number of fractional bits).
<code>dividend</code>	Value to be divided						
<code>divisor</code>	Dividing value						
<code>q_format</code>	Fixed point32_t format (i.e. number of fractional bits).						
Returns	Quotient of dividend/divisor						

15 Scalar Math Functions: Square Root

Function	dsp_math_sqrt
Description	Scalar square root. This function computes the square root of an unsigned input value using the Babylonian method of successive averaging. Error is ≤ 1 LSB and worst case performance is 96 cycles. Use dsp_math_int_sqrt() or dsp_math_int_sqrt64() for integer square roots.
Type	uq8_24 dsp_math_sqrt(uq8_24 x)
Parameters	x Unsigned 32-bit value in Q8.24 format
Returns	Unsigned 32-bit value in Q8.24 format

16 Scalar Math Functions: Sine

Function	dsp_math_sin
Description	This function returns the sine of a q8_24 fixed point number in radians. The input number has to be in the range $-\text{MIN_Q8_24} + \text{PI}$ and $\text{MIN_Q8_24} - \text{PI}$.
Type	q8_24 dsp_math_sin(q8_24 rad)
Parameters	rad input value in radians
Returns	sine(rad)

17 Scalar Math Functions: Cosine

Function	dsp_math_cos
Description	This function returns the cosine of a q8_24 fixed point number in radians. The input number has to be in the range $-\text{MIN_Q8_24} + \text{PI}$ and $\text{MIN_Q8_24} - \text{PI}$.
Type	q8_24 dsp_math_cos(q8_24 rad)
Parameters	rad input value in radians
Returns	cosine(rad)

18 Scalar Math Functions: Arctangent

Function	dsp_math_atan
Description	This function returns the arctangent of x. It uses an algorithm based on Cody and Waite pp. 194-216. The algorithm was optimised for accuracy (using improved precision and rounding) and performance (using a dedicated instruction for unsigned long division) Error compared to atan from math.h is ≤ 1 LSB. Performance is 84 cycles worst case. MIN_INT is an invalid input because the algorithm negates all negative inputs and there is no positive representation of MIN_INT Example: <pre>q8_24 x = Q24(0.005); q8_24 rad = dsp_math_atan(x);</pre>
Type	q8_24 dsp_math_atan(q8_24 x)
Parameters	x input value Q8.24 format.
Returns	arctangent of x in radians between $-\pi/2$.. $+\pi/2$

19 Scalar Math Functions: Exponential

Function	dsp_math_exp
Description	This function returns the natural exponent of a fixed point number. The input number has to be less than 4.8, otherwise the answer cannot be represented as a fixed point number. For input values larger than 3 there is a relatively large error in the last three bits of the answer.
Type	q8_24 dsp_math_exp(q8_24 x)
Parameters	x input value
Returns	e^x

20 Scalar Math Functions: Natural Logarithm

Function	dsp_math_log
Description	This function returns the natural logarithm (ln) of a fixed point number. The input number has to be positive.
Type	q8_24 dsp_math_log(uq8_24 x)
Parameters	x input value Q8.24 format.
Returns	ln(x).

21 Scalar Math Functions: Hyperbolic Sine

Function	dsp_math_sinh
Description	This function returns the hyperbolic sine (sinh) of a fixed point number. The input number has to be in the range [-5.5..5.5] in order to avoid overflow, and for values outside the [-4..4] range there are errors up to 4 bits in the result.
Type	q8_24 dsp_math_sinh(q8_24 x)
Parameters	x input value Q8.24 format.
Returns	sinh(x)

22 Scalar Math Functions: Hyperbolic Cosine

Function	dsp_math_cosh
Description	This function returns the hyperbolic cosine (cosh) of a fixed point number. The input number has to be in the range [-5.5..5.5] in order to avoid overflow, and for values outside the [-4..4] range there are errors up to 4 bits in the result.
Type	q8_24 dsp_math_cosh(q8_24 x)
Parameters	x input value Q8.24 format.
Returns	sinh(x)

23 Scalar Math Functions: Softplus

Function	dsp_math_softplus
Description	This function computes the softplus function, $\ln(1+\exp(-x))$. The input number has to be in q8_24 format, the output is a number between 0 and MAX_Q8_24 in q8_24 format.
Type	q8_24 dsp_math_softplus(q8_24 x)
Parameters	x input value Q8.24 format.
Returns	approximaton to $\log(1+\exp(x))$

24 Integer Math Functions: Square Root

Function	dsp_math_int_sqrt
Description	Integer square root 32 -> 16 bits. This function computes the square root of an unsigned input value use dsp_math_sqrt() for fixed point square roots.
Type	uint32_t dsp_math_int_sqrt(uint32_t x)
Parameters	x Unsigned 32-bit integer value
Returns	Unsigned 32-bit integer value

25 Integer Math Functions: 64-bit Square Root

Function	dsp_math_int_sqrt64
Description	Integer square root, 64 -> 32 bits. This function computes the square root of an unsigned input value use dsp_math_sqrt() for fixed point square roots.
Type	uint32_t dsp_math_int_sqrt64(uint64_t x)
Parameters	x Unsigned 64-bit integer value
Returns	Unsigned 32-bit integer value

26 Complex Math Functions: Add

Function	dsp_complex_add
Description	Function that adds two complex numbers that use the same fixed point representation.
Type	dsp_complex_t dsp_complex_add(dsp_complex_t a, dsp_complex_t b)
Parameters	a first complex number b second complex number
Returns	sum - it may overflow.

27 Complex Math Functions: Subtract

Function	dsp_complex_sub
Description	Function that subtracts two complex numbers that use the same fixed point representation.
Type	dsp_complex_t dsp_complex_sub(dsp_complex_t a, dsp_complex_t b)
Parameters	a first complex number b second complex number
Returns	difference - it may overflow.

28 Complex Math Functions: Multiplication

Function	dsp_complex_mul						
Description	Function that multiplies two complex numbers. The fixed point representation of one number has to be specified, the result will use the fixed point representation of the other number.						
Type	<code>dsp_complex_t dsp_complex_mul(dsp_complex_t a, dsp_complex_t b, uint32_t Q)</code>						
Parameters	<table> <tr> <td>a</td><td>first complex number</td></tr> <tr> <td>b</td><td>second complex number</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the numbers</td></tr> </table>	a	first complex number	b	second complex number	Q	Number of bits behind the binary point in one of the numbers
a	first complex number						
b	second complex number						
Q	Number of bits behind the binary point in one of the numbers						
Returns	product - it may overflow.						

29 Complex Math Functions: Inner Product

Function	dsp_complex_mul_conjugate						
Description	Function that multiplies one complex number with the conjugate of another complex number. The fixed point representation of one number has to be specified, the result will use the fixed point representation of the other number.						
Type	dsp_complex_t dsp_complex_mul_conjugate(dsp_complex_t a, dsp_complex_t b, uint32_t Q)						
Parameters	<table> <tr> <td>a</td><td>first complex number</td></tr> <tr> <td>b</td><td>second complex number</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the numbers</td></tr> </table>	a	first complex number	b	second complex number	Q	Number of bits behind the binary point in one of the numbers
a	first complex number						
b	second complex number						
Q	Number of bits behind the binary point in one of the numbers						
Returns	product - it may overflow.						

30 Complex Math Functions: Complex FIR Filter

Function	dsp_complex_fir										
Description	Function that computes the inner product of two complex vectors. The representation of one vector has to be specified, the result will use the fixed point representation of the other number.										
Type	<pre>dsp_complex_t dsp_complex_fir(dsp_complex_t a[], dsp_complex_t b[], uint32_t N, uint32_t offset, uint32_t Q)</pre>										
Parameters	<table> <tr> <td>a</td><td>first complex vector</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> <tr> <td>offset</td><td>starting point to use in vector a</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the vectors</td></tr> </table>	a	first complex vector	b	second complex vector	N	Length of the vectors	offset	starting point to use in vector a	Q	Number of bits behind the binary point in one of the vectors
a	first complex vector										
b	second complex vector										
N	Length of the vectors										
offset	starting point to use in vector a										
Q	Number of bits behind the binary point in one of the vectors										
Returns	inner product - it may overflow.										

31 Complex Math Functions: Element By Element Multiplication

Function	dsp_complex_mul_vector								
Description	Function that computes the element-by-element product of two complex vectors. The representation of one vector has to be specified, the result will use the fixed point representation of the other number. $a = a * b$								
Type	void dsp_complex_mul_vector(dsp_complex_t a[], dsp_complex_t b[], uint32_t N, uint32_t Q)								
Parameters	<table> <tr> <td>a</td><td>first complex vector, also output</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the vectors</td></tr> </table>	a	first complex vector, also output	b	second complex vector	N	Length of the vectors	Q	Number of bits behind the binary point in one of the vectors
a	first complex vector, also output								
b	second complex vector								
N	Length of the vectors								
Q	Number of bits behind the binary point in one of the vectors								

32 Complex Math Functions: Element By Element Inner Product

Function	dsp_complex_mul_conjugate_vector								
Description	Function that computes the element-by-element product of two complex vectors, where the complex conjugate of the second vector is used. The representation of one vector has to be specified, the result will use the fixed point representation of the other number: $a = a * \text{conj}(b)$								
Type	void dsp_complex_mul_conjugate_vector(dsp_complex_t a[], dsp_complex_t b[], uint32_t N, uint32_t Q)								
Parameters	<table> <tr> <td>a</td><td>first complex vector, also output</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the vectors</td></tr> </table>	a	first complex vector, also output	b	second complex vector	N	Length of the vectors	Q	Number of bits behind the binary point in one of the vectors
a	first complex vector, also output								
b	second complex vector								
N	Length of the vectors								
Q	Number of bits behind the binary point in one of the vectors								

33 Complex Math Functions: Element By Element Vector Product

Function	dsp_complex_mul_conjugate_vector3	
Description	Function that computes the element-by-element product of two complex vectors, where the complex conjugate of the second vector is used. The representation of one vector has to be specified, the result will use the fixed point representation of the other number: $o = a * conj(b)$	
Type	<pre>void dsp_complex_mul_conjugate_vector3(dsp_complex_t o[], dsp_complex_t a[], dsp_complex_t b[], uint32_t N, uint32_t Q)</pre>	
Parameters	o	first complex vector, also output
	a	first complex vector, also output
	b	second complex vector
	N	Length of the vectors
	Q	Number of bits behind the binary point in one of the vectors

34 Complex Math Functions: Element By Element Vector Add

Function	dsp_complex_add_vector	
Description	Function that computes the element-by-element sum of two complex vectors: $a = a + b$	
Type	<pre>void dsp_complex_add_vector(dsp_complex_t a[], dsp_complex_t b[], uint32_t N)</pre>	
Parameters	a	first complex vector, also output
	b	second complex vector
	N	Length of the vectors

35 Complex Math Functions: Element By Element Add With Shift Left Scaling

Function	dsp_complex_add_vector_shl		
Description	Function that computes the element-by-element sum of two complex vectors, scaling one vector by a fixed number of bit positions: $a = a + (b \ll \text{shift})$		
Type	<pre>void dsp_complex_add_vector_shl(dsp_complex_t a[], dsp_complex_t b[], uint32_t N, int32_t shift)</pre>		
Parameters	a	first complex vector, also output	
	b	second complex vector	
	N	Length of the vectors	
	shift	Number of bits to shift b left by prior to adding. negative values indicate right shift.	

36 Complex Math Functions: Element By Element Add With Scaling

Function	dsp_complex_add_vector_scale		
Description	Function that computes the element-by-element sum of two complex vectors, scaling one vector by a fixed number of bit positions: $a = a + ((b * \text{scalar}) \gg 24)$		
Type	<pre>void dsp_complex_add_vector_scale(dsp_complex_t a[], dsp_complex_t b[], uint32_t N, int32_t scalar)</pre>		
Parameters	a	first complex vector, also output	
	b	second complex vector	
	N	Length of the vectors	
	scalar	Number to multiply b with, in q8_24 notation. Use 0x01000000 to mimic dsp_complex_add_vector() .	

37 Complex Math Functions: Element By Element Subtraction

Function	dsp_complex_sub_vector	
Description	Function that computes the element-by-element sum of two complex vectors: $a = a - b$	
Type	<pre>void dsp_complex_sub_vector(dsp_complex_t a[], dsp_complex_t b[], uint32_t N)</pre>	
Parameters	a	first complex vector, also output
	b	second complex vector
	N	Length of the vectors

38 Complex Math Functions: Element By Element Addition Into A Third Vector

Function	dsp_complex_add_vector3								
Description	Function that computes the element-by-element difference of two complex vectors, into a third vector: $o = a + b$								
Type	void dsp_complex_add_vector3(dsp_complex_t o[], dsp_complex_t a[], dsp_complex_t b[], uint32_t N)								
Parameters	<table> <tr> <td>o</td><td>output complex vector</td></tr> <tr> <td>a</td><td>first complex vector</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> </table>	o	output complex vector	a	first complex vector	b	second complex vector	N	Length of the vectors
o	output complex vector								
a	first complex vector								
b	second complex vector								
N	Length of the vectors								

39 Complex Math Functions: Element By Element Subtraction Into A Third Vector

Function	dsp_complex_sub_vector3		
Description	Function that computes the element-by-element difference of two complex vectors, into a third vector: o = a - b		
Type	void dsp_complex_sub_vector3(dsp_complex_t o[], dsp_complex_t a[], dsp_complex_t b[], uint32_t N)		
Parameters	o	output complex vector	
	a	first complex vector	
	b	second complex vector	
	N	Length of the vectors	

40 Complex Math Functions: Element By Element Multiply Accumulate

Function	dsp_complex_macc_vector										
Description	Function that computes the element-by-element product of two complex vectors, and adds the result to a third vector: $a = a + b * c$										
Type	void dsp_complex_macc_vector(dsp_complex_t a[], dsp_complex_t b[], dsp_complex_t c[], uint32_t N, int Q)										
Parameters	<table> <tr> <td>a</td><td>first complex vector, also output</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>c</td><td>third complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the numbers</td></tr> </table>	a	first complex vector, also output	b	second complex vector	c	third complex vector	N	Length of the vectors	Q	Number of bits behind the binary point in one of the numbers
a	first complex vector, also output										
b	second complex vector										
c	third complex vector										
N	Length of the vectors										
Q	Number of bits behind the binary point in one of the numbers										

41 Complex Math Functions: Element By Element Multiply Subtract

Function	dsp_complex_nmaccc_vector										
Description	Function that computes the element-by-element product of two complex vectors, and subtracts the result from a third vector: $a = a - b * c$										
Type	void dsp_complex_nmaccc_vector(dsp_complex_t a[], dsp_complex_t b[], dsp_complex_t c[], uint32_t N, int Q)										
Parameters	<table> <tr> <td>a</td><td>first complex vector, also output</td></tr> <tr> <td>b</td><td>second complex vector</td></tr> <tr> <td>c</td><td>third complex vector</td></tr> <tr> <td>N</td><td>Length of the vectors</td></tr> <tr> <td>Q</td><td>Number of bits behind the binary point in one of the numbers</td></tr> </table>	a	first complex vector, also output	b	second complex vector	c	third complex vector	N	Length of the vectors	Q	Number of bits behind the binary point in one of the numbers
a	first complex vector, also output										
b	second complex vector										
c	third complex vector										
N	Length of the vectors										
Q	Number of bits behind the binary point in one of the numbers										

42 Complex Math Functions: Multiply Scalar And Shift Down

Function	dsp_complex_scalar_vector3	
Description	Function that multiplies a complex vector with a scalar, and shifts the data down. The scalar is a 8.24 number $a = \text{scalar} * b \gg Q$	
Type	<pre>void dsp_complex_scalar_vector3(dsp_complex_t ALIAS_PTR a, dsp_complex_t ALIAS_PTR b, uint32_t N, int32_t scalar, uint32_t Q)</pre>	
Parameters	a	first complex vector, also output
	b	second complex vector
	N	Length of the vectors
	scalar	Scalar multiplier
	Q	Number of bits behind the binary point in one of the vectors

43 Complex Math Functions: Return Magnitude With Selectable Precision

Function	dsp_complex_magnitude_vector
Description	Function that computes the magnitude of an array of complex numbers. The function allows for a trade-off to be made between accuracy and speed. The precision parameter should be set to (24-P) to request P bits accuracy in the magnitude. Hence, the value 0 requests a 24-bit accurate magnitude, whereas 23 requests a 1-bit accurate magnitude. precision must be between 0 and 23 inclusive. Execution time is ((31 + 5 N) x P) thread cycles.
Type	void dsp_complex_magnitude_vector(uint32_t magnitude[], dsp_complex_t input[], uint32_t N, uint32_t P)
Parameters	magnitude array in which magnitudes are stored input Array of complex numbers of which to compute magnitude N Number of elements in the input and output arrays P Integer that sets the precision; set to 0 for maximum precision; and to 23 for no precision; 0 <= precision <= 23.

44 Complex Math Functions: Scale By A Fraction

Function	<code>dsp_complex_scale_vector</code>
Description	Function that scales an array of complex numbers by a fraction. It requires an array of complex number, and an array of numerators and an array of denominators. It computes: $\text{array} = \text{array} * \text{numerator} / \text{denominator}$ Note that both numerator and denominator are unsigned values. The typical calling sequence below may be used to reset the magnitude of each element of a complex vector to a new value: <pre>dsp_complex_magnitude_vector(cur_magnitude, input, N, 0); dsp_complex_remagnitude_vector(input, cur_magnitude, new_mag, N);</pre>
Type	<pre>void dsp_complex_scale_vector(dsp_complex_t array[], uint32_t numerator[], uint32_t denominator[], uint32_t N)</pre>
Parameters	<code>array</code> Array of complex numbers on which to reset magnitude <code>numerator</code> Array of ints that store the current magnitude <code>denominator</code> Array of ints that store the desired magnitude <code>N</code> Number of elements in the input and output arrays

45 Complex Math Functions: Combine Separate Real And Imaginary Arrays Into Interleaved Complex

Function	dsp_complex_combine		
Description	Function that combines an array of real numbers and an array of imaginary numbers into an interleaved array of complex numbers.		
Type	<pre>void dsp_complex_combine(const int32_t re[], const int32_t im[], dsp_complex_t complex[], const uint32_t N)</pre>		
Parameters	re	Array of real numbers to combine into complex array	
	im	Array of imaginary numbers to combine into complex array	
	complex	Array of complex numbers, with interleaved real and imaginary value	
	N	Number of elements in the input and output arrays	

46 Complex Math Functions: Split Interleaved Complex Array Into Separate Real And Imaginary

Function	dsp_complex_split	
Description	Function that splits an array of complex numbers into separate arrays for the real and imaginary numbers.	
Type	<pre>void dsp_complex_split(const dsp_complex_t complex[], int32_t re[], int32_t im[], const uint32_t N)</pre>	
Parameters	complex	Array of complex numbers, with interleaved real and imaginary value
	re	Array of real numbers split from complex array
	im	Array of imaginary numbers split from complex array
	N	Number of elements in the input and output arrays

47 Vector Math Functions: Minimum Value

Function	dsp_vector_minimum
Description	Vector Minimum. Locate the vector's first occurring minimum value, returning the index of the first occurring minimum value. Example: <pre>int32_t samples[256]; int32_t index = dsp_vector_minimum(samples, 256);</pre>
Type	int32_t dsp_vector_minimum(const int32_t input_vector[], const int32_t vector_length)
Parameters	input_vector Pointer to source data array. vector_length Length of the output and input arrays.
Returns	Array index where first minimum value occurs.

48 Vector Math Functions: Maximum Value

Function	dsp_vector_maximum
Description	Vector Minimum. Locate the vector's first occurring maximum value, returning the index of the first occurring maximum value. Example: <pre>int32_t samples[256]; int32_t index = dsp_vector_maximum(samples, 256);</pre>
Type	int32_t dsp_vector_maximum(const int32_t input_vector[], const int32_t vector_length)
Parameters	input_vector Pointer to source data array. vector_length Length of the output and input arrays.
Returns	Array index where first maximum value occurs.

49 Vector Math Functions: Element Negation

Function	dsp_vector_negate
Description	Vector negation: $R[i] = -X[i]$. This function sets each result element to the negative value of the corresponding input element. Each negated element is computed by twos-compliment negation therefore the minimum negative fixed-point value can not be negated to generate its corresponding maximum positive fixed-point value. For example: -Q28(-8.0) will not result in a fixed-point value representing +8.0. Example: <pre>int32_t samples[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t result[256]; dsp_vector_negate(samples, result, 256);</pre>
Type	<pre>void dsp_vector_negate(const int32_t input_vector_X[], int32_t result_vector_R[], const int32_t vector_length)</pre>
Parameters	input_vector_X Pointer/reference to source data. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors.

50 Vector Math Functions: Element Absolute Value

Function	dsp_vector_abs
Description	Vector absolute value: $R[i] = X[i]$. This function sets each element of the result vector to the absolute value of the corresponding input vector element. Example: <pre>int32_t samples[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t result[256]; dsp_vector_abs(samples, result, 256);</pre> If an element is less than zero it is negated to compute its absolute value. Negation is computed via two's-complement negation therefore the minimum negative fixed-point value can not be negated to generate its corresponding maximum positive fixed-point value. For example: -Q28(-8.0) will not result in a fixed-point value representing +8.0.
Type	<pre>void dsp_vector_abs(const int32_t input_vector_X[], int32_t result_vector_R[], const int32_t vector_length)</pre>
Parameters	input_vector_X Pointer/reference to source data. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors.

51 Vector Math Functions: Scalar Addition

Function	dsp_vector_adds
Description	Vector / scalar addition: $R[i] = X[i] + A$. This function adds a scalar value to each vector element. 32-bit addition is used to compute the scalar plus vector element result. Therefore fixed-point value overflow conditions should be observed. The resulting values are not saturated. Example: <pre>int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_scalar_A = Q28(0.333); int32_t result_vector_R[256]; dsp_vector_adds(input_vector_X, scalar_value_A, result_vector_R, 256 ↪);</pre>
Type	<pre>void dsp_vector_adds(const int32_t input_vector_X[], int32_t input_scalar_A, int32_t result_vector_R[], const int32_t vector_length)</pre>
Parameters	input_vector_X Pointer/reference to source data array X input_scalar_A Scalar value to add to each input element result_vector_R Pointer to the resulting data array vector_length Length of the input and output vectors

52 Vector Math Functions: Scalar Multiplication

Function	dsp_vector_muls
Description	Vector / scalar multiplication: $R[i] = X[i] * A$. The elements in vector X are multiplied with the scalar A and stored in result vector R. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits. Example: <pre>int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_scalar_A = Q28(0.333); int32_t result_vector_R[256]; dsp_vector_muls(input_vector_X, scalar_value_A, result_vector_R, 256 ↪);</pre>
Type	<pre>void dsp_vector_muls(const int32_t input_vector_X[], int32_t input_scalar_A, int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	input_vector_X Pointer/reference to source data array X. input_scalar_A Scalar value to multiply each element by. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors. q_format Fixed point format, the number of bits making up fractional part.

53 Vector Math Functions: Vector Addition

Function	dsp_vector_addv
Description	Vector / vector addition: $R[i] = X[i] + Y[i]$. 32 bit addition is used to add Vector X to Vector Y. The results are stored in result vector R. They are not saturated so overflow may occur. Example: <pre> int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_vector_Y[256] = { 0, -1, -2, -3, ... not shown for ↪ brevity }; int32_t result_vector_R[256]; dsp_vector_addv(input_vector_X, input_vector_Y, result_vector_R, 256 ↪); </pre>
Type	<pre> void dsp_vector_addv(const int32_t input_vector_X[], const int32_t input_vector_Y[], int32_t result_vector_R[], const int32_t vector_length) </pre>
Parameters	input_vector_X Pointer to source data array X. input_vector_Y Pointer to source data array Y. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors.

54 Vector Math Functions: Vector Subtraction

Function	dsp_vector_subv
Description	Vector / vector subtraction: $R[i] = X[i] - Y[i]$. 32 bit subtraction is used to subtract Vector Y from Vector X. The results are stored in result vector R. They are not saturated so overflow may occur. Example: <pre> int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_vector_Y[256] = { 0, -1, -2, -3, ... not shown for ↪ brevity }; int32_t result_vector_R[256]; dsp_vector_subv(input_vector_X, input_vector_Y, result_vector_R, 256 ↪); </pre>
Type	<pre> void dsp_vector_subv(const int32_t input_vector_X[], const int32_t input_vector_Y[], int32_t result_vector_R[], const int32_t vector_length) </pre>
Parameters	input_vector_X Pointer to source data array X. input_vector_Y Pointer to source data array Y. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors.

55 Vector Math Functions: Vector Multiplication

Function	dsp_vector_mulv
Description	Vector / vector multiplication: $R[i] = X[i] * Y[i]$. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. Example: <pre>int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_vector_Y[256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t result_vector_R[256]; dsp_vector_mulv(input_vector_X, input_vector_Y, result_vector_R, 256, ↪ 28);</pre>
Type	<pre>void dsp_vector_mulv(const int32_t input_vector_X[], const int32_t input_vector_Y[], int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>input_vector_Y</code> Pointer to source data array Y. <code>result_vector_R</code> Pointer to the resulting data array. <code>vector_length</code> Length of the input and output vectors. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

56 Vector Math Functions: Vector multiplication and scalar addition

Function	dsp_vector_mulv_adds
Description	Vector multiplication and scalar addition: $R[i] = X[i] * Y[i] + A$. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. Then the scalar is added. Overflow may occur after the addition. Example: <pre>int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_vector_Y[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_scalar_A = Q28(0.333); int32_t result_vector_R[256]; dsp_vector_mulv_adds(input_vector_X, input_vector_Y, scalar_value_A, ↪ result_vector_R, 256, 28);</pre>
Type	<pre>void dsp_vector_mulv_adds(const int32_t input_vector_X[], const int32_t input_vector_Y[], int32_t input_scalar_A, int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>input_vector_Y</code> Pointer to source data array Y. <code>input_scalar_A</code> Scalar value to add to each X*Y result. <code>result_vector_R</code> Pointer to the resulting data array. <code>vector_length</code> Length of the input and output vectors. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

57 Vector Math Functions: Scalar multiplication and vector addition

Function	dsp_vector_muls_addv
Description	Scalar multiplication and vector addition: $R[i] = X[i] * A + Y[i]$. The elements in vector X are multiplied with the scalar A and then added to the corresponding element in Y. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits. Overflow may occur after the addition Y[i]. Example: <pre>int32_t input_vector_X[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_scalar_A = Q28(0.333); int32_t input_vector_Y[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t result_vector_R[256]; dsp_vector_muls_addv(input_vector_X, input_scalar_A, input_vector_Y, ↪ result_vector_R, 256, 28);</pre>
Type	<pre>void dsp_vector_muls_addv(const int32_t input_vector_X[], int32_t input_scalar_A, const int32_t input_vector_Y[], int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	input_vector_X Pointer to source data array X. input_scalar_A Scalar value to multiply each element by. input_vector_Y Pointer to source data array Y result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors q_format Fixed point format (i.e. number of fractional bits).

58 Vector Math Functions: Scalar multiplication and vector subtraction

Function	dsp_vector_muls_subv
Description	Scalar multiplication and vector subtraction: $R[i] = X[i] * A - Y[i]$. The elements in vector X are multiplied with the scalar A. The corresponding element in Y is subtracted from that. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits. Overflow may occur after the subtraction of Y[i]. Example: <pre>int32_t input_vector_X[256]; int32_t input_scalar_A = Q28(0.333); int32_t input_vector_Y[256]; int32_t result_vector_R[256]; dsp_vector_muls_subv(input_vector_X, input_scalar_A, input_vector_Y, ↪ result_vector_R, 256, 28);</pre>
Type	<pre>void dsp_vector_muls_subv(const int32_t input_vector_X[], int32_t input_scalar_A, const int32_t input_vector_Y[], int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	input_scalar_A Scalar value to multiply each element by. input_vector_X Pointer to source data array X. input_vector_Y Pointer to source data array Y. result_vector_R Pointer to the resulting data array. vector_length Length of the input and output vectors. q_format Fixed point format (i.e. number of fractional bits).

59 Vector Math Functions: Vector multiplication and vector addition

Function	dsp_vector_mulv_addv
Description	Vector multiplication and vector addition: $R[i] = X[i] * Y[i] + Z[i]$. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. Overflow may occur after the addition of <code>Z[i]</code>. Example: <pre>int32_t input_vector_X[256]; int32_t input_vector_Y[256]; int32_t input_vector_Z[256]; int32_t result_vector_R[256]; dsp_vector_mulv_addv(input_vector_X, input_vector_Y, input_vector_Z, ↪ result_vector_R, 256, 28);</pre>
Type	<pre>void dsp_vector_mulv_addv(const int32_t input_vector_X[], const int32_t input_vector_Y[], const int32_t input_vector_Z[], int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>input_vector_Y</code> Pointer to source data array Y. <code>input_vector_Z</code> Pointer to source data array Z. <code>result_vector_R</code> Pointer to the resulting data array. <code>vector_length</code> Length of the input and output vectors. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

60 Vector Math Functions: Vector multiplication and vector subtraction

Function	dsp_vector_mulv_subv
Description	Vector multiplication and vector subtraction: $R[i] = X[i] * Y[i] - Z[i]$. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by <code>q_format</code> bits. Overflow may occur after the subtraction of <code>Z[i]</code>. Example: <pre>int32_t input_vector_X[256]; int32_t input_vector_Y[256]; int32_t input_vector_Z[256]; int32_t result_vector_R[256]; dsp_vector_mulv_subv(input_vector_X, input_vector_Y, input_vector_Z, ↪ result_vector_R, 256, 28);</pre>
Type	<pre>void dsp_vector_mulv_subv(const int32_t input_vector_X[], const int32_t input_vector_Y[], const int32_t input_vector_Z[], int32_t result_vector_R[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>input_vector_Y</code> Pointer to source data array Y. <code>input_vector_Z</code> Pointer to source data array Z. <code>result_vector_R</code> Pointer to the resulting data array. <code>vector_length</code> Length of the input and output vectors. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

61 Vector Math Functions: Complex vector multiplication

Function	dsp_vector_mulv_complex
Description	Complex vector / vector multiplication: $R[i] = X[i] * Y[i]$. Vectors $X[i]$ and $Y[i]$ are complex, with separate arrays for the real and imaginary components. Each multiplication produces a 64-bit result. If overflow occurs it is saturated at the minimum/maximum value given the fixed-point format and finally shifted right by q_format bits. Example: <pre> int32_t input_vector_X_re[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_vector_X_im[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_vector_Y_re[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t input_vector_Y_im[256] = { 0, 1, 2, 3, ... not shown for brevity }; int32_t result_vector_R_re[256]; int32_t result_vector_R_im[256]; dsp_vector_mulv_complex(input_vector_X_re, input_vector_X_im, input_vector_Y_re, ↪ input_vector_Y_im, result_vector_R_re, result_vector_R_im, 256, 28); </pre>
Type	<pre> void dsp_vector_mulv_complex(const int32_t input_vector_X_re[], const int32_t input_vector_X_im[], const int32_t input_vector_Y_re[], const int32_t input_vector_Y_im[], int32_t result_vector_R_re[], int32_t result_vector_R_im[], const int32_t vector_length, const int32_t q_format) </pre>
Parameters	input_vector_X_re Pointer to real source data array X. input_vector_X_im Pointer to imaginary source data array X. input_vector_Y_re Pointer to real source data array Y. input_vector_Y_im Pointer to imaginary source data array Y. result_vector_R_re Pointer to the resulting real data array. result_vector_R_im Pointer to the resulting imaginary data array. vector_length Length of the input and output vectors. q_format Fixed point format (i.e. number of fractional bits).

62 Matrix Math Functions: Element Negation

Function	dsp_matrix_negate
Description	Matrix negation: $R[i][j] = -X[i][j]$. Each negated element is computed by twos-compliment negation therefore the minimum negative fixed-point value can not be negated to generate its corresponding maximum positive fixed-point value. For example: -Q28(-8.0) will not result in a fixed-point value representing +8.0. Example: <pre>int32_t samples[8][32]; int32_t result[8][32]; dsp_matrix_negate(samples, result, 8, 32);</pre>
Type	<pre>void dsp_matrix_negate(const int32_t input_matrix_X[], int32_t result_matrix_R[], const int32_t row_count, const int32_t column_count)</pre>
Parameters	input_matrix_X Pointer/reference to source data. result_matrix_R Pointer to the resulting 2-dimensional data array. row_count Number of rows in input and output matrices. column_count Number of columns in input and output matrices.

63 Matrix Math Functions: Scalar Addition

Function	dsp_matrix_adds
Description	Matrix / scalar addition: $R[i][j] = X[i][j] + A$. 32-bit addition is used to compute the scalar plus matrix element result. Therefore fixed-point value overflow conditions should be observed. The resulting values are not saturated. Example: <pre>int32_t input_matrix_X[8][32]; int32_t input_scalar_A = Q28(0.333); int32_t result_vector_R[8][32]; dsp_matrix_adds(input_matrix_X, scalar_matrix_A, result_matrix_R, 8, ↪ 32);</pre>
Type	<pre>void dsp_matrix_adds(const int32_t input_matrix_X[], int32_t input_scalar_A, int32_t result_matrix_R[], const int32_t row_count, const int32_t column_count)</pre>
Parameters	input_matrix_X Pointer/reference to source data. input_scalar_A Scalar value to add to each input element. result_matrix_R Pointer to the resulting 2-dimensional data array. row_count Number of rows in input and output matrices. column_count Number of columns in input and output matrices.

64 Matrix Math Functions: Scalar Multiplication

Function	dsp_matrix_muls
Description	Matrix / scalar multiplication: $R[i][j] = X[i][j] * A$. Each element of the input matrix is multiplied by a scalar value using a 32bit multiply 64-bit accumulate function therefore fixed-point multiplication and q-format adjustment overflow behavior must be considered (see behavior for the function dsp_math_multiply). Example: <pre>int32_t input_matrix_X[8][32]; int32_t input_scalar_A = Q28(0.333); int32_t result_vector_R[8][32]; dsp_matrix_muls(input_matrix_X, scalar_value_A, result_matrix_R, 256, 8, 32, 28);</pre>
Type	<pre>void dsp_matrix_muls(const int32_t input_matrix_X[], int32_t input_scalar_A, int32_t result_matrix_R[], const int32_t row_count, const int32_t column_count, const int32_t q_format)</pre>
Parameters	input_matrix_X Pointer/reference to source data X. input_scalar_A Scalar value to multiply each element by. result_matrix_R Pointer to the resulting 2-dimensional data array. row_count Number of rows in input and output matrices. column_count Number of columns in input and output matrices. q_format Fixed point format (i.e. number of fractional bits).

65 Matrix Math Functions: Matrix Addition

Function	dsp_matrix_addm
Description	Matrix / matrix addition: $R[i][j] = X[i][j] + Y[i][j]$. 32-bit addition is used to compute the result for each element. Therefore fixed-point value overflow conditions should be observed. The resulting values are not saturated. Example: <pre> int32_t input_matrix_X [256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_matrix_Y [256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t result_matrix_R[256]; dsp_matrix_addm(input_matrix_X, input_matrix_Y, result_matrix_R, 8, ↪ 32); </pre>
Type	<pre> void dsp_matrix_addm(const int32_t input_matrix_X[], const int32_t input_matrix_Y[], int32_t result_matrix_R[], const int32_t row_count, const int32_t column_count) </pre>
Parameters	input_matrix_X Pointer to source data array X. input_matrix_Y Pointer to source data array Y. result_matrix_R Pointer to the resulting 2-dimensional data array. row_count Number of rows in input and output matrices. column_count Number of columns in input and output matrices.

66 Matrix Math Functions: Matrix Subtraction

Function	dsp_matrix_subm
Description	Matrix / matrix subtraction: $R[i][j] = X[i][j] - Y[i][j]$. 32-bit subtraction is used to compute the result for each element. Therefore fixed-point value overflow conditions should be observed. The resulting values are not saturated. Example: <pre> int32_t input_matrix_X [256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t input_matrix_Y [256] = { 0, 1, 2, 3, ... not shown for brevity ↪ }; int32_t result_matrix_R[256]; dsp_matrix_subm(input_matrix_X, input_matrix_Y, result_matrix_R, 8, ↪ 32); </pre>
Type	<pre> void dsp_matrix_subm(const int32_t input_matrix_X[], const int32_t input_matrix_Y[], int32_t result_matrix_R[], const int32_t row_count, const int32_t column_count) </pre>
Parameters	input_matrix_X Pointer to source data array X. input_matrix_Y Pointer to source data array Y. result_matrix_R Pointer to the resulting 2-dimensional data array. row_count Number of rows in input and output matrices. column_count Number of columns in input and output matrices.

67 Matrix Math Functions: Matrix Multiplication

Function	dsp_matrix_mulm
Description	Matrix / matrix multiplication: $R[i][j] = X[i][j] * Y[i][j]$. Elements in each of the input matrices are multiplied together using a 32bit multiply 64-bit accumulate function therefore fixed-point multiplication and q-format adjustment overflow behavior must be considered (see behavior for the function dsp_math_multiply). The algorithm is optimised for performance using double word load and store instructions. As a result the matrices must have an even number of rows and columns. Example: $M \times N * N \times P = M \times P$ <pre>int32_t input_matrix_X[rows_X][N]; int32_t input_matrix_Y[columns_Y][N]; // transposed for better memory alignment !! int32_t result_matrix_R[rows_X][columns_Y]; dsp_matrix_mulm(input_matrix_X, input_matrix_Y, result_matrix_R, 256, 8, 32, 28);</pre>
Type	<pre>void dsp_matrix_mulm(const int32_t input_matrix_X[], const int32_t input_matrix_Y[], int32_t result_matrix_R[], const int32_t rows_X, const int32_t cols_Y, const int32_t cols_X_rows_Y, const int32_t q_format)</pre>
Parameters	input_matrix_X Pointer to source data array X. input_matrix_Y Pointer to source data array Y. result_matrix_R Pointer to the resulting 2-dimensional data array. rows_X Number of rows in input matrix X. Must be even or will trap. cols_Y Number of columns input matrix Y. Must be even or will trap. cols_X_rows_Y Number of columns in input matrix X == rows in input matrix Y. Must be even or will trap. q_format Fixed point format (i.e. number of fractional bits).

68 Statistics Functions: Vector Absolute Sum

Function	dsp_vector_abs_sum
Description	Vector absolute sum: $R = (\text{abs}(X[0]) + \text{abs}(X[1]) + \dots + \text{abs}(X[N-1]))$. This function computes the absolute sum of all vector elements. Due to successive 32-bit additions being accumulated using 64-bit arithmetic overflow during the summation process is unlikely. The final value, being effectively the result of a left-shift by <code>q_format</code> bits will potentially overflow the final fixed-point value depending on the resulting summed value and the chosen Q-format. Example: <pre>int32_t result; result = dsp_vector_abs_sum(input_vector, 256, 28);</pre>
Type	<code>int32_t</code> <pre>dsp_vector_abs_sum(const int32_t *input_vector_X, const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>vector_length</code> Length of the input vector. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

69 Statistics Functions: Vector Mean

Function	dsp_vector_mean
Description	Vector mean: $R = (X[0] + X[1] + \dots + X[N-1]) / N$. This function computes the mean of the values contained within the input vector. Due to successive 32-bit additions being accumulated using 64-bit arithmetic overflow during the summation process is unlikely. The final value, being effectively the result of a left-shift by q_format bits will potentially overflow the final fixed-point value depending on the resulting summed value and the chosen Q-format. Example: <pre>int32_t result; result = dsp_vector_mean(input_vector, 256, 28);</pre>
Type	<pre>int32_t dsp_vector_mean(const int32_t input_vector_X[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	input_vector_X Pointer to source data array X. vector_length Length of the input vector. q_format Fixed point format (i.e. number of fractional bits).

70 Statistics Functions: Vector Power (Sum-of-Squares)

Function	dsp_vector_power
Description	Vector power (sum of squares): $R = X[0]^2 + X[1]^2 + \dots + X[N-1]^2$. This function computes the power (also known as the sum-of-squares) of the values contained within the input vector. Since each element in the vector is squared the behavior for fixed-point multiplication should be considered (see behavior for the function <code>dsp_math_multiply</code>). Due to successive 32-bit additions being accumulated using 64-bit arithmetic overflow during the summation process is unlikely. The final value, being effectively the result of a left-shift by <code>q_format</code> bits will potentially overflow the final fixed-point value depending on the resulting summed value and the chosen Q-format. Example: <pre>int32_t result; result = dsp_vector_power(input_vector, 256, 28);</pre>
Type	<code>int32_t</code> <pre>dsp_vector_power(const int32_t input_vector_X[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>vector_length</code> Length of the input vector. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

71 Statistics Functions: Root Mean Square (RMS)

Function	dsp_vector_rms
Description	Vector root mean square: $R = ((X[0]^2 + X[1]^2 + \dots + X[N-1]^2) / N)^{0.5}$. This function computes the root-mean-square (RMS) of the values contained within the input vector. <pre> result = 0 for i = 0 to N-1: result += input_vector_X[i] return dsp_math_sqrt(result / vector_length) </pre> Since each element in the vector is squared the behavior for fixed-point multiplication should be considered (see behavior for the function dsp_math_multiply). Due to successive 32-bit additions being accumulated using 64-bit arithmetic overflow during the summation process is unlikely. The squareroot of the 'sum-of-squares divided by N uses the function dsp_math_sqrt; see behavior for that function. The final value, being effectively the result of a left-shift by q_format bits will potentially overflow the final fixed-point value depending on the resulting summed value and the chosen Q-format. Example: <pre> int32_t result; result = dsp_vector_rms(input_vector, 256, 28); </pre>
Type	<pre> int32_t dsp_vector_rms(const int32_t input_vector_X[], const int32_t vector_length, const int32_t q_format) </pre>
Parameters	input_vector_X Pointer to source data array X. vector_length Length (N) of the input vector. q_format Fixed point format (i.e. number of fractional bits).

72 Statistics Functions: Dot Product

Function	dsp_vector_dotprod
Description	Vector dot product: $R = X[0] * Y[0] + X[1] * Y[1] + \dots + X[N-1] * Y[N-1]$. This function computes the dot-product of two equal length vectors. The elements in the input vectors are multiplied before being summed therefore fixed-point multiplication behavior must be considered (see behavior for the function <code>dsp_math_multiply</code>). Due to successive 32-bit additions being accumulated using 64-bit arithmetic overflow during the summation process is unlikely. The final value, being effectively the result of a left-shift by <code>q_format</code> bits will potentially overflow the final fixed-point value depending on the resulting summed value and the chosen Q-format. Example: <pre>int32_t result; result = dsp_vector_dotprod(input_vector_X, input_vector_Y, 256, 28) ↵ ;</pre>
Type	<code>int32_t</code> <pre>dsp_vector_dotprod(const int32_t input_vector_X[], const int32_t input_vector_Y[], const int32_t vector_length, const int32_t q_format)</pre>
Parameters	<code>input_vector_X</code> Pointer to source data array X. <code>input_vector_Y</code> Pointer to source data array Y. <code>vector_length</code> Length of the input vectors. <code>q_format</code> Fixed point format (i.e. number of fractional bits).

73 Filter Design Functions: Notch Filter

Function	dsp_design_biquad_notch
Description	This function generates BiQuad filter coefficients for a notch filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_notch(0.25, 0.707, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_notch(double filter_frequency, double filter_Q, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter center frequency normalized to the sampling frequency. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

74 Filter Design Functions: Low-pass Filter

Function	dsp_design_biquad_lowpass
Description	This function generates BiQuad filter coefficients for a low-pass filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_lowpass(0.25, 0.707, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_lowpass(double filter_frequency, double filter_Q, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter cutoff (-3db) frequency normalized to the sampling frequency. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

75 Filter Design Functions: High-pass Filter

Function	dsp_design_biquad_highpass
Description	This function generates BiQuad filter coefficients for a high-pass filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_highpass(0.25, 0.707, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_highpass(double filter_frequency, double filter_Q, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter cutoff (-3db) frequency normalized to the sampling frequency. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

76 Filter Design Functions: All-pass Filter

Function	dsp_design_biquad_allpass
Description	This function generates BiQuad filter coefficients for an all-pass filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_allpass(0.25, 0.707, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_allpass(double filter_frequency, double filter_Q, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter center frequency normalized to the sampling frequency. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

77 Filter Design Functions: Band-pass Filter

Function	dsp_design_biquad_bandpass
Description	This function generates BiQuad filter coefficients for a band-pass filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_bandpass(0.20, 0.30, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_bandpass(double filter_frequency1, double filter_frequency2, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency1 Filter cutoff #1 (-3db) frequency normalized to the sampling frequency. $0 < \text{frequency1} < 0.5$, where 0.5 represents $F_s/2$. filter_frequency2 Filter cutoff #2 (-3db) frequency normalized to the sampling frequency. $0 < \text{frequency2} < 0.5$, where 0.5 represents $F_s/2$. Note that frequency1 must be less than frequency2. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

78 Filter Design Functions: Peaking Filter

Function	dsp_design_biquad_peaking
Description	This function generates BiQuad filter coefficients for a peaking filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_peaking(0.25, 0.707, 3.0, coeffs[5], 28);</pre>
Type	<pre>void dsp_design_biquad_peaking(double filter_frequency, double filter_Q, double peak_gain_db, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter center frequency normalized to the sampling frequency. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. peak_gain_db The filter gain in dB (positive or negative). +gain results in peaking gain (gain at peak center = gain_db). -gain results in attenuation (gain at peak center = -gain_db). biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

79 Filter Design Functions: Bass Shelving Filter

Function	dsp_design_biquad_lowshelf
Description	This function generates BiQuad filter coefficients for a bass shelving filter. The filter coefficients are stored in forward order (e.g. b0,b1,b2,-a1,-a2). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_lowshelf(0.25, 0.707, +6.0, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_lowshelf(double filter_frequency, double filter_Q, double shelf_gain_db, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter frequency (+3db or -3db point) normalized to Fs. $0 < \text{frequency} < 0.5$, where 0.5 represents Fs/2. filter_Q The filter Q-factor. shelf_gain_db The filter shelf gain in dB (positive or negative). +gain results in bass shelf with gain of 'shelf_gain_db'. -gain results in bass shelf with attenuation of 'shelf_gain_db'. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as [b0,b1,b2,-a1,-a2]. q_format Fixed point format of coefficients (i.e. number of fractional bits).

80 Filter Design Functions: Treble Shelving Filter

Function	dsp_design_biquad_highshelf
Description	This function generates BiQuad filter coefficients for a treble shelving filter. The filter coefficients are stored in forward order (e.g. $b_0, b_1, b_2, -a_1, -a_2$). The frequency specification is normalized to the Nyquist frequency therefore the frequency value must be in the range of $0.0 \leq F < 0.5$ for valid filter coefficients. Example showing a filter coefficients generation using Q28 fixed-point formatting. <pre>int32_t coeffs[5]; dsp_design_biquad_highshelf(0.25, 0.707, +6.0, coeffs, 28);</pre>
Type	<pre>void dsp_design_biquad_highshelf(double filter_frequency, double filter_Q, double shelf_gain_db, int32_t biquad_coeffs[5], const int32_t q_format)</pre>
Parameters	filter_frequency Filter frequency (+3db or -3db point) normalized to F_s. $0 < \text{frequency} < 0.5$, where 0.5 represents $F_s/2$. filter_Q The filter Q-factor. shelf_gain_db The filter shelf gain in dB (positive or negative). +gain results in bass shelf with gain of 'shelf_gain_db'. -gain results in bass shelf with attenuation of 'shelf_gain_db'. biquad_coeffs The array used to contain the resulting filter coefficients. Filter coefficients are ordered as $[b_0, b_1, b_2, -a_1, -a_2]$. q_format Fixed point format of coefficients (i.e. number of fractional bits).

81 FFT functions

Note: The method for processing two real signals with a single complex FFT was improved. It now requires only half the memory. The function `dsp_fft_split_spectrum` is used to split the combined N point output of `dsp_fft_forward` into two half-spectra of size N/2. One for each of the two real input signals. `dsp_fft_merge_spectra` is used to merge the two half-spectra into a combined spectrum that can be processed by `dsp_fft_inverse`.

Function	<code>dsp_fft_split_spectrum</code>
Description	This function splits the spectrum of the FFT of two real sequences. Takes the result of a double-packed <code>dsp_complex_t</code> array that has undergone an FFT. This function splits the result into two arrays, one for each real sequence, of length N/2. It is expected that the output will be cast by: <code>dsp_complex_t (* restrict w)[2] = (dsp_complex_t (*)[2])pts</code> ; or a C equivalent. The 2 dimensional array <code>w[2][N/2]</code> can now be used to access the frequency information of the two real sequences independently, with the first index denoting the corresponding real sequence and the second index denoting the FFT frequency bin. Note that the DC component of the imaginary output spectrum (index zero) will contain the real component for the Nyquist rate. Note the minimum N is 8.
Type	void <code>dsp_fft_split_spectrum(dsp_complex_t pts[], const uint32_t N)</code>
Parameters	pts Array of <code>dsp_complex_t</code> elements. N Number of points. Must be a power of two.

Function	<code>dsp_fft_merge_spectra</code>
Description	This function merges two split spectra. It is the exact inverse operation of <code>dsp_fft_split_spectrum</code> . Note the minimum N is 8.
Type	void <code>dsp_fft_merge_spectra(dsp_complex_t pts[], const uint32_t N)</code>
Parameters	pts Array of <code>dsp_complex_t</code> elements. N Number of points. Must be a power of two.

Function	<code>dsp_fft_short_to_long</code>
Description	This function copies an array of <code>dsp_complex_short_t</code> elements to an array of an equal number of <code>dsp_complex_t</code> elements.

Continued on next page

Type	void dsp_fft_short_to_long(const dsp_complex_short_t s[], dsp_complex_t l[], const uint32_t N)	
Parameters	l	Array of dsp_complex_t elements.
	s	Array of dsp_complex_short_t elements.
	N	Number of points.

Function	dsp_fft_long_to_short	
Description	This function copies an array of dsp_complex_t elements to an array of an equal number of dsp_complex_short_t elements.	
Type	void dsp_fft_long_to_short(const dsp_complex_t l[], dsp_complex_short_t s[], const uint32_t N)	
Parameters	s	Array of dsp_complex_short_t elements.
	l	Array of dsp_complex_t elements.
	N	Number of points.

Function	dsp_fft_bit_reverse	
Description	This function preforms index bit reversing on the the arrays around prior to computing an FFT. A calling sequence for a forward FFT involves dsp_fft_bit_reverse() followed by dsp_fft_forward() , and for an inverse FFT it involves dsp_fft_bit_reverse() followed by dsp_fft_inverse() . In some cases bit reversal can be avoided, for example when computing a convolution.	
Type	void dsp_fft_bit_reverse(dsp_complex_t pts[], const uint32_t N)	
Parameters	pts	Array of dsp_complex_t elements.
	N	Number of points. Must be a power of two.

Function	dsp_fft_forward	
Description	This function computes a forward FFT. The complex input signal is supplied in an array of real and imaginary fixed-point values. The same array is also used to store the output. The magnitude of the FFT output is right shifted $\log_2(N)$ times which corresponds to division by N as shown in EQUATION 31-5 of http://www.dspguide.com/CH31.PDF. The number of points must be a power of 2, and the array of sine values should contain a quarter sine-wave. Use one of the dsp_sine_N tables. The function does not perform bit reversed indexing of the input data. If required then dsp_fft_bit_reverse() should be called beforehand.	
Type	<pre>void dsp_fft_forward(dsp_complex_t pts[], const uint32_t N, const int32_t sine[])</pre>	
Parameters	pts	Array of dsp_complex_t elements.
	N	Number of points. Must be a power of two.
	sine	Array of $N/4+1$ sine values, each represented as a sign bit, and a 31 bit fraction. 1 should be represented as 0x7fffffff. Arrays are provided in dsp_tables.c; for example, for a 1024 point FFT use dsp_sine_1024.

Function	dsp_fft_inverse	
Description	This function computes an inverse FFT. The complex input array is supplied as two arrays of integers, with numbers represented as fixed-point values. Max input range is -0x3fffffff..0x3fffffff. Integer overflow can occur with inputs outside of this range. The number of points must be a power of 2, and the array of sine values should contain a quarter sine-wave. Use one of the dsp_sine_N tables. The function does not perform bit reversed indexing of the input data. if required then dsp_fft_bit_reverse() should be called beforehand.	
Type	<pre>void dsp_fft_inverse(dsp_complex_t pts[], const uint32_t N, const int32_t sine[])</pre>	
Parameters	pts	Array of dsp_complex_t elements.
	N	Number of points. Must be a power of two.
	sine	Array of $N/4+1$ sine values, each represented as a sign bit, and a 31 bit fraction. 1 should be represented as 0x7fffffff. Arrays are provided in dsp_tables.c; for example, for a 1024 point FFT use dsp_sine_1024.

82 DCT functions

Function	dsp_dct_forward48
Description	This function performs a 48 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward48(int32_t output[48], int32_t input[48])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward32
Description	This function performs a 32 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward32(int32_t output[32], int32_t input[32])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward24
Description	This function performs a 24 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward24(int32_t output[24], int32_t input[24])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward16
Description	This function performs a 16 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward16(int32_t output[16], int32_t input[16])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward12
Description	This function performs a 12 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward12(int32_t output[12], int32_t input[12])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward8
Description	This function performs a 8 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward8(int32_t output[8], int32_t input[8])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward6
Description	This function performs a 6 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.

Continued on next page

Type	void dsp_dct_forward6(int32_t output[6], int32_t input[6])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward4
Description	This function performs a 4 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward4(int32_t output[4], int32_t input[4])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward3
Description	This function performs a 3 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward3(int32_t output[3], int32_t input[3])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward2
Description	This function performs a 2 point DCT. The first output is the DC value, subsequent values are the values for the basis vectors of half a cosine, a whole cosine, 1.5 cosine, 2 consines, etc.
Type	void dsp_dct_forward2(int32_t output[2], int32_t input[2])
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_forward1
Description	This function performs a single point DCT; it copies its input to its output.
Type	<code>void dsp_dct_forward1(int32_t output[1], int32_t input[1])</code>
Parameters	input input values to the DCT output DCT values.

Function	dsp_dct_inverse4
Description	This function performs a 4 point inverse DCT.
Type	<code>void dsp_dct_inverse4(int32_t output[4], int32_t input[4])</code>
Parameters	input Basis vector values output output signal

Function	dsp_dct_inverse3
Description	This function performs a 3 point inverse DCT.
Type	<code>void dsp_dct_inverse3(int32_t output[3], int32_t input[3])</code>
Parameters	input Basis vector values output output signal

Function	dsp_dct_inverse2
Description	This function performs a 2 point inverse DCT.
Type	<code>void dsp_dct_inverse2(int32_t output[2], int32_t input[2])</code>
Parameters	input Basis vector values output output signal

Function	dsp_dct_inverse1
Description	This function performs a 1 point inverse DCT.
Type	<code>void dsp_dct_inverse1(int32_t output[1], int32_t input[1])</code>
Parameters	input Basis vector values output output signal

APPENDIX A - Known Issues

There are no known issues.

APPENDIX B - DSP library change log

B.1 6.3.0

- ADDED: Support for XCommon CMake build system
- Changes to dependencies:
 - lib_logging: 3.0.0 -> 3.2.0

B.2 6.2.1

- CHANGED: Jenkinsfile used for CI

B.3 6.2.0

- CHANGED: Use XMOS Public Licence Version 1
- CHANGED: dsp_vector_abs function to always return positive values

B.4 6.1.0

- ADDED: Support for XS3 architecture

B.5 6.0.2

- CHANGED: use XS2 version of platform-specific functions on XS3
- CHANGED: Pin Python package versions
- REMOVED: not necessary cpanfile

B.6 6.0.1

- CHANGED: Updated generate_window.py to work with Python 3
- CHANGED: Updated to use Jenkins shared library 0.14.1

B.7 6.0.0

- ADDED: vector add, sub and mul functions for int8, int16 and int32 in both fixed and block floating point.
- CHANGED: Functions re-implemented in C, requires the use of unsafe pointers if called from XC
- Changes to dependencies:
 - lib_logging: Added dependency 3.0.0

B.8 5.0.0

- CHANGED: Build files updated to support new “xcommon” behaviour in xwaf.

B.9 4.2.0

- Added dsp_ch_pair_t struct to represent channel pairs
- Added floating point biquads
- Changed implementation of biquads to be faster and smaller
- Added use of pipenv to set up python environment

B.10 4.1.0

- Added post-FFT Hanning windowing
- Added function to combine real and imaginary arrays into complex array
- Added function to split complex array into real and imaginary arrays
- Added 48-point DCT
- Added dsp filter FIR add sample
- Added softplus
- Added integer sqrt
- Documentation updates

B.11 4.0.0

- Removed synchronous sample rate conversion functions - now maintained in lib_src
- Fixed bug in dsp_vector_mulv_addv()
- Faster bit reverse and inverse FFT
- Added real FFT
- Added real reverse FFT
- Added FFT with top half blanked
- Logistics functions
- Block floating point functions
- Complex FIR
- Complex vector arithmetic, with optional scaling
- Added in-place complex vector scaling
- Added complex vector magnitude
- Added complex vector scaling with arithmetic shift
- Added complex negative multiply and accumulate

B.12 3.1.0

- Deprecated synchronous sample rate conversion functions - now maintained in lib_src
- Added functions to compute a fast fixed point atan2 and hypotenuse
- Added Q8 versions of the arc sine and arc cosine functions
- Added 16384 point sine table
- Improved performance of the forwards FFT function, with small reduction in memory footprint

B.13 3.0.0

- Added exponential and natural logarithm functions
- Added Hyperbolic sine and cosine
- Fixed Matrix Multiplication and improved performance
- Changed API prefix from lib_dsp_ to dsp_.
- Changed lib_dsp_fft_complex_t to dsp_complex_t and lib_dsp_fft_complex_short_t to dsp_complex_short_t
- Various fixes in API documentation
- Added complex vector multiplication
- Added synchronous sample rate conversion (downsample or upsample by factor 3)

B.14 2.0.0

- FFT interface update. Consolidated interface and improved testing.

- Halved the memory for processing two real signals with a single complex FFT.
- Renamed *_transforms to *_fft to improve naming consistency
- Improved performance and accuracy of dsp_math_sqrt. Error is ≤ 1 . Worst case performance is 96 cycles.
- int32_t and uint32_t now used more consistently.

B.15 1.0.4

- Added fixed point sine and cosine functions. Performance: 62 cycles for dsp_math_sin, 64 cycles for dsp_math_cos.
- Brute force testing of all input values proved accuracy to within one LSB (error is ≤ 1)
- Added short int complex and tworeals FFT and iFFT
- Improved Macros for converting from double to int and int to double.
- Added optimised fixed point atan function dsp_math_atan
- Most tests in math_app.xc are now self-checking. Improved error reporting.
- Option for performance measurements in 10ns cycles.

B.16 1.0.3

- Update to source code license and copyright

B.17 1.0.2

- FFT and inverse FFT for two complex short int signals

B.18 1.0.1

- FFT and inverse FFT for complex signals or two real signals.

B.19 1.0.0

- Initial version