
**Information technology — Language
independent arithmetic —**

Part 3:

**Complex integer and floating point
arithmetic and complex elementary
numerical functions**

*Technologies de l'information — Arithmétique indépendante des
langages —*

*Partie 3: Arithmétique des nombres complexes entiers et en virgule
flottante et fonctions numériques élémentaires complexes*

PDF disclaimer

This PDF file may contain embedded typefaces. In accordance with Adobe's licensing policy, this file may be printed or viewed but shall not be edited unless the typefaces which are embedded are licensed to and installed on the computer performing the editing. In downloading this file, parties accept therein the responsibility of not infringing Adobe's licensing policy. The ISO Central Secretariat accepts no liability in this area.

Adobe is a trademark of Adobe Systems Incorporated.

Details of the software products used to create this PDF file can be found in the General Info relative to the file; the PDF-creation parameters were optimized for printing. Every care has been taken to ensure that the file is suitable for use by ISO member bodies. In the unlikely event that a problem relating to it is found, please inform the Central Secretariat at the address given below.

IECNORM.COM : Click to view the full PDF of ISO/IEC 10967-3:2006

© ISO/IEC 2006

All rights reserved. Unless otherwise specified, no part of this publication may be reproduced or utilized in any form or by any means, electronic or mechanical, including photocopying and microfilm, without permission in writing from either ISO at the address below or ISO's member body in the country of the requester.

ISO copyright office
Case postale 56 • CH-1211 Geneva 20
Tel. + 41 22 749 01 11
Fax + 41 22 749 09 47
E-mail copyright@iso.org
Web www.iso.org

Published in Switzerland

Contents

Foreword	vii
Introduction	viii
1 Scope	1
1.1 Inclusions	1
1.2 Exclusions	2
2 Conformity	3
3 Normative references	4
4 Symbols and definitions	4
4.1 Symbols	4
4.1.1 Sets and intervals	4
4.1.2 Operators and relations	4
4.1.3 Mathematical functions	5
4.1.4 Exceptional values	6
4.1.5 Datatypes and special values	6
4.1.6 Complex value constructors and complex datatype constructors	8
4.2 Definitions of terms	9
5 Specifications for imaginary and complex datatypes and operations	14
5.1 Imaginary and complex integer datatypes and operations	14
5.1.1 The complex integer <i>result</i> helper function	15
5.1.2 Imaginary and complex integer operations	15
5.1.2.1 Complex integer comparisons	15
5.1.2.2 Multiplication by the imaginary unit	17
5.1.2.3 The real and imaginary parts of a complex value	17
5.1.2.4 Formation of a complex integer from two real valued integers	18
5.1.2.5 Basic complex integer arithmetic	18
5.1.2.6 Absolute value and signum of integers and imaginary integers	21
5.1.2.7 Divisibility interrogation	21
5.1.2.8 Integer division and remainder extended to imaginary and complex integers	22
5.1.2.9 Maximum and minimum	27
5.2 Imaginary and complex floating point datatypes and operations	28
5.2.1 Maximum error requirements	28
5.2.2 Sign requirements	29
5.2.3 Monotonicity requirements	30
5.2.4 The complex floating point <i>result</i> helper functions	30
5.2.5 Basic arithmetic for complex floating point	31
5.2.5.1 Complex floating point comparisons	31
5.2.5.2 Multiplication by the imaginary unit	33
5.2.5.3 The real and imaginary parts of a complex value	34
5.2.5.4 Formation of a complex floating point from two floating point values	34
5.2.5.5 Fundamental complex floating point arithmetic	34

5.2.5.6	Absolute value, phase and signum of complex floating point values	38
5.2.5.7	Floor, round, and ceiling	39
5.2.5.8	Maximum and minimum	39
5.2.6	Complex sign, multiplication, and division	40
5.2.6.1	Complex signum	41
5.2.6.2	Complex multiplication	41
5.2.6.3	Complex division	42
5.2.7	Operations for conversion from polar to Cartesian	43
5.3	Elementary transcendental imaginary and complex floating point operations	44
5.3.1	Operations for exponentiations and logarithms	44
5.3.1.1	Exponentiation of imaginary base to integer power	44
5.3.1.2	Natural exponentiation	45
5.3.1.3	Complex exponentiation of argument base	45
5.3.1.4	Complex square root	46
5.3.1.5	Natural logarithm	47
5.3.2	Operations for radian trigonometric elementary functions	49
5.3.2.1	Radian angle normalisation	49
5.3.2.2	Radian sine	49
5.3.2.3	Radian cosine	50
5.3.2.4	Radian tangent	50
5.3.2.5	Radian cotangent	51
5.3.2.6	Radian secant	52
5.3.2.7	Radian cosecant	52
5.3.2.8	Radian arc sine	53
5.3.2.9	Radian arc cosine	54
5.3.2.10	Radian arc tangent	56
5.3.2.11	Radian arc cotangent	57
5.3.2.12	Radian arc secant	58
5.3.2.13	Radian arc cosecant	59
5.3.3	Operations for hyperbolic elementary functions	60
5.3.3.1	Hyperbolic normalisation	61
5.3.3.2	Hyperbolic sine	61
5.3.3.3	Hyperbolic cosine	61
5.3.3.4	Hyperbolic tangent	61
5.3.3.5	Hyperbolic cotangent	62
5.3.3.6	Hyperbolic secant	62
5.3.3.7	Hyperbolic cosecant	62
5.3.3.8	Inverse hyperbolic sine	62
5.3.3.9	Inverse hyperbolic cosine	63
5.3.3.10	Inverse hyperbolic tangent	63
5.3.3.11	Inverse hyperbolic cotangent	64
5.3.3.12	Inverse hyperbolic secant	64
5.3.3.13	Inverse hyperbolic cosecant	65
5.4	Operations for conversion between imaginary and complex numeric datatypes	65
5.4.1	Integer to complex integer conversions	65
5.4.2	Floating point to complex floating point conversions	66
5.5	Support for imaginary and complex numerals	67

6	Notification	68
6.1	Continuation values	68
7	Relationship with language standards	68
8	Documentation requirements	69
Annex A	(normative) Partial conformity	71
A.1	Maximum error relaxation	71
A.2	Extra accuracy requirements relaxation	71
A.3	Relationships to other operations relaxation	72
A.4	Part 1 and part 2 requirements relaxation	72
Annex B	(informative) Rationale	73
B.1	Scope	73
B.1.1	Inclusions	73
B.1.2	Exclusions	73
B.2	Conformity	73
B.3	Normative references	74
B.4	Symbols and definitions	74
B.4.1	Symbols	74
B.4.1.1	Sets and intervals	74
B.4.1.2	Operators and relations	74
B.4.1.3	Mathematical functions	74
B.4.1.4	Exceptional values	75
B.4.1.5	Datatypes and special values	75
B.4.1.6	Complex value constructors and complex datatype constructors	75
B.4.2	Definitions of terms	75
B.5	Specifications for the imaginary and complex datatypes and operations	76
B.5.1	Imaginary and complex integer datatypes and operations	76
B.5.2	Imaginary and complex floating point datatypes and operations	76
B.5.2.1	Maximum error requirements	76
B.5.2.2	Sign requirements	76
B.5.2.3	Maximum error requirements	76
B.5.2.4	Basic arithmetic for complex floating point	77
B.5.3	Elementary transcendental imaginary and complex floating point operations	78
B.5.3.1	Operations for exponentiations and logarithms	78
B.5.3.2	Operations for radian trigonometric elementary functions	78
B.5.3.2.1	Radian angle normalisation	78
B.5.3.2.2	Radian sine	79
B.5.3.2.3	Radian cosine	79
B.5.3.2.4	Radian tangent	79
B.5.3.2.5	Radian cotangent	80
B.5.3.2.6	Radian secant	80
B.5.3.2.7	Radian cosecant	80
B.5.3.2.8	Radian arc sine	80
B.5.3.2.9	Radian arc cosine	81
B.5.3.2.10	Radian arc tangent	81
B.5.3.2.11	Radian arc cotangent	81

B.5.3.2.12	Radian arc secant	81
B.5.3.2.13	Radian arc cosecant	81
B.5.3.3	Operations for hyperbolic elementary functions	82
B.5.3.3.1	Hyperbolic normalisation	82
B.5.3.3.2	Hyperbolic sine	82
B.5.3.3.3	Hyperbolic cosine	82
B.5.3.3.4	Hyperbolic tangent	82
B.5.3.3.5	Hyperbolic cotangent	83
B.5.3.3.6	Hyperbolic secant	83
B.5.3.3.7	Hyperbolic cosecant	83
B.5.3.3.8	Inverse hyperbolic sine	83
B.5.3.3.9	Inverse hyperbolic cosine	83
B.5.3.3.10	Inverse hyperbolic tangent	84
B.5.3.3.11	Inverse hyperbolic cotangent	84
B.5.3.3.12	Inverse hyperbolic secant	84
B.5.3.3.13	Inverse hyperbolic cosecant	84
B.5.4	Operations for conversion between imaginary and complex numeric datatypes	84
B.5.5	Support for imaginary and complex numerals	84
B.6	Notification	84
B.6.1	Continuation values	84
B.7	Relationship with language standards	84
B.8	Documentation requirements	84
Annex C	(informative) Example bindings for specific languages	85
C.1	Ada	86
C.2	C	96
C.3	C++	105
C.4	Fortran	113
C.5	Common Lisp	122
Annex D	(informative) Bibliography	133
Annex E	(informative) Cross reference	135
Annex F	(informative) Possible changes to part 2	147

Foreword

ISO (the International Organization for Standardization) and IEC (the International Electrotechnical Commission) form the specialized system for worldwide standardization. National bodies that are members of ISO or IEC participate in the development of International Standards through technical committees established by the respective organization to deal with particular fields of technical activity. ISO and IEC technical committees collaborate in fields of mutual interest. Other international organizations, governmental and non-governmental, in liaison with ISO and IEC, also take part in the work. In the field of information technology, ISO and IEC have established a joint technical committee, ISO/IEC JTC 1.

International Standards are drafted in accordance with the rules given in the ISO/IEC Directives, Part 2.

The main task of the joint technical committee is to prepare International Standards. Draft International Standards adopted by the joint technical committee are circulated to national bodies for voting. Publication as an International Standard requires approval by at least 75 % of the national bodies casting a vote.

Attention is drawn to the possibility that some of the elements of this document may be the subject of patent rights. ISO shall not be held responsible for identifying any or all such patent rights.

ISO/IEC 10967-3 was prepared by Joint Technical Committee ISO/IEC JTC 1, *Information technology*, Subcommittee SC 22, *Programming languages, their environments and system software interfaces*.

ISO/IEC 10967 consists of the following parts, under the general title *Information technology — Language independent arithmetic*:

- *Part 1: Integer and floating point arithmetic*
- *Part 2: Elementary numerical functions*
- *Part 3: Complex integer and floating point arithmetic and complex elementary numerical functions*

Introduction

The aims

Portability is a key issue for scientific and numerical software in today's heterogeneous computing environment. Such software may be required to run on systems ranging from personal computers to high performance pipelined vector processors and massively parallel systems, and the source code may be ported between several programming languages.

Part 1 of ISO/IEC 10967 specifies the basic properties of integer and floating point types that can be relied upon in writing portable software.

Part 2 of ISO/IEC 10967 specifies a number of additional operations for integer and floating point types, in particular specifications for numerical approximations to elementary functions on reals.

The content

The content of this document is based on part 1 and part 2, and extends part 1's and part 2's specifications to also cover operations approximating imaginary-integer and complex-integer arithmetic, imaginary-real and complex-real arithmetic, as well as imaginary-real and complex-real elementary functions.

The numerical functions covered by this document are computer approximations to mathematical functions of one or more imaginary or complex arguments. Accuracy versus performance requirements often vary with the application at hand. This is recognised by recommending that implementors support more than one library of these numerical functions. Various documentation and (program available) parameters requirements are specified to assist programmers in the selection of the library best suited to the application at hand.

The benefits

Adoption and proper use of this document can lead to the following benefits.

For programming language standards it will be possible to define their arithmetic semantics more precisely without preventing the efficient implementation of the language on a wide range of machine architectures.

Programmers of numeric software will be able to assess the portability of their programs in advance. Programmers will be able to trade off program design requirements for portability in the resulting program. Programs will be able to determine (at run time) the crucial numeric properties of the implementation. They will be able to reject unsuitable implementations, and (possibly) to correctly characterize the accuracy of their own results. Programs will be able to detect (and possibly correct for) exceptions in arithmetic processing.

Procurers of numerical programs will find it easier to determine whether a (properly documented) application program is likely to execute satisfactorily on the platform used. This can be done by comparing the documented requirements of the program against the documented properties of the platform.

Finally, end users of numeric application packages will be able to rely on the correct execution of those packages. That is, for correctly programmed algorithms, the results are reliable if and only if there is no notification.

Information technology — Language independent arithmetic —

Part 3: Complex integer and floating point arithmetic and complex elementary numerical functions

1 Scope

This part of ISO/IEC 10967 specifies the properties of numerical approximations for complex arithmetic operations and many of the complex elementary numerical functions available in a variety of programming languages in common use for mathematical and numerical applications.

An implementor may choose any combination of hardware and software support to meet the specifications of this document. It is the computing environment, as seen by the programmer/user, that does or does not conform to the specifications.

The term *implementation* (of this part) denotes the total computing environment pertinent to this part, including hardware, language processors, subroutine libraries, exception handling facilities, other software, and documentation.

1.1 Inclusions

The specifications of part 1 and part 2 are included by reference in this part.

This document provides specifications for properties of complex and imaginary integer datatypes and floating point datatypes, basic operations on values of these datatypes as well as for some numerical functions for which operand or result values are of imaginary or complex integer datatypes or imaginary or complex floating point datatypes constructed from integer and floating point datatypes satisfying the requirements of part 1 (ISO/IEC 10967-1). Boundaries for the occurrence of exceptions and the maximum error allowed are prescribed for each specified operation. Also the result produced by giving a special value operand, such as an infinity, or a NaN (not-a-number), is prescribed for each specified floating point operation.

This document provides specifications for:

- a) Basic imaginary integer and complex integer operations.
- b) Non-transcendental imaginary floating point and Cartesian complex floating point operations.
- c) Exponentiation, logarithm, radian trigonometric, and hyperbolic operations for imaginary floating point and Cartesian complex floating point.

This document also provides specifications for:

- d) The results produced by an included floating point operation when one or more operand values include IEC 60559 special values.

- e) Program-visible parameters that characterise certain aspects of the operations.

1.2 Exclusions

This document provides no specifications for:

- a) Datatypes and operations for polar complex floating point. This part neither requires nor excludes the presence of such polar complex datatypes and operations.
- b) Numerical functions whose operands are of more than one datatype, except certain imaginary/complex combinations. This part neither requires nor excludes the presence of such “mixed operand” operations.
- c) A complex interval datatype, or the operations on such datatypes. This part neither requires nor excludes such datatypes or operations.
- d) A complex fixed point datatype, or the operations on such datatypes. This part neither requires nor excludes such datatypes or operations.
- e) A complex rational datatype, or the operations on such datatypes. This part neither requires nor excludes such datatypes or operations.
- f) Matrix, statistical, or symbolic operations (on suitable datatypes). This part neither requires nor excludes such operations or datatypes.
- g) The properties of complex arithmetic datatypes that are not related to the numerical process, such as the representation of values on physical media.
- h) The properties of integer and floating point datatypes that properly belong in programming language standards or other specifications. Examples include:
 - 1) the syntax of numerals and expressions in the programming language,
 - 2) the syntax used for parsed (input) or generated (output) character string forms for numerals by any specific programming language or library,
 - 3) the precedence of operators in the programming language,
 - 4) the rules for assignment, parameter passing, and returning value,
 - 5) the presence or absence of automatic datatype coercions,
 - 6) the consequences of applying an operation to values of improper datatype, or to uninitialised data.

Furthermore, this part does not provide specifications for how the operations should be implemented or which algorithms are to be used for the various operations.

2 Conformity

It is expected that the provisions of this document will be incorporated by reference and further defined in other International Standards; specifically in programming language standards and in binding standards.

A binding standard specifies the correspondence between one or more of the arithmetic datatypes, parameters, and operations specified in this document and the concrete language syntax of some programming language. More generally, a binding standard specifies the correspondence between certain datatypes, parameters, and operations and the elements of some arbitrary computing entity. A language standard that explicitly provides such binding information can serve as a binding standard.

When a binding standard for a language exists, an implementation shall be said to conform to this document if and only if it conforms to the binding standard. In case of conflict between a binding standard and this document, the specifications of the binding standard take precedence.

When a binding standard requires only a subset of the imaginary or complex integer or imaginary or complex floating point datatypes specified in this document, an implementation remains free to conform to this document with respect to other datatypes independently of that binding standard.

When a binding standard requires only a subset of the operations specified in this document, an implementation remains free to conform to this document with respect to other operations, independently of that binding standard.

When no binding standard between a language and some datatypes or operations specified in this document exists, an implementation conforms to this document if and only if it provides one or more datatypes and one or more operations that together satisfy all the requirements of clauses 5 through 8 that are relevant to those datatypes and operations. The implementation shall then document the binding.

Conformity to this document is always with respect to a specified set of datatypes and set of operations. Conformity to this document implies conformity to part 1 and part 2 of ISO/IEC 10967 for the integer and floating point datatypes and operations used. Under certain circumstances, conformity to IEC 60559 is implied by conformity to part 1 of ISO/IEC 10967.

An implementation is free to provide arithmetic datatypes and arithmetic operations that do not conform to this document, or that are beyond the scope of this document. The implementation shall not claim or imply conformity to this document for such datatypes or operations.

An implementation is permitted to have modes of operation that do not conform to this document. A conforming implementation shall specify how to select the modes of operation that ensure conformity.

NOTES

- 1 Language bindings are essential. Clause 8 requires an implementation to document a binding if no binding standard exists. See Annex ?? for an example of a conformity statement, and Annex C for suggested language bindings.
- 2 A complete binding for this document will include (explicitly or by reference) a binding for part 2 and part 1 as well, which in turn may include (explicitly or by reference) a binding for IEC 60559 as well.

- 3 This document does not require a particular set of operations to be provided. It is not possible to conform to this document without specifying to which datatypes and set of operations (and modes of operation) conformity is claimed.

3 Normative references

The following referenced documents are indispensable for the application of this document. For dated references, only the edition cited applies. For undated references, the latest edition of the referenced document (including any amendments) applies.

IEC 60559:1989, *Binary floating-point arithmetic for microprocessor systems*.

ISO/IEC 10967-1, *Information technology – Language independent arithmetic – Part 1: Integer and floating point arithmetic*.

ISO/IEC 10967-2, *Information technology – Language independent arithmetic – Part 2: Elementary numerical functions*.

4 Symbols and definitions

4.1 Symbols

4.1.1 Sets and intervals

In this document, $\mathcal{Z}$ denotes the set of mathematical integers, $\mathcal{G}$ denotes the set of complex integers. $\mathcal{R}$ denotes the set of classical real numbers, and $\mathcal{C}$ denotes the set of complex numbers over $\mathcal{R}$. Note that $\mathcal{Z} \subset \mathcal{R} \subset \mathcal{C}$, and $\mathcal{Z} \subset \mathcal{G} \subset \mathcal{C}$.

The conventional notation for set definition and manipulation is used.

The following notation for intervals is used:

$[x, z]$ designates the interval $\{y \in \mathcal{R} \mid x \leq y \leq z\}$,
 $]x, z]$ designates the interval $\{y \in \mathcal{R} \mid x < y \leq z\}$,
 $[x, z[$ designates the interval $\{y \in \mathcal{R} \mid x \leq y < z\}$, and
 $]x, z[$ designates the interval $\{y \in \mathcal{R} \mid x < y < z\}$.

NOTE – The notation using a round bracket for an open end of an interval is not used, for the risk of confusion with the notation for pairs.

4.1.2 Operators and relations

All prefix and infix operators have their conventional (exact) mathematical meaning. The conventional notation for set definition and manipulation is also used. In particular:

$\Rightarrow$ and $\Leftrightarrow$ for logical implication and equivalence
 $+$, $-$, $/$, $|x|$, conj , $\lfloor x \rfloor$, $\lceil x \rceil$, and $\text{round}(x)$ on complex values
 $\cdot$ for multiplication on complex values
 $<$, $\leq$, $\geq$, and $>$ between real values
 $=$ and $\neq$ between real, complex, as well as special values

$\cup, \cap, \in, \notin, \subset, \subseteq, \not\subset, \neq$, and $=$ with sets
 $\times$ for the Cartesian product of sets
 $\rightarrow$ for a mapping between sets
 $|$ for the divides relation between complex integer values (in $\mathcal{G}$)
 $\bar{i}$ as the imaginary unit ($\bar{i}^2 = -1$)
 $\Re$ to extract the real part of a complex value (in $\mathcal{C}$)
 $\Im$ to extract the imaginary part of a complex value (in $\mathcal{C}$)
 NOTE 1 – $\approx$ is used informally, in notes and the rationale.

For $x \in \mathcal{C}$, the notation $\lfloor x \rfloor$ designates the component-wise largest complex integer where each part is not greater than the corresponding part of x :

$$\lfloor x \rfloor \in \mathcal{G} \text{ and } \Re(x) - 1 < \Re(\lfloor x \rfloor) \leq \Re(x) \text{ and } \Im(x) - 1 < \Im(\lfloor x \rfloor) \leq \Im(x)$$

the notation $\lceil x \rceil$ designates the component-wise smallest complex integer where each part is not less than the corresponding part of x :

$$\lceil x \rceil \in \mathcal{G} \text{ and } \Re(x) \leq \Re(\lceil x \rceil) < \Re(x) + 1 \text{ and } \Im(x) \leq \Im(\lceil x \rceil) < \Im(x) + 1$$

and the notation $\text{round}(x)$ designates the complex integer closest to x :

$$\text{round}(x) \in \mathcal{G} \text{ and}$$

$$\Re(x) - 0.5 \leq \Re(\text{round}(x)) \leq \Re(x) + 0.5 \text{ and } \Im(x) - 0.5 \leq \Im(\text{round}(x)) \leq \Im(x) + 0.5$$

where in case $\Re(x)$ or $\Im(x)$ is exactly half-way between two integers, the even integer is the result component.

The *divides* relation ($|$) on complex integers tests whether a complex integer i divides a complex integer j exactly:

$$i|j \Leftrightarrow (i \neq 0 \text{ and } i \cdot n = j \text{ for some } n \in \mathcal{G})$$

NOTE 2 – $i|j$ is true exactly when j/i is defined and $j/i \in \mathcal{G}$.

4.1.3 Mathematical functions

This document specifies properties for a number of operations numerically approximating some of the elementary functions. The following ideal mathematical functions are defined in chapter 4 of the *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables* [11] (e is the Napierian base):

$$e^x, x^y, \sqrt{x}, |x|, \ln, \log_b,$$

$$\sin, \cos, \tan, \cot, \sec, \csc, \arcsin, \arccos, \arctan, \text{arccot}, \text{arcsec}, \text{arccsc},$$

$$\sinh, \cosh, \tanh, \coth, \text{sech}, \text{csch}, \text{arcsinh}, \text{arccosh}, \text{arctanh}, \text{arccoth}, \text{arcsech}, \text{arccsch}.$$

Many of the inverses above are multi-valued. The selection of which value to return, the principal value, so as to make the inverses into functions, is done in the conventional way. E.g., $\sqrt{x} \in [0, \infty[$ when $x \in [0, \infty[$.

Part 2 defines the mathematical *arc* function, which is used in this part to define the mathematical complex *signum* function.

4.1.4 Exceptional values

ISO/IEC 10967 uses the following five exceptional values:

underflow: the result has an absolute value that is smaller than the smallest positive normalised value, and the result may be inexact. This notification need not be given if the result is exact. In particular, if the result is zero, or exact and IEC 60559 is conformed to and trapping is not enabled.

overflow: the result, after rounding, has a magnitude larger than the maximum value representable in the result datatype.

infinitary: the operation result is (exactly) infinite (in its real or its imaginary part), while all the arguments (their real and imaginary parts) are finite. This is to say that the corresponding mathematical function has a pole at the finite argument point.

invalid: the operation is undefined, while not infinitary, for the given arguments.

absolute_precision_underflow: indicate that the argument is such that the density of representable argument values is too small in the neighbourhood of the given argument value for a numeric result to be considered appropriate to return. Used for operations that approximate trigonometric functions (part 2 and part 3), and hyperbolic and exponentiation functions (part 3).

The exceptional value **inexact** is not specified in ISO/IEC 10967, but IEC 60559 conforming implementations will provide it. It should then be used also for operations approximating transcendental functions, when the returned result may be approximate. This part of ISO/IEC 10967 does not specify when it is appropriate to return this exceptional value, but does specify an appropriate continuation value (see Definition 4.2.4). Thus, v is specified by ISO/IEC 10967 when v or **inexact**(v) should be returned by implementations that are based on IEC 60559, and ISO/IEC 10967 specifies **underflow**(v) when **underflow**(v) or **inexact**(**underflow**(v)) should be returned by implementations that are based on IEC 60559. Ideally, **underflow** should imply **inexact**.

For the exceptional values, a continuation value may be given in this part in parenthesis after the exceptional value.

4.1.5 Datatypes and special values

The datatype **Boolean** consists of the two values **true** and **false**.

NOTE 1 – Mathematical relations are true or false (or undefined, if an operand is undefined), which are not values. In contrast, **true** and **false** are values in **Boolean**.

Square brackets are used to write finite sequences of values. $[]$ is the sequence containing no values. $[s]$ is the sequence of one value, s . $[s_1, s_2]$ is the sequence of two values, s_1 and then s_2 , etc. The colon operator is used to prepend a value to a sequence: $x : [x_1, \dots, x_n] = [x, x_1, \dots, x_n]$. $[S]$, where S is a set, denotes the set of finite sequences, where each value in a sequence is in S .

NOTE 2 – It is always clear from context, in the text of this part, if $[X]$ is a sequence of one element, or the set of sequences with values from X . It is also clear from context if $[x_1, x_2]$ is a sequence of two values or an interval.

Integer datatypes and floating point datatypes are defined in part 1. Let I be the non-special value set for an integer datatype conforming to part 1. Let F be the non-special value set for a floating point datatype conforming to part 1 and part 2. The following symbols used in this part are defined in part 1 or part 2:

Exceptional values:

underflow, **overflow**, **infinitary**, **invalid**, and **absolute_precision_underflow**.

Integer helper function:

$result_I$.

Integer operations:

eq_I , neq_I , lss_I , leq_I , gtr_I , geq_I , neg_I , add_I , sub_I , mul_I , abs_I , max_I , min_I , max_seq_I , min_seq_I , $divides_I$, $quot_I$, mod_I , $ratio_I$, $residue_I$, $group_I$, and pad_I .

Integer conversion operation:

$convert_{I \rightarrow I'}$.

Floating point parameters:

r_F , p_F , $emin_F$, $emax_F$, $denorm_F$, and iec_559_F .

Derived floating point constants:

$fmax_F$, $fmin_F$, $fminN_F$, $fminD_F$, and $epsilon_F$.

Floating point value sets related to F :

F^* , F_D , F_N , G_F , $F^{2\cdot\pi}$, and F^u (for a given u).

Floating point helper functions:

up_F , $down_F$, $nearest_F$, $result_F$, $result_F^*$, no_result_F , and $no_result2_F$.

Floating point operations from part 1:

eq_F , neq_F , lss_F , leq_F , gtr_F , geq_F , neg_F , add_F , sub_F , mul_F , div_F , abs_F , max_F , $mmax_F$, min_F , $mmin_F$, max_seq_F , $mmax_seq_F$, min_seq_F , $mmin_seq_F$, $floor_F$, $ceiling_F$, and $rounding_F$.

Floating point operations from part 2:

$sqrt_F$, $hypot_F$, $power_{F,I}$, exp_F , $power_F$, ln_F , $logbase_F$, rad_F , sin_F , cos_F , tan_F , cot_F , sec_F , csc_F , $arcsin_F$, $arccos_F$, $arctan_F$, $arccot_F$, $arcsec_F$, $arccsc_F$, arc_F , $arcu_F$, $sinu_F$, $cosu_F$, $sinh_F$, $cosh_F$, $tanh_F$, $coth_F$, $sech_F$, $csch_F$, $arcsinh_F$, $arccosh_F$, $arctanh_F$, $arcoth_F$, $arcsech_F$, and $arccsch_F$.

Angular parameters and maximum error parameters from part 2:

$big_angle_r_F$, $big_angle_u_F$, $max_error_tan_F$, and $max_error_tanu_F$.

Floating point conversion operations:

$convert_{F \rightarrow F'}$, $convert_{F \rightarrow D}$, and $convert_{D \rightarrow F}$.

Approximation helper functions from part 2:

exp_F^* , $power_F^*$, ln_F^* , $logbase_F^*$, sin_F^* , cos_F^* , tan_F^* , cot_F^* , sec_F^* , csc_F^* , $arcsin_F^*$, $arccos_F^*$, $arctan_F^*$, $arccot_F^*$, $arcsec_F^*$, $arccsc_F^*$, $sinh_F^*$, $cosh_F^*$, $tanh_F^*$, $coth_F^*$, $sech_F^*$, $csch_F^*$, $arcsinh_F^*$, $arccosh_F^*$, $arctanh_F^*$, $arcoth_F^*$, $arcsech_F^*$, and $arccsch_F^*$.

Floating point datatypes that conform to part 1 shall, for use with this part, like for part 2, have a value for the parameter p_F such that $p_F \geq 2 \cdot \max\{1, \log_{r_F}(2 \cdot \pi)\}$, and have a value for the parameter $emin_F$ such that $emin_F \leq -p_F - 1$.

NOTES

- 3 This implies that $fminN_F < 0.5 \cdot epsilon_F / r_F$ in this part, rather than just $fminN_F \leq epsilon_F$.

- 4 These extra requirements, which do not limit the use of any existing floating point datatype, are made so that angles in radians are not too degenerate within the first two cycles, plus and minus, when represented in F .
- 5 F should also be such that $p_F \geq 2 + \log_{r_F}(1000)$, to allow for a not too coarse angle resolution anywhere in the interval $[-big_angle_r_F, big_angle_r_F]$ with the default value for $big_angle_r_F$. See Clause 5.3.9 of part 2.

The following symbols represent special values defined in IEC 60559 and are used in this part:

-0, **+∞**, **-∞**, **qNaN**, and **sNaN**.

These floating point values are not part of the set F , but if iec_559_F has the value **true**, these values are included in the floating point datatype corresponding to F .

NOTE 6 – This part uses the above five special values for compatibility with IEC 60559. In particular, the symbol **-0** (in bold) is not the application of (mathematical) unary $-$ to the value 0, and is a value logically distinct from 0.

The specifications cover the results to be returned by an operation if given one or more of the IEC 60559 special values **-0**, **+∞**, **-∞**, or NaNs as input values. These specifications apply only to systems which provide and support these special values. If an implementation is not capable of representing a **-0** result or continuation value, 0 shall be used as the actual result or continuation value. If an implementation is not capable of representing a prescribed result or continuation value of the IEC 60559 special values **+∞**, **-∞**, or **qNaN**, the actual result or continuation value is binding or implementation defined.

If and only if an implementation is not capable of representing **-0**:

- a) a 0 as the imaginary part of a complex argument (in $c(F)$, see Clause 4.1.6) shall be interpreted as if it was **-0** if and only if the real part of that complex argument is greater than or equal to zero, and
- b) a 0 as the real part of a complex argument (in $c(F)$, see Clause 4.1.6) shall be interpreted as if it was **-0** if and only if the imaginary part of the complex argument is less than zero.

NOTES

- 7 Reinterpreting 0 as **-0** as required above is needed to follow the sign rules for inverse trigonometric and inverse hyperbolic operations, as well as the exact relations between trigonometric and hyperbolic operations also for argument parts (real and imaginary) that have a zero as value.
- 8 The rule above is sometimes referred to as *continuous when approaching an axis in a counterclockwise path*. This fits both with Common Lisp and C requirements when zeroes don't have a distinguishable sign.
- 9 For consistency, this rule also has implications for the operations that implicitly or explicitly extract an implicit real or implicit imaginary part (see for example the specifications for the $re_{i(F)}$ and im_F operations in Clause 5.2.5).

4.1.6 Complex value constructors and complex datatype constructors

Let X be a set containing values in $\mathcal{R}$, and possibly also containing special values (such as IEC 60559 special values).

$i(X)$ is a subset of values in an imaginary datatype, constructed from the datatype corresponding to X . $\hat{i}\cdot$ is a prefix constructor that takes one parameter.

$$i(X) = \{\hat{i}\cdot y \mid y \in X\}$$

$c(X)$ is a subset of values in a complex datatype, constructed from the datatype corresponding to X . $\mathbf{+i\cdot}$ is an infix constructor that takes two parameters.

$$c(X) = \{x \mathbf{+i\cdot} y \mid x, y \in X\}$$

NOTES

- 1 While $\mathbf{i\cdot}$ and $\mathbf{+i\cdot}$ (note that they are written in bold) have an appearance of being the imaginary unit together with the plus and times operators, that is not the case. For instance, $\mathbf{i\cdot} 2$ is an element of $i(X)$ (if $2 \in X$), but not of $\mathcal{G}$ or $\mathcal{C}$. $\tilde{\mathbf{i\cdot}} 2$, on the other hand, is an expression that denotes an element of $\mathcal{G}$ (and $\mathcal{C}$), but neither of $i(X)$ nor $c(X)$. Further, e.g., $4 + \tilde{\mathbf{i\cdot}} 0 = 4$, but $4 \mathbf{+i\cdot} 0 \neq 4$.
- 2 A constructor that takes one argument is a one-tuple tag. A constructor that takes two arguments is a two-tuple (pair) tag. The arguments are part of the resulting value.
- 3 The tuple tags need not be explicitly represented in implementations. But if represented, there should be different tags for different argument types (which is not needed in this text).

Some of the helper function signatures use $\mathcal{C}_F$, where

$$\mathcal{C}_F = \{x + \tilde{\mathbf{i\cdot}} y \mid x, y \in F\}$$

where $F \subset \mathcal{R}$.

4.2 Definitions of terms

For the purposes of this document, the following terms and definitions apply.

4.2.1

accuracy

The closeness between the true mathematical result and a computed result.

4.2.2

arithmetic datatype

A datatype whose non-special values are members of $\mathcal{Z}$, $i(\mathcal{Z})$, $c(\mathcal{Z})$, $\mathcal{R}$, $i(\mathcal{R})$, or $c(\mathcal{R})$.

NOTE 1 – $i(\mathcal{Z})$ corresponds to imaginary integer values in $\mathcal{G}$. $c(\mathcal{Z})$ corresponds to complex integer values in $\mathcal{G}$. $i(\mathcal{R})$ corresponds to imaginary values in $\mathcal{C}$. $c(\mathcal{R})$ corresponds to complex values in $\mathcal{C}$.

4.2.3

binding

Documentation that specifies which syntax (most often, operations or functions) of a programming language, or programming library, that is used for an operation specified in this part.

4.2.4

continuation value

A computational value used as the result of an arithmetic operation when an exception occurs. Continuation values are intended to be used in subsequent arithmetic processing. A continuation value can be a (in the datatype representable) value in $\mathcal{Z}$, $i(\mathcal{Z})$, $c(\mathcal{Z})$, $\mathcal{R}$, $i(\mathcal{R})$, $c(\mathcal{R})$, or a value where one or both parts of the value is an IEC 60559 special value (e.g. $-\mathbf{0}$, $\mathbf{i\cdot}(-\infty)$), or $+\infty + \mathbf{i\cdot} \mathbf{qNaN}$). (Contrast with *exceptional value*. See Clause 6.1.2 of part 1.)

4.2.5**denormalisation loss**

A larger than normal rounding error caused by the fact that subnormal values have less than full precision. (See Clause 5.2.5 of part 1 for a full definition.)

4.2.6**error**

⟨in computed value⟩ The difference between a computed value and the mathematically correct value. Used in phrases like “rounding error” or “error bound”.

4.2.7**error**

⟨computation gone awry⟩ A synonym for *exception* in phrases like “error message” or “error output”. Error and exception are not synonyms in any other contexts.

4.2.8**exception**

The inability of an operation to return a suitable finite numeric result from finite arguments. This might arise because no such finite result exists mathematically, or because the mathematical result cannot be represented with sufficient accuracy.

4.2.9**exceptional value**

A non-numeric value produced by an arithmetic operation to indicate the occurrence of an exception. Exceptional values are not used in subsequent arithmetic processing. (See Clause 5 of part 1.)

NOTES

- 2 Exceptional values are used as part of the defining formalism only. With respect to this part, they do not represent values of any of the datatypes described. There is no requirement that they be represented or stored in the computing system.
- 3 Exceptional values are not to be confused with the NaNs and infinities defined in IEC 60559. Contrast this definition with that of *continuation value* above.

4.2.10**helper function**

A function used solely to aid in the expression of a requirement. Helper functions are not visible to the programmer, and are not required to be part of an implementation.

4.2.11**implementation** (of this part)

The total arithmetic environment presented to a programmer, including hardware, language processors, exception handling facilities, subroutine libraries, other software, and all pertinent documentation.

4.2.12**literal**

A syntactic entity denoting a constant value without having proper sub-entities that are expressions.

4.2.13**monotonic approximation**

An approximation helper function $h : \dots \times S \times \dots \rightarrow \mathcal{R}$, where the other arguments are kept constant, and where $S \subseteq \mathcal{R}$, is a monotonic approximation of a predetermined mathematical function $f : \mathcal{R} \rightarrow \mathcal{R}$ if, for every $a \in S$ and $b \in S$, where $a < b$,

- a) f is monotonic non-decreasing on $[a, b]$ implies $h(\dots, a, \dots) \leq h(\dots, b, \dots)$,
- b) f is monotonic non-increasing on $[a, b]$ implies $h(\dots, a, \dots) \geq h(\dots, b, \dots)$.

4.2.14**monotonic non-decreasing**

A function $f : \mathcal{R} \rightarrow \mathcal{R}$ is monotonic non-decreasing on a real interval $[a, b]$ if for every x and y such that $a \leq x \leq y \leq b$, $f(x)$ and $f(y)$ are well-defined and $f(x) \leq f(y)$.

4.2.15**monotonic non-increasing**

A function $f : \mathcal{R} \rightarrow \mathcal{R}$ is monotonic non-increasing on a real interval $[a, b]$ if for every x and y such that $a \leq x \leq y \leq b$, $f(x)$ and $f(y)$ are well-defined and $f(x) \geq f(y)$.

4.2.16**normalised**

The non-zero values of a floating point type F that provide the full precision allowed by that type. (See F_N in Clause 5.2 of part 1 for a full definition.)

4.2.17**notification**

The process by which a program (or that program's end user) is informed that an arithmetic exception has occurred. For example, dividing 2 by 0 results in a notification. (See Clause 6 of part 1)

4.2.18**numeral**

A numeric literal. It may denote a value in $\mathcal{Z}$, $i(\mathcal{Z})$, $c(\mathcal{Z})$, $\mathcal{R}$, $i(\mathcal{R})$, or $c(\mathcal{R})$, a value which is -0 , an infinity, or a NaN.

4.2.19**numerical function**

A computer routine or other mechanism for the approximate evaluation of a mathematical function.

4.2.20**operation**

A function directly available to the programmer, as opposed to helper functions or theoretical mathematical functions.

4.2.21**pole**

A mathematical function f has a pole at x_0 if x_0 is finite, f is defined, finite, monotone, and continuous in at least part of the neighbourhood of x_0 , and the imaginary or real part of $\lim_{x \rightarrow x_0} f(x)$ is infinite via at least one path to x_0 .

4.2.22**precision**

The number of digits in the fraction of a floating point number. (See Clause 5.2 of part 1.)

4.2.23**rounding**

The act of computing a representable final result for an operation that is close to the exact (but unrepresentable in the result datatype) result for that operation. Note that a suitable representable result may not exist (see Clause 5.2.6 of part 1). (See also Annex A.5.2.6 of part 1 for some examples.)

4.2.24**rounding function**

Any function $rnd : \mathcal{R} \rightarrow X$ (where X is a given discrete and unlimited subset of $\mathcal{R}$) that maps each element of X to itself, and is monotonic non-decreasing. Formally, if x and y are in $\mathcal{R}$,

$$\begin{aligned} x \in X &\Rightarrow rnd(x) = x \\ x < y &\Rightarrow rnd(x) \leq rnd(y) \end{aligned}$$

Note that if $u \in \mathcal{R}$ is between two adjacent values in X , $rnd(u)$ selects one of those adjacent values.

4.2.25**round to nearest**

The property of a rounding function rnd that when $u \in \mathcal{R}$ is between two adjacent values in X , $rnd(u)$ selects the one nearest u . If the adjacent values are equidistant from u , either may be chosen deterministically.

4.2.26**round toward minus infinity**

The property of a rounding function rnd that when $u \in \mathcal{R}$ is between two adjacent values in X , $rnd(u)$ selects the one less than u .

4.2.27**round toward plus infinity**

The property of a rounding function rnd that when $u \in \mathcal{R}$ is between two adjacent values in X , $rnd(u)$ selects the one greater than u .

4.2.28**shall**

A verbal form used to indicate requirements strictly to be followed in order to conform to the standard and from which no deviation is permitted. (Quoted from the directives [1].)

4.2.29**should**

A verbal form used to indicate that among several possibilities one is recommended as particularly suitable, without mentioning or excluding others; or that (in the negative form) a certain possibility is deprecated but not prohibited. (Quoted from the directives [1].)

4.2.30**signature** (of a function or operation)

A summary of information about an operation or function. A signature includes the function or operation name; a subset of allowed argument values to the function or operation; and a superset of results from the function or operation (including exceptional values if any), if the argument is in the subset of argument values given in the signature.

The signature

$$add_I : I \times I \rightarrow I \cup \{\mathbf{overflow}\}$$

states that the operation named add_I shall accept any pair of I values as input, and (when given such input) shall return either a single I value as its output or the exceptional value **overflow**.

A signature for an operation or function does not forbid the operation from accepting a wider range of arguments, nor does it guarantee that every value in the result range will actually be returned for some input. An operation given an argument outside the stipulated argument domain may produce a result outside the stipulated result range.

4.2.31**subnormal**

denormal (obsolete)

The values of a floating point datatype F , as well as $-\mathbf{0}$, whose absolute values are strictly less than the smallest positive normal value in F .

NOTE 4 – Compare the definition of F_D in Clause 5.2 of part 1. In the first edition (1994) of part 1 and in IEC 60559:1989 this concept, then excepting the zero values, was called denormal.

4.2.32**ulp**

The value of one “unit in the last place” of a floating point number. This value depends on the exponent, the radix, and the precision used in representing the number. Thus, the ulp of a

normalised value x (in F), with exponent t , precision p , and radix r , is r^{t-p} , and the ulp of a subnormal value is $fminD_F$. (See Clause 5.2 of part 1.)

5 Specifications for imaginary and complex datatypes and operations

This clause specifies imaginary and complex integer datatypes, imaginary and complex floating point datatypes and a number of helper functions and operations for imaginary and complex integer as well as imaginary and complex floating point datatypes.

Each operation is given a signature and is further specified by a number of cases. These cases may refer to other operations (specified in this document, in part 1, or in part 2), to mathematical functions (textbook elementary functions, or functions defined in ISO/IEC 10967), and to helper functions (specified in this document, in part 1, or in part 2). They also use special abstract values ($-\infty$, $+\infty$, -0 , **qNaN**, **sNaN**). For each datatype, two of these abstract values, **qNaN** and **sNaN**, may represent several actual values each. Finally, the specifications may refer to exceptional values.

The signatures in the specifications in this clause specify only all non-special values as input values, and indicate as output values a superset of all non-special, special, and exceptional values that may result from these (non-special) input values. Exceptional and special values that can never result from non-special input values are not included in the signatures given. Also, signatures that, for example, include IEC 60559 special values as arguments are not given in the specifications below. This does not exclude such signatures from being valid for these operations.

NOTE – For instance, the *realpart* operation on complex floating point is given with the following signature:

$$re_{c(F)} : c(F) \rightarrow F$$

But the following signature is also valid, and takes some special values into account

$$re_{c(F)} : c(F \cup \{-\infty, -0, +\infty\}) \rightarrow F \cup \{-\infty, -0, +\infty\}$$

The following signature is also valid

$$re_{c(F)} : c(F \cup \{-\infty, -0, +\infty, \mathbf{qNaN}, \mathbf{sNaN}\}) \rightarrow F \cup \{-\infty, -0, +\infty, \mathbf{qNaN}, \mathbf{invalid}\}$$

5.1 Imaginary and complex integer datatypes and operations

Clause 5.1 of part 1 and clause 5.1 of part 2 specify integer datatypes and a number of operations on values of an integer datatype. In this clause imaginary and complex integer datatypes and operations on values of an imaginary or complex integer datatype are specified.

A complex integer datatype is constructed from an integer datatype. There should be at least one imaginary integer datatype and at least one complex integer datatype for each provided integer datatype.

I is the set of non-special values, $I \subseteq \mathcal{Z}$, for an integer datatype conforming to part 1. Integer datatypes conforming to part 1 often do not contain any NaN or infinity values, even though they may do so. Therefore this clause has no specifications for such values as arguments or results.

$i(I)$ (see Clause 4.1.6) is the set of non-special values in an imaginary integer datatype, constructed from the integer datatype corresponding to non-special value set I .

$c(I)$ (see Clause 4.1.6) is the set of non-special values in a complex integer datatype, constructed from the integer datatype corresponding to the non-special value set I .

NOTE – The operations that return zero for certain cases, according to the specifications below, may in a subset of those cases return negative zero instead, if negative zero can be represented. Compare the specifications for corresponding complex floating point operations in Clause 5.2.

5.1.1 The complex integer *result* helper function

The $result_{c(I)}$ helper function:

$$result_{c(I)} : \mathcal{G} \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$result_{c(I)}(z) = result_I(\Re(z)) + \mathbf{i} \cdot result_I(\Im(z))$$

NOTE – If one or both of the $result_I$ (defined in part 2) function applications on the right side returns **overflow**, then the $result_{c(I)}$ application returns **overflow**. The continuation values used when **overflow** occurs are to be specified by the binding or implementation. The same holds also for other exceptional values and also for the specifications below that do not use $result_{c(I)}$ but specify the result parts directly.

5.1.2 Imaginary and complex integer operations

5.1.2.1 Complex integer comparisons

$$eq_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$eq_{i(I)}(\mathbf{i} \cdot y, \mathbf{i} \cdot w) = eq_I(y, w)$$

$$eq_{I,i(I)} : I \times i(I) \rightarrow \mathbf{Boolean}$$

$$eq_{I,i(I)}(x, \mathbf{i} \cdot w) = eq_{c(I)}(x + \mathbf{i} \cdot 0, 0 + \mathbf{i} \cdot w)$$

$$eq_{i(I),I} : i(I) \times I \rightarrow \mathbf{Boolean}$$

$$eq_{i(I),I}(\mathbf{i} \cdot y, z) = eq_{c(I)}(0 + \mathbf{i} \cdot y, z + \mathbf{i} \cdot 0)$$

$$eq_{I,c(I)} : I \times c(I) \rightarrow \mathbf{Boolean}$$

$$eq_{I,c(I)}(x, z + \mathbf{i} \cdot w) = eq_{c(I)}(x + \mathbf{i} \cdot 0, z + \mathbf{i} \cdot w)$$

$$eq_{c(I),I} : c(I) \times I \rightarrow \mathbf{Boolean}$$

$$eq_{c(I),I}(x + \mathbf{i} \cdot y, z) = eq_{c(I)}(x + \mathbf{i} \cdot y, z + \mathbf{i} \cdot 0)$$

$$eq_{i(I),c(I)} : i(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$eq_{i(I),c(I)}(\mathbf{i} \cdot y, z + \mathbf{i} \cdot w) = eq_{c(I)}(0 + \mathbf{i} \cdot y, z + \mathbf{i} \cdot w)$$

$$eq_{c(I),i(I)} : c(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$eq_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = eq_{c(I)}(x + \hat{\mathbf{i}} \cdot y, 0 + \hat{\mathbf{i}} \cdot w)$$

$$eq_{c(I)} : c(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$eq_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = \mathbf{true} \quad \text{if } eq_I(x, z) = \mathbf{true} \text{ and } eq_I(y, w) = \mathbf{true} \\ = \mathbf{false} \quad \text{if } eq_I(x, z) = \mathbf{false} \text{ and } eq_I(y, w) = \mathbf{true} \\ = \mathbf{false} \quad \text{if } eq_I(x, z) = \mathbf{true} \text{ and } eq_I(y, w) = \mathbf{false} \\ = \mathbf{false} \quad \text{if } eq_I(x, z) = \mathbf{false} \text{ and } eq_I(y, w) = \mathbf{false}$$

$$neq_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$neq_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = neq_I(y, w)$$

$$neq_{I,i(I)} : I \times i(I) \rightarrow \mathbf{Boolean}$$

$$neq_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) = neq_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, 0 + \hat{\mathbf{i}} \cdot w)$$

$$neq_{i(I),I} : i(I) \times I \rightarrow \mathbf{Boolean}$$

$$neq_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = neq_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0)$$

$$neq_{I,c(I)} : I \times c(I) \rightarrow \mathbf{Boolean}$$

$$neq_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = neq_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w)$$

$$neq_{c(I),I} : c(I) \times I \rightarrow \mathbf{Boolean}$$

$$neq_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\ = neq_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0)$$

$$neq_{i(I),c(I)} : i(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$neq_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = neq_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w)$$

$$neq_{c(I),i(I)} : c(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$neq_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = neq_{c(I)}(x + \hat{\mathbf{i}} \cdot y, 0 + \hat{\mathbf{i}} \cdot w)$$

$$neq_{c(I)} : c(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$neq_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = \mathbf{true} \quad \text{if } neq_I(x, z) = \mathbf{true} \text{ and } neq_I(y, w) = \mathbf{true} \\ = \mathbf{true} \quad \text{if } neq_I(x, z) = \mathbf{false} \text{ and } neq_I(y, w) = \mathbf{true} \\ = \mathbf{true} \quad \text{if } neq_I(x, z) = \mathbf{true} \text{ and } neq_I(y, w) = \mathbf{false} \\ = \mathbf{false} \quad \text{if } neq_I(x, z) = \mathbf{false} \text{ and } neq_I(y, w) = \mathbf{false}$$

$$lss_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$lss_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = lss_I(y, w)$$

$$leq_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$leq_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = leq_I(y, w)$$

$$gtr_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$gtr_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = gtr_I(y, w)$$

$$geq_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$geq_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = geq_I(y, w)$$

5.1.2.2 Multiplication by the imaginary unit

$$itimes_{I \rightarrow i(I)} : I \rightarrow i(I)$$

$$itimes_{I \rightarrow i(I)}(x) = \hat{\mathbf{i}} \cdot x$$

$$itimes_{i(I) \rightarrow I} : i(I) \rightarrow I \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} itimes_{i(I) \rightarrow I}(\hat{\mathbf{i}} \cdot y) \\ = neg_I(y) \end{aligned}$$

$$itimes_{c(I)} : c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} itimes_{c(I)}(x \mathbf{+} \hat{\mathbf{i}} \cdot y) \\ = neg_I(y) \mathbf{+} \hat{\mathbf{i}} \cdot x \end{aligned}$$

5.1.2.3 The real and imaginary parts of a complex value

$$re_I : I \rightarrow I$$

$$re_I(x) = x \quad \text{if } x \in I$$

$$re_{i(I)} : i(I) \rightarrow \{0\}$$

$$re_{i(I)}(\hat{\mathbf{i}} \cdot y) = 0 \quad \text{if } y \in I$$

$$re_{c(I)} : c(I) \rightarrow I$$

$$re_{c(I)}(x \mathbf{+} \hat{\mathbf{i}} \cdot y) = x \quad \text{if } x \in I$$

$$im_I : I \rightarrow \{0\}$$

$$im_I(x) = 0 \quad \text{if } x \in I$$

$$\begin{aligned}
im_{i(I)} : i(I) &\rightarrow I \\
im_{i(I)}(\hat{\mathbf{i}} \cdot y) &= y && \text{if } y \in I
\end{aligned}$$

$$\begin{aligned}
im_{c(I)} : c(I) &\rightarrow I \\
im_{c(I)}(x + \hat{\mathbf{i}} \cdot y) &= y && \text{if } y \in I
\end{aligned}$$

5.1.2.4 Formation of a complex integer from two real valued integers

$$\begin{aligned}
plusitimes_I : I \times I &\rightarrow c(I) \\
plusitimes_I(x, z) &= x + \hat{\mathbf{i}} \cdot z
\end{aligned}$$

5.1.2.5 Basic complex integer arithmetic

$$\begin{aligned}
neg_{i(I)} : i(I) &\rightarrow i(I) \cup \{\mathbf{overflow}\} \\
neg_{i(I)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot neg_I(y)
\end{aligned}$$

$$\begin{aligned}
neg_{c(I)} : c(I) &\rightarrow c(I) \cup \{\mathbf{overflow}\} \\
neg_{c(I)}(x + \hat{\mathbf{i}} \cdot y) &= neg_I(x) + \hat{\mathbf{i}} \cdot neg_I(y)
\end{aligned}$$

$$\begin{aligned}
conj_I : I &\rightarrow I \\
conj_I(x) &= x
\end{aligned}$$

$$\begin{aligned}
conj_{i(I)} : i(I) &\rightarrow i(I) \cup \{\mathbf{overflow}\} \\
conj_{i(I)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot neg_I(y)
\end{aligned}$$

$$\begin{aligned}
conj_{c(I)} : c(I) &\rightarrow c(I) \cup \{\mathbf{overflow}\} \\
conj_{c(I)}(x + \hat{\mathbf{i}} \cdot y) &= x + \hat{\mathbf{i}} \cdot neg_I(y)
\end{aligned}$$

$$\begin{aligned}
add_{i(I)} : i(I) \times i(I) &\rightarrow i(I) \cup \{\mathbf{overflow}\} \\
add_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) &= \hat{\mathbf{i}} \cdot add_I(y, w)
\end{aligned}$$

$$\begin{aligned}
add_{I, i(I)} : I \times i(I) &\rightarrow c(I) \\
add_{I, i(I)}(x, \hat{\mathbf{i}} \cdot w) &= x + \hat{\mathbf{i}} \cdot w
\end{aligned}$$

$$\begin{aligned}
add_{i(I), I} : i(I) \times I &\rightarrow c(I) \\
add_{i(I), I}(\hat{\mathbf{i}} \cdot y, z) &= z + \hat{\mathbf{i}} \cdot y
\end{aligned}$$

$$\begin{aligned}
& add_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& add_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\
& \quad = add_I(x, z) + \hat{\mathbf{i}} \cdot w
\end{aligned}$$

$$\begin{aligned}
& add_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& add_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\
& \quad = add_I(x, z) + \hat{\mathbf{i}} \cdot y
\end{aligned}$$

$$\begin{aligned}
& add_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& add_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
& \quad = z + \hat{\mathbf{i}} \cdot add_I(y, w)
\end{aligned}$$

$$\begin{aligned}
& add_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& add_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
& \quad = x + \hat{\mathbf{i}} \cdot add_I(y, w)
\end{aligned}$$

$$\begin{aligned}
& add_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& add_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
& \quad = add_I(x, z) + \hat{\mathbf{i}} \cdot add_I(y, w)
\end{aligned}$$

$$\begin{aligned}
& sub_{i(I)} : i(I) \times i(I) \rightarrow i(I) \cup \{\mathbf{overflow}\} \\
& sub_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = \hat{\mathbf{i}} \cdot sub_I(y, w)
\end{aligned}$$

$$\begin{aligned}
& sub_{I,i(I)} : I \times i(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& sub_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) = x + \hat{\mathbf{i}} \cdot neg_I(w)
\end{aligned}$$

$$\begin{aligned}
& sub_{i(I),I} : i(I) \times I \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& sub_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = neg_I(z) + \hat{\mathbf{i}} \cdot y
\end{aligned}$$

$$\begin{aligned}
& sub_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& sub_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\
& \quad = sub_I(x, z) + \hat{\mathbf{i}} \cdot neg_I(w)
\end{aligned}$$

$$\begin{aligned}
& sub_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& sub_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\
& \quad = sub_I(x, z) + \hat{\mathbf{i}} \cdot y
\end{aligned}$$

$$\begin{aligned}
& sub_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\} \\
& sub_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
& \quad = neg_I(z) + \hat{\mathbf{i}} \cdot sub_I(y, w)
\end{aligned}$$

$$sub_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} sub_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = x + \hat{\mathbf{i}} \cdot sub_I(y, w) \end{aligned}$$

$$sub_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} sub_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = sub_I(x, z) + \hat{\mathbf{i}} \cdot sub_I(y, w) \end{aligned}$$

$$mul_{i(I)} : i(I) \times i(I) \rightarrow I \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = result_I(-(y \cdot w)) \quad \text{if } y, w \in I \end{aligned}$$

$$mul_{I,i(I)} : I \times i(I) \rightarrow i(I) \cup \{\mathbf{overflow}\}$$

$$mul_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) = \hat{\mathbf{i}} \cdot mul_I(x, w)$$

$$mul_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\mathbf{overflow}\}$$

$$mul_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = \hat{\mathbf{i}} \cdot mul_I(y, z)$$

$$mul_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = mul_I(x, z) + \hat{\mathbf{i}} \cdot mul_I(x, w) \end{aligned}$$

$$mul_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\ = mul_I(x, z) + \hat{\mathbf{i}} \cdot mul_I(y, z) \end{aligned}$$

$$mul_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = result_I(-(y \cdot w)) + \hat{\mathbf{i}} \cdot mul_I(y, z) \\ \text{if } y, z, w \in I \end{aligned}$$

$$mul_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = result_I(-(y \cdot w)) + \hat{\mathbf{i}} \cdot mul_I(x, w) \\ \text{if } x, y, w \in I \end{aligned}$$

$$mul_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}\}$$

$$\begin{aligned} mul_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = result_{c(I)}((x + (\tilde{\mathbf{i}} \cdot y)) \cdot (z + (\tilde{\mathbf{i}} \cdot w))) \\ \text{if } x, y, z, w \in I \end{aligned}$$

5.1.2.6 Absolute value and signum of integers and imaginary integers

NOTE – abs_I is specified in part 1.

$$abs_{i(I)} : i(I) \rightarrow I \cup \{\text{overflow}\}$$

$$abs_{i(I)}(\hat{\mathbf{i}} \cdot y) = abs_I(y)$$

$$signum_I : I \rightarrow \{-1, 1\}$$

$$\begin{aligned} signum_I(x) &= 1 && \text{if } (x \in I \text{ and } x \geq 0) \\ &= -1 && \text{if } (x \in I \text{ and } x < 0) \end{aligned}$$

$$signum_{i(I)} : i(I) \rightarrow \{\hat{\mathbf{i}} \cdot (-1), \hat{\mathbf{i}} \cdot 1\}$$

$$signum_{i(I)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot signum_I(y)$$

5.1.2.7 Divisibility interrogation

$$divides_{i(I)} : i(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ &= divides_I(y, w) \end{aligned}$$

$$divides_{I,i(I)} : I \times i(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) \\ &= divides_I(x, w) \end{aligned}$$

$$divides_{i(I),I} : i(I) \times I \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) \\ &= divides_I(y, z) \end{aligned}$$

$$divides_{I,c(I)} : I \times c(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\ &= divides_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$divides_{c(I),I} : c(I) \times I \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\ &= divides_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0) \end{aligned}$$

$$divides_{i(I),c(I)} : i(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} divides_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ &= divides_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$divides_{c(I),i(I)} : c(I) \times i(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} \text{divides}_{c(I),i(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \text{divides}_{c(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, 0 \mathrel{+} \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$\text{divides}_{c(I)} : c(I) \times c(I) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} \text{divides}_{c(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\ = \mathbf{true} & \quad \text{if } x, y, z, w \in I \text{ and } (x + (\tilde{\mathbf{i}} \cdot y)) \mid (z + (\tilde{\mathbf{i}} \cdot w)) \\ = \mathbf{false} & \quad \text{if } x, y, z, w \in I \text{ and not } (x + (\tilde{\mathbf{i}} \cdot y)) \mid (z + (\tilde{\mathbf{i}} \cdot w)) \end{aligned}$$

5.1.2.8 Integer division and remainder extended to imaginary and complex integers

For these operations, I shall be signed.

NOTE – Even though the integer division operations in principle could be included also for unsigned integer datatypes, that would result in operations that would overflow (mathematically have negative subresults) for so many argument values, that the inclusion of those operations would be pointless. The same goes for the remainder operations.

$$\text{quot}_{i(I)} : i(I) \times i(I) \rightarrow I \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} \text{quot}_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \text{quot}_I(y, w) \end{aligned}$$

$$\text{quot}_{I,i(I)} : I \times i(I) \rightarrow i(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} \text{quot}_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot \text{minint}_I & \quad \text{if } x = \text{minint}_I \text{ and } w = -1 \\ = \hat{\mathbf{i}} \cdot \text{neg}_I(\text{group}_I(x, w)) & \quad \text{otherwise} \end{aligned}$$

$$\text{quot}_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\text{quot}_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = \hat{\mathbf{i}} \cdot \text{quot}_I(y, z)$$

$$\text{quot}_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} \text{quot}_{I,c(I)}(x, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\ = \text{quot}_{c(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot 0, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$\text{quot}_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} \text{quot}_{c(I),I}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, z) \\ = \text{quot}_I(x, z) \mathrel{+} \hat{\mathbf{i}} \cdot \text{quot}_I(y, z) \end{aligned}$$

$$\text{quot}_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} \text{quot}_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\ = \text{quot}_{c(I)}(0 \mathrel{+} \hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$\text{quot}_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned}
quot_{c(I),i(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= neg_I(y) \mathrel{+} \hat{\mathbf{i}} \cdot minint_I \quad \text{if } x = minint_I \text{ and } w = -1 \\
&= quot_I(y, w) \mathrel{+} \hat{\mathbf{i}} \cdot neg_I(group_I(x, w)) \\
&\quad \text{otherwise}
\end{aligned}$$

$$\begin{aligned}
quot_{c(I)} : c(I) \times c(I) &\rightarrow c(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\} \\
quot_{c(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\
&= result_{c(I)}(\lfloor (x + (\tilde{\mathbf{i}} \cdot y)) / (z + (\tilde{\mathbf{i}} \cdot w)) \rfloor) \\
&\quad \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\
&= quot_I(x, 0) \mathrel{+} \hat{\mathbf{i}} \cdot quot_I(y, 0) \\
&\quad \text{otherwise}
\end{aligned}$$

$$\begin{aligned}
mod_{i(I)} : i(I) \times i(I) &\rightarrow i(I) \cup \{\mathbf{invalid}\} \\
mod_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= \hat{\mathbf{i}} \cdot mod_I(y, w)
\end{aligned}$$

$$\begin{aligned}
mod_{I,i(I)} : I \times i(I) &\rightarrow I \cup \{\mathbf{invalid}\} \\
mod_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) \\
&= neg_I(pad_I(x, w))
\end{aligned}$$

$$\begin{aligned}
mod_{i(I),I} : i(I) \times I &\rightarrow i(I) \cup \{\mathbf{invalid}\} \\
mod_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) &= \hat{\mathbf{i}} \cdot mod_I(y, z)
\end{aligned}$$

$$\begin{aligned}
mod_{I,c(I)} : I \times c(I) &\rightarrow c(I) \cup \{\mathbf{invalid}\} \\
mod_{I,c(I)}(x, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\
&= mod_{c(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot 0, z \mathrel{+} \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$\begin{aligned}
mod_{c(I),I} : c(I) \times I &\rightarrow c(I) \cup \{\mathbf{invalid}\} \\
mod_{c(I),I}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, z) \\
&= mod_I(x, z) \mathrel{+} \hat{\mathbf{i}} \cdot mod_I(y, z)
\end{aligned}$$

$$\begin{aligned}
mod_{i(I),c(I)} : i(I) \times c(I) &\rightarrow c(I) \cup \{\mathbf{invalid}\} \\
mod_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\
&= mod_{c(I)}(0 \mathrel{+} \hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$\begin{aligned}
mod_{c(I),i(I)} : c(I) \times i(I) &\rightarrow c(I) \cup \{\mathbf{invalid}\} \\
mod_{c(I),i(I)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= neg_I(pad_I(x, w)) \mathrel{+} \hat{\mathbf{i}} \cdot mod_I(y, w)
\end{aligned}$$

$$mod_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned}
& \text{mod}_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
&= \text{result}_{c(I)}((x + (\tilde{\mathbf{i}} \cdot y)) - (\lfloor (x + (\tilde{\mathbf{i}} \cdot y)) / (z + (\tilde{\mathbf{i}} \cdot w)) \rfloor \cdot (z + (\tilde{\mathbf{i}} \cdot w)))) \\
&\quad \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\
&= \text{invalid}(\text{qNaN} + \hat{\mathbf{i}} \cdot \text{qNaN}) \\
&\quad \text{otherwise}
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{i(I)} : i(I) \times i(I) \rightarrow I \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= \text{ratio}_I(y, w)
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{I,i(I)} : I \times i(I) \rightarrow i(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) \\
&= \hat{\mathbf{i}} \cdot \text{minint}_I \quad \text{if } x = \text{minint}_I \text{ and } w = -1 \\
&= \hat{\mathbf{i}} \cdot \text{neg}_I(\text{ratio}_I(x, w)) \quad \text{otherwise}
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = \hat{\mathbf{i}} \cdot \text{ratio}_I(y, z)
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\
&= \text{ratio}_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\
&= \text{ratio}_I(x, z) + \hat{\mathbf{i}} \cdot \text{ratio}_I(y, z)
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
&= \text{ratio}_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= \text{neg}_I(y) + \hat{\mathbf{i}} \cdot \text{minint}_I \quad \text{if } x = \text{minint}_I \text{ and } w = -1 \\
&= \text{ratio}_I(y, w) + \hat{\mathbf{i}} \cdot \text{neg}_I(\text{ratio}_I(x, w)) \\
&\quad \text{otherwise}
\end{aligned}$$

$$\begin{aligned}
& \text{ratio}_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\text{overflow}, \text{infinitary}, \text{invalid}\} \\
& \text{ratio}_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
&= \text{result}_{c(I)}(\text{round}((x + (\tilde{\mathbf{i}} \cdot y)) / (z + (\tilde{\mathbf{i}} \cdot w)))) \\
&\quad \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\
&= \text{ratio}_I(x, 0) + \hat{\mathbf{i}} \cdot \text{ratio}_I(y, 0) \\
&\quad \text{otherwise}
\end{aligned}$$

$$residue_{i(I)} : i(I) \times i(I) \rightarrow i(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot residue_I(y, w) \end{aligned}$$

$$residue_{I,i(I)} : I \times i(I) \rightarrow I \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) \\ = residue_I(x, w) \end{aligned}$$

$$residue_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) \\ = \hat{\mathbf{i}} \cdot residue_I(y, z) \end{aligned}$$

$$residue_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = residue_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$residue_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\ = residue_I(x, z) + \hat{\mathbf{i}} \cdot residue_I(y, z) \end{aligned}$$

$$residue_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = residue_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$residue_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = residue_I(x, w) + \hat{\mathbf{i}} \cdot residue_I(y, w) \end{aligned}$$

$$residue_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} residue_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = result_{c(I)}((x + (\tilde{\mathbf{i}} \cdot y)) - (\text{round}((x + (\tilde{\mathbf{i}} \cdot y))/(z + (\tilde{\mathbf{i}} \cdot w))) \cdot (z + (\tilde{\mathbf{i}} \cdot w)))) \\ \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\ = \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) \\ \text{otherwise} \end{aligned}$$

$$group_{i(I)} : i(I) \times i(I) \rightarrow I \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned} group_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = group_I(y, w) \end{aligned}$$

$$group_{I,i(I)} : I \times i(I) \rightarrow i(I) \cup \{\mathbf{overflow}, \mathbf{infinitary}, \mathbf{invalid}\}$$

$$\begin{aligned}
group_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) &= \hat{\mathbf{i}} \cdot minint_I && \text{if } x = minint_I \text{ and } w = -1 \\
&= \hat{\mathbf{i}} \cdot neg_I(quot_I(x, w)) && \text{otherwise}
\end{aligned}$$

$$group_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) \\
= \hat{\mathbf{i}} \cdot group_I(y, z)
\end{aligned}$$

$$group_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\
= group_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$group_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\
= group_I(x, z) + \hat{\mathbf{i}} \cdot group_I(y, z)
\end{aligned}$$

$$group_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
= group_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w)
\end{aligned}$$

$$group_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
= neg_I(y) + \hat{\mathbf{i}} \cdot minint_I &&& \text{if } x = minint_I \text{ and } w = -1 \\
= group_I(y, w) + \hat{\mathbf{i}} \cdot neg_I(quot_I(x, w)) &&& \text{otherwise}
\end{aligned}$$

$$group_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\text{overflow, infinitary, invalid}\}$$

$$\begin{aligned}
group_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
= result_{c(I)}(\lceil (x + (\tilde{\mathbf{i}} \cdot y)) / (z + (\tilde{\mathbf{i}} \cdot w)) \rceil) &&& \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\
= group_I(x, 0) + \hat{\mathbf{i}} \cdot group_I(y, 0) &&& \text{otherwise}
\end{aligned}$$

$$pad_{i(I)} : i(I) \times i(I) \rightarrow i(I) \cup \{\text{invalid}\}$$

$$pad_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = \hat{\mathbf{i}} \cdot pad_I(y, w)$$

$$pad_{I,i(I)} : I \times i(I) \rightarrow I \cup \{\text{invalid}\}$$

$$pad_{I,i(I)}(x, \hat{\mathbf{i}} \cdot w) = neg_I(mod_I(x, w))$$

$$pad_{i(I),I} : i(I) \times I \rightarrow i(I) \cup \{\text{invalid}\}$$

$$pad_{i(I),I}(\hat{\mathbf{i}} \cdot y, z) = \hat{\mathbf{i}} \cdot pad_I(y, z)$$

$$pad_{I,c(I)} : I \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} pad_{I,c(I)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = pad_{c(I)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$pad_{c(I),I} : c(I) \times I \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} pad_{c(I),I}(x + \hat{\mathbf{i}} \cdot y, z) \\ = pad_I(x, z) + \hat{\mathbf{i}} \cdot pad_I(y, z) \end{aligned}$$

$$pad_{i(I),c(I)} : i(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} pad_{i(I),c(I)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = pad_{c(I)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$pad_{c(I),i(I)} : c(I) \times i(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} pad_{c(I),i(I)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = neg_I(mod_I(x, w)) + \hat{\mathbf{i}} \cdot pad_I(y, w) \end{aligned}$$

$$pad_{c(I)} : c(I) \times c(I) \rightarrow c(I) \cup \{\mathbf{invalid}\}$$

$$\begin{aligned} pad_{c(I)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = result_{c(I)}(((x + (\tilde{\mathbf{i}} \cdot y)) / (z + (\tilde{\mathbf{i}} \cdot w))) \cdot (z + (\tilde{\mathbf{i}} \cdot w))) - (x + (\tilde{\mathbf{i}} \cdot y))) \\ \quad \text{if } x, y, z, w \in I \text{ and } z + (\tilde{\mathbf{i}} \cdot w) \neq 0 \\ = \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) \\ \quad \text{otherwise} \end{aligned}$$

5.1.2.9 Maximum and minimum

$$max_{i(I)} : i(I) \times i(I) \rightarrow i(I)$$

$$\begin{aligned} max_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot max_I(y, w) \end{aligned}$$

$$min_{i(I)} : i(I) \times i(I) \rightarrow i(I)$$

$$\begin{aligned} min_{i(I)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot min_I(y, w) \end{aligned}$$

$$max_seq_{i(I)} : [i(I)] \rightarrow i(I) \cup \{\mathbf{infinitary}\}$$

$$\begin{aligned} max_seq_{i(I)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\ = \hat{\mathbf{i}} \cdot max_seq_I([y_1, \dots, y_n]) \end{aligned}$$

$$min_seq_{i(I)} : [i(I)] \rightarrow i(I) \cup \{\mathbf{infinitary}\}$$

$$\begin{aligned} min_seq_{i(I)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\ = \hat{\mathbf{i}} \cdot min_seq_I([y_1, \dots, y_n]) \end{aligned}$$

5.2 Imaginary and complex floating point datatypes and operations

Clause 5.2 of part 1 and Clause 5.2 of part 2 of ISO/IEC 10967 specify floating point datatypes and a number of operations on values of a floating point datatype. In this clause imaginary and complex floating point datatypes and operations on values of an imaginary or complex floating point datatype are specified.

NOTE – Further operations on values of an imaginary or complex floating point datatype, for elementary complex floating point numerical functions, are specified in Clause 5.3.

An imaginary or complex floating point datatype is constructed from a floating point datatype. There should be at least one imaginary floating point datatype and at least one complex floating point datatype for each provided floating point datatype.

F is the non-special value set, $F \subset \mathcal{R}$, for a floating point datatype conforming to part 1 of ISO/IEC 10967. Floating point datatypes conforming to part 1 often do contain $-\mathbf{0}$, infinity, and NaN values. Therefore, in this clause there are specifications for such values as arguments.

$i(F)$ (see Clause 4.1.6) is the set of non-special values in an imaginary floating point datatype, constructed from the floating point datatype corresponding to the non-special value set F .

$c(F)$ (see Clause 4.1.6) is the set of non-special values in a complex floating point datatype, constructed from the floating point datatype corresponding to the non-special value set F .

5.2.1 Maximum error requirements

Some of the operations are exact, such as taking the imaginary part of a complex value. Some operations are approximate, with maximum error requirements implied by their exact relationships with operations defined in part 1, part 2, or this part of this International Standard. Some other approximate operations have new error parameters associated with them as detailed in the specifications.

The approximation helper functions for the individual operations in these subclauses have maximum error parameters that describe the maximum *relative* error, in ulps, of the helper function composed with $nearest_F$, for non-subnormal results. The maximum error parameters also describe the maximum *absolute* error, in ulps, for $-fminN_F$, $fminN_F$, subnormal results and underflow continuation values if $denorm_F = \mathbf{true}$. All maximum error parameters shall have a value that is ≥ 0.5 . For the maximum value of the maximum error parameters, see the specification of each of the maximum error parameters. See also Annex A, on partial conformity. The relevant maximum error parameters shall be made available to programs.

When the maximum error for an approximation helper function $h_{c(F)}$, approximating f , is $max_error_op_{c(F)}$, then for all arguments $z, \dots \in \mathcal{C}_F \times \dots$ the following equations shall hold:

$$\begin{aligned} |\Re(f(z, \dots)) - nearest_F(\Re(h_{c(F)}(z, \dots)))| &\leq max_error_op_{c(F)} \cdot r_F^{e_F(\Re(f(z, \dots))) - p_F} \\ |\Im(f(z, \dots)) - nearest_F(\Im(h_{c(F)}(z, \dots)))| &\leq max_error_op_{c(F)} \cdot r_F^{e_F(\Im(f(z, \dots))) - p_F} \end{aligned}$$

Some operations have a **Boolean** parameter $box_error_mode_op_{c(F)}$. If such a parameter is present for an operation and that parameter has the value **true**, then the sign requirements need not be fulfilled (see below) and the maximum error requirements are modified to the following:

$$|\Re(f(z, \dots)) - \text{nearest}_F(\Re(h_{c(F)}(z, \dots)))| \leq \text{max_error_op}_{c(F)} \cdot r_F^{e_F(|f(z, \dots)|) - p_F}$$

$$|\Im(f(z, \dots)) - \text{nearest}_F(\Im(h_{c(F)}(z, \dots)))| \leq \text{max_error_op}_{c(F)} \cdot r_F^{e_F(|f(z, \dots)|) - p_F}$$

The default value for a *box_error_mode_op_{c(F)}* parameter should be **false**.

NOTES

- 1 Partially conforming implementations may have greater values for maximum error parameters than stipulated below. See Annex A.
- 2 Multiplication and division of complex values have the box error mode parameter. See 5.2.6.
- 3 The relative error requirement results in a ‘rectangular’ error bound, which can only span over a zero (in either dimension) when the returned result is very close to 0 (but see ‘Sign requirements’ below).
- 4 The box error requirement results in a ‘square-formed’ error bound, which for cancellation cases can span over zero in either or both dimensions. Relative error requirements in each axis can be fulfilled also for multiplication and division, but that is often considered too inefficient, and an implementation that may suffer from cancellation is often used instead.

5.2.2 Sign requirements

The following sign requirements shall hold:

- a) The real or imaginary part of the result of an approximation helper function shall be zero exactly at the points where the real or imaginary part (respectively) of the approximated mathematical function is exactly zero.
- b) At a point where an approximation helper function result part is not zero, the result part shall have the same sign as the corresponding result part of the approximated mathematical function at that point.

However, the following exceptions are made:

- a) For the trigonometric helper functions, these zero and sign requirements are imposed only for arguments, $x + (\tilde{i} \cdot y)$, such that $|x| \leq \text{big_angle_r}_F$ (see Clause 5.3.2; *big_angle_r_F* is specified in part 2).
- b) If there is a *box_error_mode_op_{c(F)}* parameter for an operation and that parameter has the value **true**, then the sign requirements are not imposed for that operation.

NOTES

- 1 For the operations, the continuation value after an **underflow** may be zero (including negative zero) as given by *result_F*^{*} (see part 2), even though the approximation helper function is not zero at that point. Such zero results are required to be accompanied by an **underflow** notification. When appropriate, zero may also be returned for IEC 60559 infinities arguments. See the individual specifications.
- 2 Multiplication and division of complex values have modified sign requirements. See above and 5.2.6.

5.2.3 Monotonicity requirements

For this document, each approximation helper function shall be a monotonic approximation, for each real and imaginary argument part individually, and the real and imaginary result parts individually, to the mathematical function it is approximating.

For the trigonometric approximation helper functions, the monotonic approximation requirement is imposed only for arguments, $x + (\tilde{\mathbf{i}} \cdot y)$, such that $|x| \leq \text{big_angle_r}_F$ (see Clause 5.3.2). Similarly, for the exponentiation approximation helper functions, the monotonic approximation requirement is imposed only for arguments, $x + (\tilde{\mathbf{i}} \cdot y)$, such that $|y| \leq \text{big_angle_r}_F$ (see Clause 5.3.1).

NOTES

- 1 As in part 2, the monotonicity requirement applies to each real dimension individually. For the complex operations, it thus applies to each real and imaginary part of the argument(s) individually, and individually to the real and imaginary parts of the result.
- 2 As in part 2, the monotonicity requirement applies individually to each monotonic interval of the approximated mathematical function.
- 3 There is an exact *relationship* between the trigonometric and hyperbolic *operations* specified in this document, and thus there are no complex hyperbolic approximation helper functions used in this document.

5.2.4 The complex floating point *result* helper functions

$$\text{result}_{\mathbf{c}(F)}^* : \mathcal{C} \times (\mathcal{R} \rightarrow F^*) \rightarrow \mathbf{c}(F) \cup \{\mathbf{underflow}, \mathbf{overflow}\}$$

$$\begin{aligned} \text{result}_{\mathbf{c}(F)}^*(z, \text{rnd}) \\ = \text{result}_F^*(\Re(z), \text{rnd}) + \hat{\mathbf{i}} \cdot \text{result}_F^*(\Im(z), \text{rnd}) \end{aligned}$$

NOTES

- 1 **overflow** and **underflow** can both occur for a single application of a complex operation. Likewise, e.g., one part may overflow, while the other produced a value in F . This is not exposed accurately by the mathematical framework used in this document. However, handling that accurately, would add complexity to the framework with little gain.
- 2 result_F^* is defined in part 2.

The $\text{no_result}_{\mathbf{c}(F)}$, $\text{no_result}_{\mathbf{i}(F)}$, $\text{no_result}_{F \rightarrow \mathbf{c}(F)}$, $\text{no_result}_{\mathbf{i}(F) \rightarrow \mathbf{c}(F)}$, and $\text{no_result}_{2\mathbf{c}(F)}$ helper functions are defined as follows:

$$\text{no_result}_{\mathbf{c}(F)} : \mathbf{c}(F) \rightarrow \{\mathbf{invalid}\}$$

$$\begin{aligned} \text{no_result}_{\mathbf{c}(F)}(x + \hat{\mathbf{i}} \cdot y) \\ = \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) & \text{ if } x, y \in F \cup \{-\infty, -\mathbf{0}, +\infty\} \\ = \mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN} & \text{ if at least one of } x \text{ and } y \text{ is a quiet NaN and} \\ & \text{neither is a signalling NaN} \\ = \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) & \text{ if } x \text{ is a signalling NaN or } y \text{ is a signalling NaN} \end{aligned}$$

$$\text{no_result}_{\mathbf{i}(F)} : \mathbf{i}(F) \rightarrow \{\mathbf{invalid}\}$$

$$\begin{aligned}
no_result_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \mathbf{invalid}(\hat{\mathbf{i}} \cdot \mathbf{qNaN}) && \text{if } y \in F \cup \{-\infty, -0, +\infty\} \\
&= \hat{\mathbf{i}} \cdot \mathbf{qNaN} && \text{if } y \text{ is a quiet NaN} \\
&= \mathbf{invalid}(\hat{\mathbf{i}} \cdot \mathbf{qNaN}) && \text{if } y \text{ is a signalling NaN}
\end{aligned}$$

$$\begin{aligned}
no_result_{F \rightarrow c(F)} : F &\rightarrow \{\mathbf{invalid}\} \\
no_result_{F \rightarrow c(F)}(x) &= no_result_{c(F)}(x + \hat{\mathbf{i}} \cdot im_F(x))
\end{aligned}$$

$$\begin{aligned}
no_result_{i(F) \rightarrow c(F)} : i(F) &\rightarrow \{\mathbf{invalid}\} \\
no_result_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= no_result_{c(F)}(re_{i(F)}(y) + \hat{\mathbf{i}} \cdot y)
\end{aligned}$$

$$\begin{aligned}
no_result_{2c(F)} : c(F) \times c(F) &\rightarrow \{\mathbf{invalid}\} \\
no_result_{2c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) &= \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) \\
&&& \text{if } x, y, z, w \in F \cup \{-\infty, -0, +\infty\} \\
&= \mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN} && \text{if at least one of } x, y, z, \text{ and } w \text{ is a quiet NaN} \\
&&& \text{and neither is a signalling NaN} \\
&= \mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN}) && \text{if at least one of } x, y, z, \text{ and } w \text{ is a signalling NaN}
\end{aligned}$$

These helper functions are used to specify both NaN argument handling and to handle non-NaN-argument cases where $\mathbf{invalid}(\mathbf{qNaN} + \hat{\mathbf{i}} \cdot \mathbf{qNaN})$ or $\mathbf{invalid}(\hat{\mathbf{i}} \cdot \mathbf{qNaN})$ is the appropriate result.

NOTES

- 3 The handling of other special values, if available, is left unspecified by this document. Other special values may be exact representations for π or e . At a lower level, such special values may be represented by IEC 60559 signalling NaNs, but which aren't regarded as NaNs at the LIA level.
- 4 $re_{i(F)}$ and im_F are defined below.

5.2.5 Basic arithmetic for complex floating point

5.2.5.1 Complex floating point comparisons

$$\begin{aligned}
eq_{i(F)} : i(F) \times i(F) &\rightarrow \mathbf{Boolean} \\
eq_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) &= eq_F(y, w) \\
eq_{F, i(F)} : F \times i(F) &\rightarrow \mathbf{Boolean} \\
eq_{F, i(F)}(x, \hat{\mathbf{i}} \cdot w) &= eq_{c(F)}(x + \hat{\mathbf{i}} \cdot 0, 0 + \hat{\mathbf{i}} \cdot w) \\
eq_{i(F), F} : i(F) \times F &\rightarrow \mathbf{Boolean} \\
eq_{i(F), F}(\hat{\mathbf{i}} \cdot y, z) &= eq_{c(F)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0) \\
eq_{F, c(F)} : F \times c(F) &\rightarrow \mathbf{Boolean}
\end{aligned}$$

$$\begin{aligned} eq_{F,c(F)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = eq_{c(F)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$eq_{c(F),F} : c(F) \times F \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} eq_{c(F),F}(x + \hat{\mathbf{i}} \cdot y, z) \\ = eq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0) \end{aligned}$$

$$eq_{i(F),c(F)} : i(F) \times c(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} eq_{i(F),c(F)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = eq_{c(F)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$eq_{c(F),i(F)} : c(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} eq_{c(F),i(F)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = eq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, 0 + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$eq_{c(F)} : c(F) \times c(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} eq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = \mathbf{true} & \quad \text{if } eq_F(x, z) = \mathbf{true} \text{ and } eq_F(y, w) = \mathbf{true} \\ = \mathbf{false} & \quad \text{if } eq_F(x, z) = \mathbf{false} \text{ and } eq_F(y, w) = \mathbf{true} \\ = \mathbf{false} & \quad \text{if } eq_F(x, z) = \mathbf{true} \text{ and } eq_F(y, w) = \mathbf{false} \\ = \mathbf{false} & \quad \text{if } eq_F(x, z) = \mathbf{false} \text{ and } eq_F(y, w) = \mathbf{false} \\ = \mathbf{invalid}(\mathbf{false}) & \quad \text{otherwise} \end{aligned}$$

$$neq_{i(F)} : i(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$neq_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = neq_F(y, w)$$

$$neq_{F,i(F)} : F \times i(F) \rightarrow \mathbf{Boolean}$$

$$neq_{F,i(F)}(x, \hat{\mathbf{i}} \cdot w) = neq_{c(F)}(x + \hat{\mathbf{i}} \cdot 0, 0 + \hat{\mathbf{i}} \cdot w)$$

$$neq_{i(F),F} : i(F) \times F \rightarrow \mathbf{Boolean}$$

$$neq_{i(F),F}(\hat{\mathbf{i}} \cdot y, z) = neq_{c(F)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0)$$

$$neq_{F,c(F)} : F \times c(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} neq_{F,c(F)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = neq_{c(F)}(x + \hat{\mathbf{i}} \cdot 0, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$neq_{c(F),F} : c(F) \times F \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} neq_{c(F),F}(x + \hat{\mathbf{i}} \cdot y, z) \\ = neq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot 0) \end{aligned}$$

$$neq_{i(F),c(F)} : i(F) \times c(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} neq_{i(F),c(F)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = neq_{c(F)}(0 + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$neq_{c(F),i(F)} : c(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} neq_{c(F),i(F)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = neq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, 0 + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$neq_{c(F)} : c(F) \times c(F) \rightarrow \mathbf{Boolean}$$

$$\begin{aligned} neq_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = \mathbf{true} & \text{ if } neq_F(x, z) = \mathbf{true} \text{ and } neq_F(y, w) = \mathbf{true} \\ = \mathbf{true} & \text{ if } neq_F(x, z) = \mathbf{false} \text{ and } neq_F(y, w) = \mathbf{true} \\ = \mathbf{true} & \text{ if } neq_F(x, z) = \mathbf{true} \text{ and } neq_F(y, w) = \mathbf{false} \\ = \mathbf{false} & \text{ if } neq_F(x, z) = \mathbf{false} \text{ and } neq_F(y, w) = \mathbf{false} \\ = \mathbf{invalid}(\mathbf{true}) & \text{ otherwise} \end{aligned}$$

$$lss_{i(F)} : i(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$lss_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = lss_F(y, w)$$

$$leq_{i(F)} : i(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$leq_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = leq_F(y, w)$$

$$gtr_{i(F)} : i(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$gtr_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = gtr_F(y, w)$$

$$geq_{i(F)} : i(F) \times i(F) \rightarrow \mathbf{Boolean}$$

$$geq_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = geq_F(y, w)$$

5.2.5.2 Multiplication by the imaginary unit

$$itimes_{F \rightarrow i(F)} : F \rightarrow i(F)$$

$$itimes_{F \rightarrow i(F)}(x) = \hat{\mathbf{i}} \cdot x$$

$$itimes_{i(F) \rightarrow F} : i(F) \rightarrow F \cup \{-0\}$$

$$\begin{aligned} itimes_{i(F) \rightarrow F}(\hat{\mathbf{i}} \cdot y) \\ = neg_F(y) \end{aligned}$$

$$itimes_{c(F)} : c(F) \rightarrow c(F \cup \{-0\})$$

$$\begin{aligned} itimes_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\ = neg_F(y) + \hat{\mathbf{i}} \cdot x \end{aligned}$$

5.2.5.3 The real and imaginary parts of a complex value

$$re_F : F \rightarrow F$$

$$\begin{aligned} re_F(x) &= x && \text{if } x \in F \cup \{-\infty, -0, +\infty\} \\ &= no_result_F(x) && \text{otherwise} \end{aligned}$$

$$re_{i(F)} : i(F) \rightarrow \{-0, 0\}$$

$$\begin{aligned} re_{i(F)}(\hat{\mathbf{i}} \cdot y) &= 0 && \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\ &= -0 && \text{if } (y \in F \text{ and } y < 0) \text{ or } y \in \{-\infty, -0\} \\ &= no_result_F(y) && \text{otherwise} \end{aligned}$$

$$re_{c(F)} : c(F) \rightarrow F$$

$$\begin{aligned} re_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= x && \text{if } x \in F \cup \{-\infty, -0, +\infty\} \\ &= no_result_F(x) && \text{otherwise} \end{aligned}$$

$$im_F : F \rightarrow \{-0, 0\}$$

$$\begin{aligned} im_F(x) &= -0 && \text{if } (x \in F \text{ and } x \geq 0) \text{ or } x = +\infty \\ &= 0 && \text{if } (x \in F \text{ and } x < 0) \text{ or } x \in \{-\infty, -0\} \\ &= no_result_F(x) && \text{otherwise} \end{aligned}$$

$$im_{i(F)} : i(F) \rightarrow F$$

$$\begin{aligned} im_{i(F)}(\hat{\mathbf{i}} \cdot y) &= y && \text{if } y \in F \cup \{-\infty, -0, +\infty\} \\ &= no_result_F(y) && \text{otherwise} \end{aligned}$$

$$im_{c(F)} : c(F) \rightarrow F$$

$$\begin{aligned} im_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= y && \text{if } y \in F \cup \{-\infty, -0, +\infty\} \\ &= no_result_F(y) && \text{otherwise} \end{aligned}$$

5.2.5.4 Formation of a complex floating point from two floating point values

$$plusitimes_F : F \times F \rightarrow c(F)$$

$$\begin{aligned} plusitimes_F(x, z) \\ &= x + \hat{\mathbf{i}} \cdot z \end{aligned}$$

5.2.5.5 Fundamental complex floating point arithmetic

NOTE 1 – The addition and subtraction operations never underflow for an IEC 60559 implementation.

$$neg_{i(F)} : i(F) \rightarrow i(F \cup \{-0\})$$

$$neg_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot neg_F(y)$$

$$\text{neg}_{c(F)} : c(F) \rightarrow c(F \cup \{-0\})$$

$$\text{neg}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) = \text{neg}_F(x) \mathrel{+} \hat{\mathbf{i}} \cdot \text{neg}_F(y)$$

$$\text{conj}_F : F \rightarrow F$$

$$\text{conj}_F(x) = x$$

$$\text{conj}_{i(F)} : i(F) \rightarrow i(F \cup \{-0\})$$

$$\text{conj}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{neg}_F(y)$$

$$\text{conj}_{c(F)} : c(F) \rightarrow c(F \cup \{-0\})$$

$$\begin{aligned} \text{conj}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\ = x \mathrel{+} \hat{\mathbf{i}} \cdot \text{neg}_F(y) \end{aligned}$$

$$\text{add}_{i(F)} : i(F) \times i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} \text{add}_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot \text{add}_F(y, w) \end{aligned}$$

$$\text{add}_{F, i(F)} : F \times i(F) \rightarrow c(F)$$

$$\begin{aligned} \text{add}_{F, i(F)}(x, \hat{\mathbf{i}} \cdot w) \\ = x \mathrel{+} \hat{\mathbf{i}} \cdot w \end{aligned}$$

$$\text{add}_{i(F), F} : i(F) \times F \rightarrow c(F)$$

$$\text{add}_{i(F), F}(\hat{\mathbf{i}} \cdot y, z) = z \mathrel{+} \hat{\mathbf{i}} \cdot y$$

$$\text{add}_{F, c(F)} : F \times c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} \text{add}_{F, c(F)}(x, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\ = \text{add}_F(x, z) \mathrel{+} \hat{\mathbf{i}} \cdot w \end{aligned}$$

$$\text{add}_{c(F), F} : c(F) \times F \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} \text{add}_{c(F), F}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, z) \\ = \text{add}_F(x, z) \mathrel{+} \hat{\mathbf{i}} \cdot y \end{aligned}$$

$$\text{add}_{i(F), c(F)} : i(F) \times c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} \text{add}_{i(F), c(F)}(\hat{\mathbf{i}} \cdot y, z \mathrel{+} \hat{\mathbf{i}} \cdot w) \\ = z \mathrel{+} \hat{\mathbf{i}} \cdot \text{add}_F(y, w) \end{aligned}$$

$$\text{add}_{c(F), i(F)} : c(F) \times i(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} \text{add}_{c(F), i(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = x \mathrel{+} \hat{\mathbf{i}} \cdot \text{add}_F(y, w) \end{aligned}$$

$$add_{c(F)} : c(F) \times c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} add_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = add_F(x, z) + \hat{\mathbf{i}} \cdot add_F(y, w) \end{aligned}$$

$$sub_{i(F)} : i(F) \times i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$sub_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) = \hat{\mathbf{i}} \cdot sub_F(y, w)$$

$$sub_{F,i(F)} : F \times i(F) \rightarrow c(F \cup \{-0\})$$

$$sub_{F,i(F)}(x, \hat{\mathbf{i}} \cdot w) = x + \hat{\mathbf{i}} \cdot neg_F(w)$$

$$sub_{i(F),F} : i(F) \times F \rightarrow c(F \cup \{-0\})$$

$$sub_{i(F),F}(\hat{\mathbf{i}} \cdot y, z) = neg_F(z) + \hat{\mathbf{i}} \cdot y$$

$$sub_{F,c(F)} : F \times c(F) \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} sub_{F,c(F)}(x, z + \hat{\mathbf{i}} \cdot w) \\ = sub_F(x, z) + \hat{\mathbf{i}} \cdot neg_F(w) \end{aligned}$$

$$sub_{c(F),F} : c(F) \times F \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} sub_{c(F),F}(x + \hat{\mathbf{i}} \cdot y, z) \\ = sub_F(x, z) + \hat{\mathbf{i}} \cdot y \end{aligned}$$

$$sub_{i(F),c(F)} : i(F) \times c(F) \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} sub_{i(F),c(F)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = neg_F(z) + \hat{\mathbf{i}} \cdot sub_F(y, w) \end{aligned}$$

$$sub_{c(F),i(F)} : c(F) \times i(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} sub_{c(F),i(F)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = x + \hat{\mathbf{i}} \cdot sub_F(y, w) \end{aligned}$$

$$sub_{c(F)} : c(F) \times c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} sub_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\ = sub_F(x, z) + \hat{\mathbf{i}} \cdot sub_F(y, w) \end{aligned}$$

$$mul_{i(F)} : i(F) \times i(F) \rightarrow F \cup \{-0, \text{underflow}, \text{overflow}\}$$

$$\begin{aligned} mul_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = neg_F(mul_F(y, w)) \end{aligned}$$

$$mul_{F,i(F)} : F \times i(F) \rightarrow i(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\}$$

$$\begin{aligned} mul_{F,i(F)}(x, \hat{\mathbf{i}} \cdot w) \\ = \hat{\mathbf{i}} \cdot mul_F(x, w) \end{aligned}$$

$$\begin{aligned} mul_{i(F),F} : i(F) \times F &\rightarrow i(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\ mul_{i(F),F}(\hat{\mathbf{i}} \cdot y, z) &= \hat{\mathbf{i}} \cdot mul_F(y, z) \end{aligned}$$

$$\begin{aligned} mul_{F,c(F)} : F \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\ mul_{F,c(F)}(x, z + \hat{\mathbf{i}} \cdot w) &= mul_F(x, z) + \hat{\mathbf{i}} \cdot mul_F(x, w) \end{aligned}$$

$$\begin{aligned} mul_{c(F),F} : c(F) \times F &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\ mul_{c(F),F}(x + \hat{\mathbf{i}} \cdot y, z) &= mul_F(x, z) + \hat{\mathbf{i}} \cdot mul_F(y, z) \end{aligned}$$

$$\begin{aligned} mul_{i(F),c(F)} : i(F) \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\ mul_{i(F),c(F)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) &= neg_F(mul_F(y, w)) + \hat{\mathbf{i}} \cdot mul_F(y, z) \end{aligned}$$

$$\begin{aligned} mul_{c(F),i(F)} : c(F) \times i(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\ mul_{c(F),i(F)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) &= neg_F(mul_F(y, w)) + \hat{\mathbf{i}} \cdot mul_F(x, w) \end{aligned}$$

NOTE 2 – $mul_{c(F)}$ is specified in clause 5.2.6.

$$\begin{aligned} div_{i(F)} : i(F) \times i(F) &\rightarrow F \cup \{-0, \text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) &= div_F(y, w) \end{aligned}$$

$$\begin{aligned} div_{F,i(F)} : F \times i(F) &\rightarrow i(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{F,i(F)}(x, \hat{\mathbf{i}} \cdot w) &= \hat{\mathbf{i}} \cdot neg_F(div_F(x, w)) \end{aligned}$$

$$\begin{aligned} div_{i(F),F} : i(F) \times F &\rightarrow i(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{i(F),F}(\hat{\mathbf{i}} \cdot y, z) &= \hat{\mathbf{i}} \cdot div_F(y, z) \end{aligned}$$

$$\begin{aligned} div_{F,c(F)} : F \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{F,c(F)}(x, z + \hat{\mathbf{i}} \cdot w) &= div_{c(F)}(x + \hat{\mathbf{i}} \cdot im_F(x), z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$\begin{aligned} div_{c(F),F} : c(F) \times F &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{c(F),F}(x + \hat{\mathbf{i}} \cdot y, z) &= div_F(x, z) + \hat{\mathbf{i}} \cdot div_F(y, z) \end{aligned}$$

$$\begin{aligned} div_{i(F),c(F)} : i(F) \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\ div_{i(F),c(F)}(\hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) &= div_{c(F)}(re_{i(F)}(\hat{\mathbf{i}} \cdot y) + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \end{aligned}$$

$$\text{div}_{c(F),i(F)} : c(F) \times i(F) \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\}$$

$$\begin{aligned} \text{div}_{c(F),i(F)}(x + \hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\ = \text{div}_F(y, w) + \hat{\mathbf{i}} \cdot \text{neg}_F(\text{div}_F(x, w)) \end{aligned}$$

NOTE 3 – $\text{div}_{c(F)}$ is specified in clause 5.2.6.

5.2.5.6 Absolute value, phase and signum of complex floating point values

NOTES

1 This clause also includes signum_F .

2 abs_F is specified in part 1.

$$\text{abs}_{i(F)} : i(F) \rightarrow F$$

$$\text{abs}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \text{abs}_F(y)$$

$$\text{abs}_{c(F)} : c(F) \rightarrow F \cup \{\text{underflow}, \text{overflow}\}$$

$$\text{abs}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) = \text{hypot}_F(x, y)$$

$$\text{phase}_F : F \rightarrow F$$

$$\text{phase}_F(x) = \text{arc}_F(x, \text{im}_F(x))$$

NOTE 3 – The arc_F (and similarly arcu_F) operation as specified in clause 5.3.8.15 of the *first edition* (issued in 2001) of part 2 has a minor flaw; it is not properly limited. See annex F for a corrected version. Implementations are recommended to follow this improvement, if possible, for arc_F , arcu_F , and all operations that reference those specifications.

$$\text{phase}_{i(F)} : i(F) \rightarrow F$$

$$\text{phase}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \text{arc}_F(\text{re}_{i(F)}(\hat{\mathbf{i}} \cdot y), y)$$

$$\text{phase}_{c(F)} : c(F) \rightarrow F \cup \{\text{underflow}\}$$

$$\begin{aligned} \text{phase}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\ = \text{arc}_F(x, y) \end{aligned}$$

$$\text{phaseu}_F : F \times F \rightarrow F \cup \{\text{invalid}\}$$

$$\text{phaseu}_F(u, x) = \text{arcu}_F(u, x, \text{im}_F(x))$$

$$\text{phaseu}_{i(F)} : F \times i(F) \rightarrow F \cup \{\text{invalid}\}$$

$$\begin{aligned} \text{phaseu}_{i(F)}(u, \hat{\mathbf{i}} \cdot y) \\ = \text{arcu}_F(u, \text{re}_{i(F)}(\hat{\mathbf{i}} \cdot y), y) \end{aligned}$$

$$\text{phaseu}_{c(F)} : F \times c(F) \rightarrow F \cup \{\text{underflow}, \text{invalid}\}$$

$$\begin{aligned} \text{phaseu}_{c(F)}(u, x + \hat{\mathbf{i}} \cdot y) \\ = \text{arcu}_F(u, x, y) \end{aligned}$$

$$\begin{aligned}
\text{signum}_F : F &\rightarrow \{-1, 1\} \\
\text{signum}_F(x) &= 1 && \text{if } (x \in F \text{ and } x \geq 0) \text{ or } x = +\infty \\
&= -1 && \text{if } (x \in F \text{ and } x < 0) \text{ or } x \in \{-\infty, -0\} \\
&= \text{no_result}_F(x) && \text{otherwise}
\end{aligned}$$

$$\text{signum}_{i(F)} : i(F) \rightarrow \{\hat{\mathbf{i}} \cdot (-1), \hat{\mathbf{i}} \cdot 1\}$$

$$\text{signum}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{signum}_F(y)$$

NOTE 4 – $\text{signum}_{c(F)}$ is specified in clause 5.2.6.

5.2.5.7 Floor, round, and ceiling

$$\text{floor}_{i(F)} : i(F) \rightarrow i(F)$$

$$\text{floor}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{floor}_F(y)$$

$$\text{floor}_{c(F)} : c(F) \rightarrow c(F)$$

$$\begin{aligned}
\text{floor}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\
= \text{floor}_F(x) \mathrel{+} \hat{\mathbf{i}} \cdot \text{floor}_F(y)
\end{aligned}$$

$$\text{rounding}_{i(F)} : i(F) \rightarrow i(F)$$

$$\begin{aligned}
\text{rounding}_{i(F)}(\hat{\mathbf{i}} \cdot y) \\
= \hat{\mathbf{i}} \cdot \text{rounding}_F(y)
\end{aligned}$$

$$\text{rounding}_{c(F)} : c(F) \rightarrow c(F)$$

$$\begin{aligned}
\text{rounding}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\
= \text{rounding}_F(x) \mathrel{+} \hat{\mathbf{i}} \cdot \text{rounding}_F(y)
\end{aligned}$$

$$\text{ceiling}_{i(F)} : i(F) \rightarrow i(F)$$

$$\text{ceiling}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{ceiling}_F(y)$$

$$\text{ceiling}_{c(F)} : c(F) \rightarrow c(F)$$

$$\begin{aligned}
\text{ceiling}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\
= \text{ceiling}_F(x) \mathrel{+} \hat{\mathbf{i}} \cdot \text{ceiling}_F(y)
\end{aligned}$$

5.2.5.8 Maximum and minimum

$$\text{max}_{i(F)} : i(F) \times i(F) \rightarrow i(F)$$

$$\begin{aligned}
\text{max}_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
= \hat{\mathbf{i}} \cdot \text{max}_F(y, w)
\end{aligned}$$

$$\text{mma}_{i(F)} : i(F) \times i(F) \rightarrow i(F)$$

$$\begin{aligned}
\text{mma}_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
= \hat{\mathbf{i}} \cdot \text{mma}_F(y, w)
\end{aligned}$$

$$\begin{aligned}
min_{i(F)} : i(F) \times i(F) &\rightarrow i(F) \\
min_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= \hat{\mathbf{i}} \cdot min_F(y, w)
\end{aligned}$$

$$\begin{aligned}
mmin_{i(F)} : i(F) \times i(F) &\rightarrow i(F) \\
mmin_{i(F)}(\hat{\mathbf{i}} \cdot y, \hat{\mathbf{i}} \cdot w) \\
&= \hat{\mathbf{i}} \cdot mmin_F(y, w)
\end{aligned}$$

$$\begin{aligned}
max_seq_{i(F)} : [i(F)] &\rightarrow i(F) \cup \{\mathbf{infinitary}\} \\
max_seq_{i(F)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\
&= \hat{\mathbf{i}} \cdot max_seq_F([y_1, \dots, y_n])
\end{aligned}$$

$$\begin{aligned}
mmax_seq_{i(F)} : [i(F)] &\rightarrow i(F) \cup \{\mathbf{infinitary}\} \\
mmax_seq_{i(F)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\
&= \hat{\mathbf{i}} \cdot mmax_seq_F([y_1, \dots, y_n])
\end{aligned}$$

$$\begin{aligned}
min_seq_{i(F)} : [i(F)] &\rightarrow i(F) \cup \{\mathbf{infinitary}\} \\
min_seq_{i(F)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\
&= \hat{\mathbf{i}} \cdot min_seq_F([y_1, \dots, y_n])
\end{aligned}$$

$$\begin{aligned}
mmin_seq_{i(F)} : [i(F)] &\rightarrow i(F) \cup \{\mathbf{infinitary}\} \\
mmin_seq_{i(F)}([\hat{\mathbf{i}} \cdot y_1, \dots, \hat{\mathbf{i}} \cdot y_n]) \\
&= \hat{\mathbf{i}} \cdot mmin_seq_F([y_1, \dots, y_n])
\end{aligned}$$

5.2.6 Complex sign, multiplication, and division

This clause gives specifications for some operations that have been deferred above, due to the possibility of lesser accuracy.

There shall be a box error mode parameter each for the multiplication and division operations on a complex datatype:

$$\begin{aligned}
box_error_mode_mul_{c(F)} &\in \mathbf{Boolean} \\
box_error_mode_div_{c(F)} &\in \mathbf{Boolean}
\end{aligned}$$

There shall be a maximum error parameter each for multiplication and division on a complex datatype:

$$\begin{aligned}
max_error_mul_{c(F)} &\in F \\
max_error_div_{c(F)} &\in F
\end{aligned}$$

The $max_error_mul_{c(F)}$ parameter shall have a value that is ≤ 5 . The $max_error_div_{c(F)}$ parameter shall have a value that is ≤ 13 .

For use in the specification below, define the mathematical complex sign function

$$signum : \mathcal{C} \rightarrow \mathcal{C}$$

The *signum* function is defined by

$$\text{signum}(z) = \cos(\text{arc}(\Re(z), \Im(z))) + (\tilde{\mathbf{i}} \cdot \sin(\text{arc}(\Re(z), \Im(z))))$$

NOTES

- 1 The *arc* function is defined in part 2, Clause 5.3.7.
- 2 Note that $z = |z| \cdot \text{signum}(z)$ if $z \in \mathcal{C}$, and $\text{signum}(x + (\tilde{\mathbf{i}} \cdot y)) = e^{\tilde{\mathbf{i}} \cdot \text{arc}(x, y)}$ if $x, y \in \mathcal{R}$.

5.2.6.1 Complex signum

The $\text{signum}_{\mathbf{c}(F)}^*$ approximation helper function:

$$\text{signum}_{\mathbf{c}(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{signum}_{\mathbf{c}(F)}^*(z)$ returns a close approximation to $\text{signum}(z)$ in $\mathcal{C}$, with maximum error max_error_tan_F (interpreted as an error parameter for a complex floating point operation).

Further requirements on the $\text{signum}_{\mathbf{c}(F)}^*$ approximation helper function are:

$$\begin{aligned} \text{signum}_{\mathbf{c}(F)}^*(\text{conj}(z)) &= \text{conj}(\text{signum}_{\mathbf{c}(F)}^*(z)) && \text{if } z \in \mathcal{C}_F \\ \text{signum}_{\mathbf{c}(F)}^*(-z) &= -\text{signum}_{\mathbf{c}(F)}^*(z) && \text{if } z \in \mathcal{C}_F \\ \Re(\text{signum}_{\mathbf{c}(F)}^*(x + (\tilde{\mathbf{i}} \cdot y))) &= \Im(\text{signum}_{\mathbf{c}(F)}^*(y + (\tilde{\mathbf{i}} \cdot x))) && \text{if } x, y \in F \\ \text{signum}_{\mathbf{c}(F)}^*(x) &= 1 && \text{if } x \in F \text{ and } x > 0 \\ \text{signum}_{\mathbf{c}(F)}^*(x + (\tilde{\mathbf{i}} \cdot x)) &= (1/\sqrt{2}) + (\tilde{\mathbf{i}} \cdot (1/\sqrt{2})) && \text{if } x \in F \text{ and } x > 0 \end{aligned}$$

The $\text{signum}_{\mathbf{c}(F)}$ operation:

$$\begin{aligned} \text{signum}_{\mathbf{c}(F)} : \mathbf{c}(F) &\rightarrow \mathbf{c}(F) \cup \{\mathbf{underflow}\} \\ \text{signum}_{\mathbf{c}(F)}(x + \tilde{\mathbf{i}} \cdot y) &= \text{result}_{\mathbf{c}(F)}^*(\text{signum}_{\mathbf{c}(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) && \text{if } x, y \in F \text{ and } x + (\tilde{\mathbf{i}} \cdot y) \neq 0 \\ &= \text{conj}_{\mathbf{c}(F)}(\text{signum}_{\mathbf{c}(F)}(\text{conj}_{\mathbf{c}(F)}(x + \tilde{\mathbf{i}} \cdot y))) && \text{if } y \in \{-\infty, -0\} \\ &= \text{neg}_{\mathbf{c}(F)}(\text{signum}_{\mathbf{c}(F)}(\text{neg}_{\mathbf{c}(F)}(x + \tilde{\mathbf{i}} \cdot y))) && \text{if } x \in \{-\infty, -0\} \text{ and } y \notin \{-\infty, -0\} \\ &= \sin_F(\text{arc}_F(y, x)) + \tilde{\mathbf{i}} \cdot \sin_F(\text{arc}_F(x, y)) && \text{otherwise} \end{aligned}$$

NOTE – $\text{signum}_{\mathbf{c}(F)}(x + \tilde{\mathbf{i}} \cdot y) \approx \cos_F(\text{arc}_F(x, y)) + \tilde{\mathbf{i}} \cdot \sin_F(\text{arc}_F(x, y))$. The specification is more complicated, in order to allow higher accuracy, including the sign of zero result parts.

5.2.6.2 Complex multiplication

The $\text{mul}_{\mathbf{c}(F)}^*$ approximation helper function:

$$\text{mul}_{\mathbf{c}(F)}^* : \mathcal{C}_F \times \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{mul}_{\mathbf{c}(F)}^*(x, z)$ returns a close approximation to $x \cdot z$ in $\mathcal{C}$ with maximum error $\text{max_error_mul}_{\mathbf{c}(F)}$, interpreted according to the value of $\text{box_error_mode_mul}_{\mathbf{c}(F)}$.

Further requirements on the $\text{mul}_{\mathbf{c}(F)}^*$ approximation helper function are:

$$\begin{aligned}
mul_{c(F)}^*(conj(z), conj(z')) &= conj(mul_{c(F)}^*(z, z')) \\
&\quad \text{if } z, z' \in \mathcal{C}_F \\
mul_{c(F)}^*(-z, z') &= -mul_{c(F)}^*(z, z') \\
&\quad \text{if } z, z' \in \mathcal{C}_F \\
mul_{c(F)}^*(z, z') &= mul_{c(F)}^*(z', z) \\
&\quad \text{if } z, z' \in \mathcal{C}_F
\end{aligned}$$

The $mul_{c(F)}$ operation:

$$\begin{aligned}
mul_{c(F)} : c(F) \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}\} \\
mul_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
&= result_{c(F)}^*(mul_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y), z + (\tilde{\mathbf{i}} \cdot w)), nearest_F) \\
&\quad \text{if } x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w \in c(F) \text{ and} \\
&\quad x \neq 0 \text{ and } y \neq 0 \text{ and } z \neq 0 \text{ and } w \neq 0 \\
&= sub_F(mul_F(x, z), mul_F(y, w)) + \hat{\mathbf{i}} \cdot add_F(mul_F(y, z), mul_F(x, w)) \\
&\quad \text{otherwise}
\end{aligned}$$

NOTE – NaN (and **invalid**) is not avoided for the result in the “otherwise” case here. Note in particular cases like $mul_{c(F)}(2 + \hat{\mathbf{i}} \cdot (-0), 3 + \hat{\mathbf{i}} \cdot (+\infty))$ which is **invalid** with a continuation value of $\mathbf{qNaN} + \hat{\mathbf{i}} \cdot (+\infty)$. However, $mul_{F, c(F)}(2, 3 + \hat{\mathbf{i}} \cdot (+\infty))$ returns $6 + \hat{\mathbf{i}} \cdot (+\infty)$, due to the strong implicit zero (see B.5.2.4).

5.2.6.3 Complex division

The $div_{c(F)}^*$ approximation helper function:

$$div_{c(F)}^* : \mathcal{C}_F \times \mathcal{C}_F \rightarrow \mathcal{C}$$

$div_{c(F)}^*(x, z)$ returns a close approximation to x/z in $\mathcal{C}$ with maximum error $max_error_div_{c(F)}$, interpreted according to the value of $box_error_mode_div_{c(F)}$.

Further requirements on the $div_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}
div_{c(F)}^*(conj(z), conj(z')) &= conj(div_{c(F)}^*(z, z')) \\
&\quad \text{if } z, z' \in \mathcal{C}_F \text{ and } z' \neq 0 \\
div_{c(F)}^*(-z, z') &= -div_{c(F)}^*(z, z') \\
&\quad \text{if } z, z' \in \mathcal{C}_F \text{ and } z' \neq 0 \\
div_{c(F)}^*(z, -z') &= -div_{c(F)}^*(z, z') \\
&\quad \text{if } z, z' \in \mathcal{C}_F \text{ and } z' \neq 0
\end{aligned}$$

The $div_{c(F)}$ operation:

$$\begin{aligned}
div_{c(F)} : c(F) \times c(F) &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{invalid}\} \\
div_{c(F)}(x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w) \\
&= result_{c(F)}^*(div_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y), z + (\tilde{\mathbf{i}} \cdot w)), nearest_F) \\
&\quad \text{if } x + \hat{\mathbf{i}} \cdot y, z + \hat{\mathbf{i}} \cdot w \in c(F) \text{ and} \\
&\quad x \neq 0 \text{ and } y \neq 0 \text{ and } z \neq 0 \text{ and } w \neq 0 \\
&= div_{c(F), F}(add_F(mul_F(x, z), mul_F(y, w)) + \hat{\mathbf{i}} \cdot sub_F(mul_F(y, z), mul_F(x, w)), \\
&\quad add_F(mul_F(z, z), mul_F(w, w))) \\
&\quad \text{otherwise}
\end{aligned}$$

5.2.7 Operations for conversion from polar to Cartesian

The $\text{polar}_{c(F)}^*$ approximation helper function:

$$\text{polar}_{c(F)}^* : F \times F^{2 \cdot \pi} \rightarrow \mathcal{C}$$

$\text{polar}_{c(F)}^*(x, z)$ returns a close approximation to $x \cdot e^{\tilde{i} \cdot z}$ in $\mathcal{C}$ with maximum error max_error_tan_F (interpreted as an error parameter for a complex floating point operation).

Further requirements on the $\text{polar}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \text{polar}_{c(F)}^*(-x, z) &= -\text{polar}_{c(F)}^*(x, z) & \text{if } x, z \in F \text{ and } x \geq 0 \\ \text{polar}_{c(F)}^*(x, -z) &= \text{conj}(\text{polar}_{c(F)}^*(x, z)) & \text{if } x, z \in F \end{aligned}$$

$$\begin{aligned} \text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi) &= x & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \Im(\text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + \pi/6)) &= x/2 & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \Re(\text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + \pi/3)) &= x/2 & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + \pi/2) &= \tilde{i} \cdot x & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \Re(\text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + 2 \cdot \pi/3)) &= -x/2 & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \Im(\text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + 5 \cdot \pi/6)) &= x/2 & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \\ \text{polar}_{c(F)}^*(x, n \cdot 2 \cdot \pi + \pi) &= -x & \text{if } x \in F \text{ and } n \in \mathbb{Z} \text{ and } |n \cdot 2 \cdot \pi| \leq \text{big_angle_r}_F \end{aligned}$$

The $\text{polar}_{c(F)}$ operation:

$$\begin{aligned} \text{polar}_{c(F)} : F \times F &\rightarrow c(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\} \\ \text{polar}_{c(F)}(x, z) &= \text{result}_{c(F)}^*(\text{polar}_{c(F)}^*(x, z), \text{nearest}_F) \\ &\quad \text{if } x, z \in F \text{ and } |z| \leq \text{big_angle_r}_F \\ &\quad \text{and } x \neq 0 \text{ and } z \neq 0 \\ &= \text{mul}_F(x, \cos_F(z)) + \tilde{i} \cdot \text{mul}_F(x, \sin_F(z)) \\ &\quad \text{otherwise} \end{aligned}$$

The $\text{polaru}_{c(F)}^*$ approximation helper function:

$$\text{polaru}_{c(F)}^* : (u : F) \times F \times F^u \rightarrow \mathcal{C}$$

$\text{polaru}_{c(F)}^*(u, x, z)$ returns a close approximation to $x \cdot e^{\tilde{i} \cdot 2 \cdot \pi \cdot z / u}$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_F(u)$ (interpreted as an error parameter for a complex floating point operation).

Further requirements on the $\text{polaru}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \text{polaru}_{c(F)}^*(u, -x, z) &= -\text{polaru}_{c(F)}^*(u, x, z) & \text{if } u, x \in F \text{ and } z \in F^u \text{ and } u \neq 0 \text{ and } x \geq 0 \\ \text{polaru}_{c(F)}^*(u, x, -z) &= \text{conj}(\text{polaru}_{c(F)}^*(u, x, z)) & \text{if } u, x \in F \text{ and } z \in F^u \text{ and } u \neq 0 \\ \text{polaru}_{c(F)}^*(-u, x, z) &= \text{polaru}_{c(F)}^*(u, x, -z) & \text{if } u, x \in F \text{ and } z \in F^u \text{ and } u \neq 0 \\ \text{polaru}_{c(F)}^*(u, x, (n \cdot u) + z) &= \text{polaru}_{c(F)}^*(u, x, z) & \text{if } u, x \in F \text{ and } z, (n \cdot u) + z \in F^u \text{ and } \\ & & u \neq 0 \text{ and } n \in \mathbb{Z} \end{aligned}$$

$$\begin{aligned} \text{polaru}_{c(F)}^*(u, x, 0) &= x & \text{if } u, x \in F \text{ and } u \neq 0 \\ \Im(\text{polaru}_{c(F)}^*(u, x, u/12)) &= x/2 & \text{if } u, x \in F \text{ and } u \neq 0 \\ \Re(\text{polaru}_{c(F)}^*(u, x, u/6)) &= x/2 & \text{if } u, x \in F \text{ and } u \neq 0 \\ \text{polaru}_{c(F)}^*(u, x, u/4) &= \tilde{i} \cdot x & \text{if } u, x \in F \text{ and } u \neq 0 \\ \Re(\text{polaru}_{c(F)}^*(u, x, u/3)) &= -x/2 & \text{if } u, x \in F \text{ and } u \neq 0 \end{aligned}$$

$$\begin{aligned} \Im(\text{polaru}_{c(F)}^*(u, x, 5 \cdot u/12)) &= x/2 && \text{if } u, x \in F \text{ and } u \neq 0 \\ \text{polaru}_{c(F)}^*(u, x, n \cdot u/2) &= -x && \text{if } u, x \in F \text{ and } u \neq 0 \end{aligned}$$

The $\text{polaru}_{c(F)}$ operation:

$$\begin{aligned} \text{polaru}_{c(F)} : F \times F \times F &\rightarrow c(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}, \text{invalid}\} \\ \text{polaru}_{c(F)}(u, x, z) &= \text{result}_{c(F)}^*(\text{polaru}_{c(F)}^*(u, x, z), \text{nearest}_F) \\ &\quad \text{if } u \in G_F \text{ and } x, z \in F \text{ and } |z/u| \leq \text{big_angle_}u_F \\ &\quad \text{and } x \neq 0 \text{ and } z \neq 0 \\ &= \text{mul}_F(x, \cos u_F(u, z)) + \hat{\mathbf{i}} \cdot \text{mul}_F(x, \sin u_F(u, z)) \\ &\quad \text{otherwise} \end{aligned}$$

NOTE – G_F is defined in part 2.

5.3 Elementary transcendental imaginary and complex floating point operations

Clause 5.3 of part 2 specifies a number of transcendental floating point operations. In this clause a number of transcendental imaginary and complex floating point operations are specified.

Several of the specifications below include relationships to specifications for operations specified in part 2. The requirements implied by these relationships and the requirements from part 2 shall hold even if none or only some of the mentioned operations are included in an associated library (which may be the same library as for the complex operations) for real-valued operations or even if there is no associated library for real-valued operations.

5.3.1 Operations for exponentiations and logarithms

There shall be two maximum error parameters for complex exponentiations and logarithms.

$$\begin{aligned} \text{max_error_exp}_{c(F)} &\in F \\ \text{max_error_power}_{c(F)} &\in F \end{aligned}$$

The $\text{max_error_exp}_{c(F)}$ parameter shall have a value that is ≤ 7 . The $\text{max_error_power}_{c(F)}$ parameter shall have a value that is ≤ 15 .

5.3.1.1 Exponentiation of imaginary base to integer power

The $\text{power}_{i(F), I}$ operation:

$$\begin{aligned} \text{power}_{i(F), I} : i(F) \times I &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}\} \\ \text{power}_{i(F), I}(\hat{\mathbf{i}} \cdot x, y) &= \text{re}_{i(F)}(\hat{\mathbf{i}} \cdot x) + \hat{\mathbf{i}} \cdot \text{power}_{F, I}(x, y) \\ &\quad \text{if } y \in I \text{ and } 4|(y+3) \\ &= \text{neg}_F(\text{power}_{F, I}(\text{abs}_F(x), y)) + \hat{\mathbf{i}} \cdot 0 \\ &\quad \text{if } y \in I \text{ and } 4|(y+2) \\ &= \text{neg}_{c(F)}(\text{re}_{i(F)}(\hat{\mathbf{i}} \cdot x) + \hat{\mathbf{i}} \cdot \text{power}_{F, I}(x, y)) \\ &\quad \text{if } y \in I \text{ and } 4|(y+1) \\ &= \text{power}_{F, I}(\text{abs}_F(x), y) + \hat{\mathbf{i}} \cdot (-0) \\ &\quad \text{if } y \in I \text{ and } 4|y \end{aligned}$$

5.3.1.2 Natural exponentiation

The $\exp_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned} \exp_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\} \\ \exp_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= \cos_F(y) + \hat{\mathbf{i}} \cdot \sin_F(y) \end{aligned}$$

NOTE 1 – Some programming languages have the operation **cis**. **cis**(x) is $\exp_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot x)$.

The $\exp_{c(F)}^*$ approximation helper function:

$$\exp_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\exp_{c(F)}^*(z)$ returns a close approximation to e^z in $\mathcal{C}$ with maximum error $\max_error_exp_{c(F)}$.

A further requirement on the $\exp_{c(F)}^*$ approximation helper function is:

$$\exp_{c(F)}^*(\text{conj}(z)) = \text{conj}(\exp_{c(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F$$

The relationships to the $\cos_F^*$, $\sin_F^*$, and $\exp_F^*$ approximation helper functions for the $\cos_F$, $\sin_F$, and $\exp_F$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned} \exp_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \cos_F^*(y) + (\tilde{\mathbf{i}} \cdot \sin_F^*(y)) & \text{if } y \in F \\ \exp_{c(F)}^*(x) &= \exp_F^*(x) & \text{if } x \in F \end{aligned}$$

The $\exp_{c(F)}$ operation:

$$\begin{aligned} \exp_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\} \\ \exp_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{result}_{c(F)}^*(\exp_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |y| \leq \text{big_angle_r}_F \\ &= \exp_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) & \text{if } x = -0 \\ &= \text{conj}_{c(F)}(\exp_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) & \text{if } y = -0 \text{ and } x \neq -0 \\ &= \text{mul}_F(0, \cos_F(y)) + \hat{\mathbf{i}} \cdot \text{mul}_F(0, \sin_F(y)) & \text{if } x = -\infty \text{ and } y \in F \text{ and } |y| \leq \text{big_angle_r}_F \\ &= \text{mul}_F(+\infty, \cos_F(y)) + \hat{\mathbf{i}} \cdot \text{mul}_F(+\infty, \sin_F(y)) & \text{if } x = +\infty \text{ and } y \in F \text{ and } |y| \leq \text{big_angle_r}_F \text{ and } y \neq 0 \\ &= (+\infty) + \hat{\mathbf{i}} \cdot 0 & \text{if } x = +\infty \text{ and } y \in F \text{ and } y = 0 \\ &= \text{radh}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) & \text{otherwise} \end{aligned}$$

NOTES

2 $\text{radh}_{c(F)}$ is specified in Clause 5.3.3.1.

3 **invalid** is avoided here for the cases $\exp_{c(F)}((+\infty) + \hat{\mathbf{i}} \cdot 0)$ and $\exp_{c(F)}((+\infty) + \hat{\mathbf{i}} \cdot (-0))$.

5.3.1.3 Complex exponentiation of argument base

The $\text{power}_{F \rightarrow c(F)}^*$ helper function:

$$\text{power}_{F \rightarrow c(F)}^* : F \times F \rightarrow \mathcal{C}$$

$\text{power}_{F \rightarrow c(F)}^*(x, z)$ returns a close approximation to x^z in $\mathcal{C}$ with maximum error $\max_error_power_{c(F)}$.

NOTE – $x^z = (-1 \cdot (-x))^z = (-1)^z \cdot (-x)^z = e^{\tilde{\mathbf{i}} \cdot \pi \cdot z} \cdot (-x)^z$ if $x < 0$

Further requirements on the $\text{power}_{F \rightarrow c(F)}^*$ are:

$$\begin{aligned}\Re(power_{F \rightarrow c(F)}^*(x, z)) &= \Im(power_{F \rightarrow c(F)}^*(x, z)) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } z - 0.25 \in \mathcal{Z} \\ \Re(power_{F \rightarrow c(F)}^*(x, z)) &= -\Im(power_{F \rightarrow c(F)}^*(x, z)) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } z - 0.75 \in \mathcal{Z}\end{aligned}$$

The relationships to the $power_F^*$, $sinu_F^*$, and $cosu_F^*$ helper functions for the $power_F$, $sinu_F$, and $cosu_F$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned}power_{F \rightarrow c(F)}^*(x, z) &= power_F^*(x, z) && \text{if } x, z \in F \text{ and } x \geq 0 \\ power_{F \rightarrow c(F)}^*(x, z) &= power_F^*(-x, z) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } z/2 \in \mathcal{Z} \\ power_{F \rightarrow c(F)}^*(x, z) &= \tilde{1} \cdot power_F^*(-x, z) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } (z - 0.5)/2 \in \mathcal{Z} \\ power_{F \rightarrow c(F)}^*(x, z) &= -power_F^*(-x, z) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } (z - 1)/2 \in \mathcal{Z} \\ power_{F \rightarrow c(F)}^*(x, z) &= -\tilde{1} \cdot power_F^*(-x, z) && \text{if } x, z \in F \text{ and } x < 0 \text{ and } (z - 1.5)/2 \in \mathcal{Z} \\ power_{F \rightarrow c(F)}^*(-1, z) &= cosu_F^*(2, z) + (\tilde{1} \cdot sinu_F^*(2, z))\end{aligned}$$

The $power_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}power_{F \rightarrow c(F)} &: F \times F \rightarrow c(F) \cup \{\text{underflow, overflow, absolute_precision_underflow, infinitary, invalid}\} \\ power_{F \rightarrow c(F)}(x, z) &= result_{c(F)}^*(power_{F \rightarrow c(F)}^*(x, z), nearest_F) \\ &\quad \text{if } x, z \in F \text{ and } (x \geq 0 \text{ or } (x < 0 \text{ and } |z/2| \leq big_angle_u_F)) \\ &= power_F(x, z) + \hat{1} \cdot 0 && \text{if } x \in F \text{ and } x \geq 0 \text{ and } z \in \{-\infty, -0, +\infty\} \\ &= power_F(x, z) + \hat{1} \cdot 0 && \text{if } x = +\infty \text{ and } z \in F \cup \{-\infty, -0, +\infty\} \\ &= mul_F(power_F(-x, z), cosu_F(2, z)) + \hat{1} \cdot sinu_F(2, z) && \text{if } x \in F \text{ and } x < 0 \text{ and } z \in \{-\infty, -0, +\infty\} \\ &= mul_F(power_F(neg_F(x), z), cosu_F(2, z)) + \hat{1} \cdot sinu_F(2, z) && \text{if } x \in \{-\infty, -0\} \text{ and } z \in \{-\infty, -0, +\infty\} \\ &= mul_{F, c(F)}(power_F(neg_F(x), z), cosu_F(2, z) + \hat{1} \cdot sinu_F(2, z)) && \text{if } x \in \{-\infty, -0\} \text{ and } z \in F \text{ and } |z/2| \leq big_angle_u_F \\ &= \text{absolute_precision_underflow}(qNaN + \hat{1} \cdot qNaN) && \text{if } x \in F \text{ and } x < 0 \text{ and } z \in F \text{ and } |z/2| > big_angle_u_F \\ &= \text{absolute_precision_underflow}(qNaN + \hat{1} \cdot qNaN) && \text{if } x \in \{-\infty, -0\} \text{ and } z \in F \text{ and } |z/2| > big_angle_u_F \\ &= no_result_{c(F)}(x, z) && \text{otherwise}\end{aligned}$$

5.3.1.4 Complex square root

The $sqrt_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}sqrt_{F \rightarrow c(F)} &: F \rightarrow c(F) \\ sqrt_{F \rightarrow c(F)}(x) &= 0 + \hat{1} \cdot sqrt_F(neg_F(x)) && \text{if } (x \in F \text{ and } x < 0) \text{ or } x \in \{-\infty, -0\} \\ &= sqrt_F(x) + \hat{1} \cdot (-0) && \text{if } (x \in F \text{ and } x \geq 0) \text{ or } x = +\infty \\ &= no_result_{c(F)}(x + \hat{1} \cdot im_F(x)) && \text{otherwise}\end{aligned}$$

The $\text{sqrt}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}\text{sqrt}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \\ \text{sqrt}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= \text{sqrt}_{c(F)}(\text{re}_{i(F)}(\hat{\mathbf{i}} \cdot y) + \hat{\mathbf{i}} \cdot y)\end{aligned}$$

The $\text{sqrt}_{c(F)}^*$ approximation helper function:

$$\text{sqrt}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{sqrt}_{c(F)}^*(z)$ returns a close approximation to $\sqrt{z}$ in $\mathcal{C}$ with maximum error $\text{max_error_exp}_{c(F)}$.

Further requirements on the $\text{sqrt}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}\text{sqrt}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{sqrt}_{c(F)}^*(z)) && \text{if } z \in \mathcal{C}_F \\ \text{sqrt}_{c(F)}^*(x) &= \sqrt{x} && \text{if } x \in F \text{ and } x \geq 0 \\ \text{sqrt}_{c(F)}^*(x) &= \tilde{\mathbf{i}} \cdot \sqrt{-x} && \text{if } x \in F \text{ and } x < 0 \\ \Re(\text{sqrt}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y)) &= \Im(\text{sqrt}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y)) && \text{if } y \in F \text{ and } y \geq 0\end{aligned}$$

The $\text{sqrt}_{c(F)}$ operation:

$$\begin{aligned}\text{sqrt}_{c(F)} : c(F) &\rightarrow c(F) \\ \text{sqrt}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{result}_{c(F)}^*(\text{sqrt}_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \\ &= \text{sqrt}_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) && \text{if } x = -0 \\ &= \text{conj}_{c(F)}(\text{sqrt}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } x \in F \cup \{-\infty, +\infty\} \text{ and } y = -0 \\ &= (+\infty) + \hat{\mathbf{i}} \cdot (+\infty) && \text{if } x \in F \cup \{-\infty, +\infty\} \text{ and } y = +\infty \\ &= (+\infty) + \hat{\mathbf{i}} \cdot 0 && \text{if } x = +\infty \text{ and } y \in F \text{ and } y \geq 0 \\ &= (+\infty) + \hat{\mathbf{i}} \cdot (-0) && \text{if } x = +\infty \text{ and } y \in F \text{ and } y < 0 \\ &= (+\infty) + \hat{\mathbf{i}} \cdot (-\infty) && \text{if } x \in F \cup \{-\infty, +\infty\} \text{ and } y = -\infty \\ &= 0 + \hat{\mathbf{i}} \cdot (+\infty) && \text{if } x = -\infty \text{ and } y \in F \text{ and } y \geq 0 \\ &= 0 + \hat{\mathbf{i}} \cdot (-\infty) && \text{if } x = -\infty \text{ and } y \in F \text{ and } y < 0 \\ &= \text{no_result}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) && \text{otherwise}\end{aligned}$$

NOTE – The inverse of complex square is multi-valued. The principal result is given by $\sqrt{b} = e^{0.5 \cdot \ln(b)}$. The $\sqrt{}$ function branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } x < 0\}$. Thus $\text{sqrt}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq \text{sqrt}_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $x < 0$.

5.3.1.5 Natural logarithm

The $\text{ln}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}\text{ln}_{F \rightarrow c(F)} : F &\rightarrow c(F) \cup \{\text{infinitary}\} \\ \text{ln}_{F \rightarrow c(F)}(x) &= \text{ln}_F(\text{abs}_F(x)) + \hat{\mathbf{i}} \cdot \text{arc}_F(x, \text{im}_F(x))\end{aligned}$$

NOTE 1 – The arc_F (and similarly arcu_F) operation as specified in Clause 5.3.8.15 of the first edition (issued in 2001) of part 2 has a minor flaw as noted in Annex F of this document.

The $\text{ln}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}\text{ln}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \cup \{\text{infinitary}\} \\ \text{ln}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= \text{ln}_F(\text{abs}_F(y)) + \hat{\mathbf{i}} \cdot \text{arc}_F(\text{re}_{i(F)}(\hat{\mathbf{i}} \cdot y), y)\end{aligned}$$

The $\text{ln}_{c(F)}^*$ approximation helper function:

$$ln_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$ln_{c(F)}^*(z)$ returns a close approximation to $\ln(z)$ in $\mathcal{C}$ with maximum error $max_error_exp_{c(F)}$ in the real part.

NOTE 2 – Since the imaginary part of the result of $ln_{c(F)}(x + \hat{\mathbf{i}} \cdot y)$ is $arc_F(x, y)$, the maximum error in the imaginary part is $max_error_tan_F$. Thus $max_error_exp_{c(F)}$ reflects the maximum error in the real part of the result.

A further requirement on the $ln_{c(F)}^*$ approximation helper function is:

$$ln_{c(F)}^*(\text{conj}(z)) = \text{conj}(ln_{c(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F$$

The relationships to the arc_F^* and ln_F^* approximation helper functions for the arc_F and ln_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned} \Im(ln_{c(F)}^*(x + (\hat{\mathbf{i}} \cdot y))) &= arc_F^*(x, y) & \text{if } x, y \in F \text{ and } x + (\hat{\mathbf{i}} \cdot y) \neq 0 \\ \Re(ln_{c(F)}^*(x)) &= ln_F^*(|x|) & \text{if } x \in F \text{ and } x \neq 0 \\ \Re(ln_{c(F)}^*(\hat{\mathbf{i}} \cdot y)) &= ln_F^*(|y|) & \text{if } y \in F \text{ and } y \neq 0 \end{aligned}$$

The $ln_{c(F)}^\#$ range limitation helper function (for $z \in \mathcal{C}_F$):

$$\begin{aligned} ln_{c(F)}^\#(z) &= \Re(ln_{c(F)}^*(z)) + (\hat{\mathbf{i}} \cdot \max\{up_F(-\pi), \min\{\Im(ln_{c(F)}^*(z)), down_F(-\pi/2)\}\}) \\ &\quad \text{if } \Re(z) < 0 \text{ and } \Im(z) < 0 \\ &= \Re(ln_{c(F)}^*(z)) + (\hat{\mathbf{i}} \cdot \max\{up_F(-\pi/2), \min\{\Im(ln_{c(F)}^*(z)), down_F(\pi/2)\}\}) \\ &\quad \text{if } \Re(z) \geq 0 \\ &= \Re(ln_{c(F)}^*(z)) + (\hat{\mathbf{i}} \cdot \max\{up_F(\pi/2), \min\{\Im(ln_{c(F)}^*(z)), down_F(\pi)\}\}) \\ &\quad \text{if } \Re(z) < 0 \text{ and } \Im(z) \geq 0 \end{aligned}$$

The $ln_{c(F)}$ operation:

$$\begin{aligned} ln_{c(F)} : c(F) &\rightarrow c(F) \cup \{\mathbf{infinitary}\} \\ ln_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= result_{c(F)}^*(ln_{c(F)}^\#(x + (\hat{\mathbf{i}} \cdot y)), nearest_F) \\ &\quad \text{if } x, y \in F \text{ and } (x \neq 0 \text{ or } y \neq 0) \\ &= \mathbf{infinitary}((-\infty) + \hat{\mathbf{i}} \cdot arc_F(x, y)) \\ &\quad \text{if } x \in \{-0, 0\} \text{ and } y = 0 \\ &= conj_{c(F)}(ln_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -0 \\ &= ln_F(y) + \hat{\mathbf{i}} \cdot up_F(\pi/2) & \text{if } x = -0 \text{ and } ((y \in F \text{ and } y > 0) \text{ or } y = +\infty) \\ &= ln_F(neg_F(y)) + \hat{\mathbf{i}} \cdot down_F(-\pi/2) \\ &\quad \text{if } x = -0 \text{ and } ((y \in F \text{ and } y < 0) \text{ or } y = -\infty) \\ &= (+\infty) + \hat{\mathbf{i}} \cdot arc_F(x, y) & \text{if } x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\} \\ &= (+\infty) + \hat{\mathbf{i}} \cdot arc_F(x, y) & \text{if } x \in F \text{ and } y \in \{-\infty, +\infty\} \\ &= no_result_{c(F)}(x + \hat{\mathbf{i}} \cdot y) & \text{otherwise} \end{aligned}$$

NOTES

- 3 The inverse of natural exponentiation is multi-valued: the imaginary part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The $\ln$ function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } x < 0\}$, is continuous on the rest of $\mathcal{C}$, and $\ln(z) \in \mathcal{R}$ if $x \in \mathcal{R}$ and $x > 0$. Thus $ln_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq ln_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $x < 0$.
- 4 $re_{c(F)}(ln_{c(F)}(x + \hat{\mathbf{i}} \cdot y)) \approx ln_F(hypot_F(x, y))$ and $im_{c(F)}(ln_{c(F)}(x + \hat{\mathbf{i}} \cdot y)) = arc_F(x, y)$ when there is no notification (if the specification of arc_F is corrected as noted in Clause 5.2.5).

5.3.2 Operations for radian trigonometric elementary functions

There shall be two maximum error parameters for complex trigonometric operations.

$$\text{max_error_sin}_{\text{c}(F)} \in F$$

$$\text{max_error_tan}_{\text{c}(F)} \in F$$

The $\text{max_error_sin}_{\text{c}(F)}$ parameter shall have a value that is ≤ 11 . The $\text{max_error_tan}_{\text{c}(F)}$ parameter shall have a value that is ≤ 14 .

5.3.2.1 Radian angle normalisation

$$\text{rad}_{\text{i}(F)} : \text{i}(F) \rightarrow \text{i}(F)$$

$$\text{rad}_{\text{i}(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot y$$

$$\text{rad}_{\text{c}(F)} : \text{c}(F) \rightarrow \text{c}(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\}$$

$$\text{rad}_{\text{c}(F)}(x + \hat{\mathbf{i}} \cdot y) = \text{rad}_F(x) + \hat{\mathbf{i}} \cdot y$$

5.3.2.2 Radian sine

The $\text{sin}_{\text{i}(F)}$ operation:

$$\text{sin}_{\text{i}(F)} : \text{i}(F) \rightarrow \text{i}(F) \cup \{\text{overflow}\}$$

$$\text{sin}_{\text{i}(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{sinh}_F(y)$$

The $\text{sin}_{\text{c}(F)}^*$ approximation helper function:

$$\text{sin}_{\text{c}(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{sin}_{\text{c}(F)}^*(z)$ returns a close approximation to $\sin(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_sin}_{\text{c}(F)}$.

Further requirements on the $\text{sin}_{\text{c}(F)}^*$ approximation helper function are:

$$\text{sin}_{\text{c}(F)}^*(\text{conj}(z)) = \text{conj}(\text{sin}_{\text{c}(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F$$

$$\text{sin}_{\text{c}(F)}^*(-z) = -\text{sin}_{\text{c}(F)}^*(z) \quad \text{if } z \in \mathcal{C}_F$$

The relationships to the sin_F^* and sinh_F^* approximation helper functions for sin_F and sinh_F operations in an associated library for real-valued operations shall be:

$$\text{sin}_{\text{c}(F)}^*(x) = \text{sin}_F^*(x) \quad \text{if } x \in F$$

$$\text{sin}_{\text{c}(F)}^*(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{sinh}_F^*(y) \quad \text{if } y \in F$$

The $\text{sin}_{\text{c}(F)}$ operation:

$$\text{sin}_{\text{c}(F)} : \text{c}(F) \rightarrow \text{c}(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\}$$

$$\begin{aligned} \text{sin}_{\text{c}(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{result}_{\text{c}(F)}^*(\text{sin}_{\text{c}(F)}^*(x + (\hat{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |x| \leq \text{big_angle_r}_F \end{aligned}$$

$$\begin{aligned} &= \text{conj}_{\text{c}(F)}(\text{sin}_{\text{c}(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -0 \end{aligned}$$

$$\begin{aligned} &= \text{neg}_{\text{c}(F)}(\text{sin}_{\text{c}(F)}(0 + \hat{\mathbf{i}} \cdot \text{neg}_F(y))) \\ &\quad \text{if } x = -0 \text{ and } y \neq -0 \end{aligned}$$

$$\begin{aligned} &= \text{mul}_F(\text{sin}_F(x), +\infty) + \hat{\mathbf{i}} \cdot \text{mul}_F(\text{cos}_F(x), y) \\ &\quad \text{if } x \notin \{-0, 0\} \text{ and } y \in \{-\infty, +\infty\} \end{aligned}$$

$$= 0 + \hat{\mathbf{i}} \cdot y \quad \text{if } x = 0 \text{ and } y \in \{-\infty, +\infty\}$$

$$= \text{rad}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \quad \text{otherwise}$$

5.3.2.3 Radian cosine

The $\text{cos}_{i(F)}$ operation:

$$\begin{aligned} \text{cos}_{i(F)} : i(F) &\rightarrow F \cup \{\text{overflow}\} \\ \text{cos}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \text{cosh}_F(y) \end{aligned}$$

The $\text{cos}_{c(F)}^*$ approximation helper function:

$$\text{cos}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{cos}_{c(F)}^*(z)$ returns a close approximation to $\cos(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_sin}_{c(F)}$.

Further requirements on the $\text{cos}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \text{cos}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{cos}_{c(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \\ \text{cos}_{c(F)}^*(-z) &= \text{cos}_{c(F)}^*(z) & \text{if } z \in \mathcal{C}_F \end{aligned}$$

The relationships to the cos_F^* and cosh_F^* approximation helper functions for cos_F and cosh_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned} \text{cos}_{c(F)}^*(x) &= \text{cos}_F^*(x) & \text{if } x \in F \\ \text{cos}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \text{cosh}_F^*(y) & \text{if } y \in F \end{aligned}$$

The $\text{cos}_{c(F)}$ operation:

$$\begin{aligned} \text{cos}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\} \\ \text{cos}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{result}_{c(F)}^*(\text{cos}_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |x| \leq \text{big_angle_r}_F \\ &= \text{conj}_{c(F)}(\text{cos}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) & \text{if } y = -0 \\ &= \text{cos}_{c(F)}(0 + \hat{\mathbf{i}} \cdot \text{neg}_F(y)) & \text{if } x = -0 \text{ and } y \neq -0 \\ &= \text{mul}_F(\text{cos}_F(x), +\infty) + \hat{\mathbf{i}} \cdot \text{mul}_F(\text{sin}_F(x), \text{neg}_F(y)) \\ &\quad \text{if } x \notin \{-0, 0\} \text{ and } y \in \{-\infty, +\infty\} \\ &= (+\infty) + \hat{\mathbf{i}} \cdot \text{div}_F(-1, y) & \text{if } x = 0 \text{ and } y \in \{-\infty, +\infty\} \\ &= \text{rad}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) & \text{otherwise} \end{aligned}$$

5.3.2.4 Radian tangent

The $\text{tan}_{i(F)}$ operation:

$$\begin{aligned} \text{tan}_{i(F)} : i(F) &\rightarrow i(F) \\ \text{tan}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{tanh}_F(y) \end{aligned}$$

The $\text{tan}_{c(F)}^*$ approximation helper function:

$$\text{tan}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{tan}_{c(F)}^*(z)$ returns a close approximation to $\tan(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_{c(F)}$.

Further requirements on the $\text{tan}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \text{tan}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{tan}_{c(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \\ \text{tan}_{c(F)}^*(-z) &= -\text{tan}_{c(F)}^*(z) & \text{if } z \in \mathcal{C}_F \end{aligned}$$

The relationships to the $\tan_F^*$ and $\tanh_F^*$ approximation helper functions for $\tan_F$ and $\tanh_F$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned}\tan_{c(F)}^*(x) &= \tan_F^*(x) && \text{if } x \in F \\ \tan_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \tilde{\mathbf{i}} \cdot \tanh_F^*(y) && \text{if } y \in F\end{aligned}$$

The $\tan_{c(F)}$ operation:

$$\begin{aligned}\tan_{c(F)} : c(F) &\rightarrow c(F) \cup \{\mathbf{underflow}, \mathbf{overflow}, \mathbf{absolute_precision_underflow}\} \\ \tan_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \mathit{result}_{c(F)}^*(\tan_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \mathit{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |x| \leq \mathit{big_angle_r}_F \\ &= \mathit{conj}_{c(F)}(\tan_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -\mathbf{0} \\ &= \mathit{neg}_F(\tan_{c(F)}(0 + \hat{\mathbf{i}} \cdot \mathit{neg}_F(y))) \\ &\quad \text{if } x = -\mathbf{0} \text{ and } y \neq -\mathbf{0} \\ &= \mathit{mul}_F(\tan_F(x), 0) + \hat{\mathbf{i}} \cdot \mathit{signum}_F(y) \\ &\quad \text{if } x \neq -\mathbf{0} \text{ and } y \in \{-\infty, +\infty\} \\ &= \mathit{rad}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) && \text{otherwise}\end{aligned}$$

5.3.2.5 Radian cotangent

The $\cot_{i(F)}$ operation:

$$\begin{aligned}\cot_{i(F)} : i(F) &\rightarrow i(F) \cup \{\mathbf{overflow}, \mathbf{infinitary}\} \\ \cot_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \mathit{neg}_F(\mathit{coth}_F(y))\end{aligned}$$

The $\cot_{c(F)}^*$ approximation helper function:

$$\cot_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\cot_{c(F)}^*(z)$ returns a close approximation to $\cot(z)$ in $\mathcal{C}$ with maximum error $\mathit{max_error_tan}_{c(F)}$.

Further requirements on the $\cot_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}\cot_{c(F)}^*(\mathit{conj}(z)) &= \mathit{conj}(\cot_{c(F)}^*(z)) && \text{if } z \in \mathcal{C}_F \\ \cot_{c(F)}^*(-z) &= -\cot_{c(F)}^*(z) && \text{if } z \in \mathcal{C}_F\end{aligned}$$

The relationships to the $\cot_F^*$ and coth_F^* approximation helper functions for $\cot_F$ and coth_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned}\cot_{c(F)}^*(x) &= \cot_F^*(x) && \text{if } x \in F \\ \cot_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \mathit{coth}_F^*(y) && \text{if } y \in F\end{aligned}$$

The $\cot_{c(F)}$ operation:

$$\begin{aligned}\cot_{c(F)} : c(F) &\rightarrow c(F) \cup \{\mathbf{underflow}, \mathbf{overflow}, \mathbf{infinitary}, \mathbf{absolute_precision_underflow}\} \\ \cot_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \mathit{result}_{c(F)}^*(\cot_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \mathit{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |x| \leq \mathit{big_angle_r}_F \text{ and } (x \neq 0 \text{ or } y \neq 0) \\ &= \mathbf{infinitary}((+\infty) + \hat{\mathbf{i}} \cdot (-\infty)) \\ &\quad \text{if } x = 0 \text{ and } y = 0 \\ &= \mathit{conj}_{c(F)}(\cot_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -\mathbf{0} \\ &= \mathit{neg}_F(\cot_{c(F)}(0 + \hat{\mathbf{i}} \cdot \mathit{neg}_F(y))) \\ &\quad \text{if } x = -\mathbf{0} \text{ and } y \neq -\mathbf{0}\end{aligned}$$

$$\begin{aligned}
&= \text{mul}_F(\tan_F(x), 0) + \hat{\mathbf{i}} \cdot \text{neg}_F(\text{signum}_F(y)) \\
&\quad \text{if } x \neq -0 \text{ and } y \in \{-\infty, +\infty\} \\
&= \text{rad}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \quad \text{otherwise}
\end{aligned}$$

NOTE – The reference to $\tan_F$ (instead of $\cot_F$) for the case when x is different from -0 and y is an infinity, is in order to avoid **invalid** for $\cot_{c(F)}$.

5.3.2.6 Radian secant

The $\text{sec}_{i(F)}$ operation:

$$\begin{aligned}
\text{sec}_{i(F)} : i(F) &\rightarrow F \cup \{\text{underflow}\} \\
\text{sec}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \text{sech}_F(y)
\end{aligned}$$

The $\text{sec}_{c(F)}^*$ approximation helper function:

$$\text{sec}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{sec}_{c(F)}^*(z)$ returns a close approximation to $\sec(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_{c(F)}$.

Further requirements on the $\text{sec}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}
\text{sec}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{sec}_{c(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \\
\text{sec}_{c(F)}^*(-z) &= \text{sec}_{c(F)}^*(z) & \text{if } z \in \mathcal{C}_F
\end{aligned}$$

The relationships to the sec_F^* and sech_F^* approximation helper functions for sec_F and sech_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned}
\text{sec}_{c(F)}^*(x) &= \text{sec}_F^*(x) & \text{if } x \in F \\
\text{sec}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \text{sech}_F^*(y) & \text{if } y \in F
\end{aligned}$$

The $\text{sec}_{c(F)}$ operation:

$$\begin{aligned}
\text{sec}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\} \\
\text{sec}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{result}_{c(F)}^*(\text{sec}_{c(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\
&\quad \text{if } x, y \in F \text{ and } |x| \leq \text{big_angle_r}_F \\
&= \text{conj}_{c(F)}(\text{sec}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\
&\quad \text{if } y = -0 \\
&= \text{sec}_{c(F)}(0 + \hat{\mathbf{i}} \cdot \text{neg}_F(y)) & \text{if } x = -0 \text{ and } y \neq -0 \\
&= \text{mul}_F(\cos_F(x), 0) + \hat{\mathbf{i}} \cdot \text{div}_F(\sin_F(x), y) \\
&\quad \text{if } x \neq -0 \text{ and } y \in \{-\infty, +\infty\} \\
&= \text{rad}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) & \text{otherwise}
\end{aligned}$$

5.3.2.7 Radian cosecant

The $\text{csc}_{i(F)}$ operation:

$$\begin{aligned}
\text{csc}_{i(F)} : i(F) &\rightarrow i(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}\} \\
\text{csc}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{neg}_F(\text{csch}_F(y))
\end{aligned}$$

The $\text{csc}_{c(F)}^*$ approximation helper function:

$$\text{csc}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{csc}_{c(F)}^*(z)$ returns a close approximation to $\csc(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_{c(F)}$.

Further requirements on the $\text{csc}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \csc_{\mathcal{C}(F)}^*(\text{conj}(z)) &= \text{conj}(\csc_{\mathcal{C}(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \\ \csc_{\mathcal{C}(F)}^*(-z) &= -\csc_{\mathcal{C}(F)}^*(z) & \text{if } z \in \mathcal{C}_F \end{aligned}$$

The relationships to the $\csc_F^*$ and $\csc_{c(F)}^*$ approximation helper functions for $\csc_F$ and $\csc_{c(F)}$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned} \csc_{\mathcal{C}(F)}^*(x) &= \csc_F^*(x) & \text{if } x \in F \\ \csc_{\mathcal{C}(F)}^*(\tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \csc_F^*(y) & \text{if } y \in F \end{aligned}$$

The $\csc_{\mathcal{C}(F)}$ operation:

$$\begin{aligned} \csc_{\mathcal{C}(F)} : \mathcal{C}(F) &\rightarrow \mathcal{C}(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{absolute_precision_underflow}\} \\ \csc_{\mathcal{C}(F)}(x + \tilde{\mathbf{i}} \cdot y) &= \text{result}_{\mathcal{C}(F)}^*(\csc_{\mathcal{C}(F)}^*(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } |x| \leq \text{big_angle_r}_F \text{ and } (x \neq 0 \text{ or } y \neq 0) \\ &= \text{infinitary}((+\infty) + \tilde{\mathbf{i}} \cdot (-\infty)) \\ &\quad \text{if } x = 0 \text{ and } y = 0 \\ &= \text{conj}_{\mathcal{C}(F)}(\csc_{\mathcal{C}(F)}(x + \tilde{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -0 \\ &= \text{neg}_F(\csc_{\mathcal{C}(F)}(0 + \tilde{\mathbf{i}} \cdot \text{neg}_F(y))) \\ &\quad \text{if } x = -0 \text{ and } y \neq -0 \\ &= \text{mul}_F(\sin_F(x), 0) + \tilde{\mathbf{i}} \cdot \text{div}_F(\cos_F(x), \text{neg}_F(y)) \\ &\quad \text{if } x \neq -0 \text{ and } y \in \{-\infty, +\infty\} \\ &= \text{rad}_{\mathcal{C}(F)}(x + \tilde{\mathbf{i}} \cdot y) & \text{otherwise} \end{aligned}$$

5.3.2.8 Radian arc sine

The $\arcsin_{F \rightarrow \mathcal{C}(F)}$ operation:

$$\begin{aligned} \arcsin_{F \rightarrow \mathcal{C}(F)} : F &\rightarrow \mathcal{C}(F \cup \{-0\}) \\ \arcsin_{F \rightarrow \mathcal{C}(F)}(x) &= \text{up}_F(-\pi/2) + \tilde{\mathbf{i}} \cdot \text{arccosh}_F(\text{neg}_F(x)) \\ &\quad \text{if } (x \in F \text{ and } x < -1) \text{ or } x = -\infty \\ &= \arcsin_F(x) + \tilde{\mathbf{i}} \cdot \text{mul}_F(-0, x) \\ &\quad \text{if } (x \in F \text{ and } |x| \leq 1) \text{ or } x = -0 \\ &= \text{down}_F(\pi/2) + \tilde{\mathbf{i}} \cdot \text{neg}_F(\text{arccosh}_F(x)) \\ &\quad \text{if } (x \in F \text{ and } x > 1) \text{ or } x = +\infty \\ &= \text{no_result}_{F \rightarrow \mathcal{C}(F)}(x) & \text{otherwise} \end{aligned}$$

The $\arcsin_{\mathbf{i}(F)}$ operation:

$$\begin{aligned} \arcsin_{\mathbf{i}(F)} : \mathbf{i}(F) &\rightarrow \mathbf{i}(F) \\ \arcsin_{\mathbf{i}(F)}(\tilde{\mathbf{i}} \cdot y) &= \tilde{\mathbf{i}} \cdot \text{arcsinh}_F(y) \end{aligned}$$

The $\arcsin_{\mathcal{C}(F)}^*$ approximation helper function:

$$\arcsin_{\mathcal{C}(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\arcsin_{\mathcal{C}(F)}^*(z)$ returns a close approximation to $\arcsin(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_sin}_{\mathcal{C}(F)}$.

Further requirements on the $\arcsin_{\mathcal{C}(F)}^*$ approximation helper function are:

$$\begin{aligned} \arcsin_{\mathcal{C}(F)}^*(\text{conj}(z)) &= \text{conj}(\arcsin_{\mathcal{C}(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \text{ and } (|\Re(z)| \leq 1 \text{ or } \Im(z) \neq 0) \\ \arcsin_{\mathcal{C}(F)}^*(-z) &= -\arcsin_{\mathcal{C}(F)}^*(z) & \text{if } z \in \mathcal{C}_F \text{ and } (|\Re(z)| \leq 1 \text{ or } \Im(z) \neq 0) \end{aligned}$$

The relationships to the $\arcsin_F^*$, $\operatorname{arsinh}_F^*$, and $\operatorname{arccosh}_F^*$ approximation helper functions for $\arcsin_F$, arsinh_F , and $\operatorname{arccosh}_F$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned}\arcsin_{c(F)}^*(x) &= (-\pi/2) + (\tilde{\imath} \cdot \operatorname{arccosh}_F^*(-x)) && \text{if } x \in F \text{ and } x < -1 \\ \arcsin_{c(F)}^*(x) &= \arcsin_F^*(x) && \text{if } x \in F \text{ and } |x| \leq 1 \\ \arcsin_{c(F)}^*(x) &= (\pi/2) + (\tilde{\imath} \cdot \operatorname{arccosh}_F^*(x)) && \text{if } x \in F \text{ and } x > 1 \\ \arcsin_{c(F)}^*(\tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \operatorname{arsinh}_F^*(y) && \text{if } y \in F\end{aligned}$$

The $\arcsin_{c(F)}^\#$ range limitation helper function:

$$\arcsin_{c(F)}^\#(z) = \max\{up_F(-\pi/2), \min\{\Re(\arcsin_{c(F)}^*(z)), down_F(\pi/2)\}\} + (\tilde{\imath} \cdot \Im(\arcsin_{c(F)}^*(z)))$$

The $\arcsin_{c(F)}$ operation:

$$\begin{aligned}\arcsin_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}\} \\ \arcsin_{c(F)}(x + \tilde{\imath} \cdot y) &= result_{c(F)}^*(\arcsin_{c(F)}^\#(x + (\tilde{\imath} \cdot y)), nearest_F) && \text{if } x, y \in F \\ &= neg_{c(F)}(\arcsin_{c(F)}(0 + \tilde{\imath} \cdot neg_F(y))) && \text{if } x = -0 \\ &= conj_{c(F)}(\arcsin_{c(F)}(x + \tilde{\imath} \cdot 0)) && \text{if } y = -0 \text{ and } x \neq -0 \\ &= arc_F(abs_F(y), x) + \tilde{\imath} \cdot mul_F(signum_F(y), +\infty) && \text{if } (x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\}) \text{ or } \\ &\quad (y \in \{-\infty, +\infty\} \text{ and } x \in F) \\ &= no_result_{c(F)}(x + \tilde{\imath} \cdot y) && \text{otherwise}\end{aligned}$$

NOTE – The inverse of sin is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arcsin function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| > 1\}$. Thus $\arcsin_{c(F)}(x + \tilde{\imath} \cdot 0) \neq \arcsin_{c(F)}(x + \tilde{\imath} \cdot (-0))$ when $|x| > 1$.

5.3.2.9 Radian arc cosine

The $\operatorname{arccos}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}\operatorname{arccos}_{F \rightarrow c(F)} : F &\rightarrow c(F \cup \{-0\}) \\ \operatorname{arccos}_{F \rightarrow c(F)}(x) &= down_F(\pi) + \tilde{\imath} \cdot neg_F(\operatorname{arccosh}_F(neg_F(x))) && \text{if } (x \in F \text{ and } x < -1) \text{ or } x = -\infty \\ &= \operatorname{arccos}_F(x) + \tilde{\imath} \cdot (-0) && \text{if } (x \in F \text{ and } |x| \leq 1) \text{ or } x = -0 \\ &= 0 + \tilde{\imath} \cdot \operatorname{arccosh}_F(x) && \text{if } (x \in F \text{ and } x > 1) \text{ or } x = +\infty \\ &= no_result_{F \rightarrow c(F)}(x) && \text{otherwise}\end{aligned}$$

The $\operatorname{arccos}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}\operatorname{arccos}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \\ \operatorname{arccos}_{i(F) \rightarrow c(F)}(\tilde{\imath} \cdot y) &= up_F(\pi/2) + \tilde{\imath} \cdot neg_F(\operatorname{arsinh}_F(y)) && \text{if } (y \in F \text{ and } y < 0) \text{ or } y \in \{-\infty, -0\}\end{aligned}$$

$$\begin{aligned}
&= \text{down}_F(\pi/2) + \mathbf{i} \cdot \text{neg}_F(\text{arcsinh}_F(y)) \\
&\quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
&= \text{no_result}_{i(F) \rightarrow c(F)}(\mathbf{i} \cdot y) \text{ otherwise}
\end{aligned}$$

The $\text{arccos}_{c(F)}^*$ approximation helper function:

$$\text{arccos}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{arccos}_{c(F)}^*(z)$ returns a close approximation to $\text{arccos}(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_sin}_{c(F)}$.

Further requirements on the $\text{arccos}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}
\text{arccos}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{arccos}_{c(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F \text{ and } (|\Re(z)| \leq 1 \text{ or } \Im(z) \neq 0) \\
\Im(\text{arccos}_{c(F)}^*(-z)) &= -\Im(\text{arccos}_{c(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F \text{ and } (|\Re(z)| \leq 1 \text{ or } \Im(z) \neq 0)
\end{aligned}$$

The relationships to the arccos_F^* , arccosh_F^* , and arcsinh_F^* approximation helper functions for arccos_F , arccosh_F , and arcsinh_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned}
\text{arccos}_{c(F)}^*(x) &= \pi - (\mathbf{i} \cdot \text{arccosh}_F^*(-x)) \quad \text{if } x \in F \text{ and } x \leq -1 \\
\text{arccos}_{c(F)}^*(x) &= \text{arccos}_F^*(x) \quad \text{if } x \in F \text{ and } |x| < 1 \\
\text{arccos}_{c(F)}^*(x) &= -(\mathbf{i} \cdot \text{arccosh}_F^*(x)) \quad \text{if } x \in F \text{ and } x \geq 1 \\
\text{arccos}_{c(F)}^*(\mathbf{i} \cdot y) &= (\pi/2) - (\mathbf{i} \cdot \text{arcsinh}_F^*(y)) \quad \text{if } y \in F
\end{aligned}$$

The $\text{arccos}_{c(F)}^\#$ range limitation helper function:

$$\begin{aligned}
\text{arccos}_{c(F)}^\#(z) &= \max\{\text{up}_F(\pi/2), \min\{\Re(\text{arccos}_{c(F)}^*(z)), \text{down}_F(\pi)\}\} + (\mathbf{i} \cdot \Im(\text{arccos}_{c(F)}^*(z))) \\
&\quad \text{if } \Re(z) < 0 \\
&= \min\{\Re(\text{arccos}_{c(F)}^*(z)), \text{down}_F(\pi/2)\} + (\mathbf{i} \cdot \Im(\text{arccos}_{c(F)}^*(z))) \\
&\quad \text{if } \Re(z) \geq 0
\end{aligned}$$

The $\text{arccos}_{c(F)}$ operation:

$$\begin{aligned}
&\text{arccos}_{c(F)} : c(F) \rightarrow c(F \cup \{-0\}) \\
&\text{arccos}_{c(F)}(x + \mathbf{i} \cdot y) \\
&\quad = \text{result}_{c(F)}^*(\text{arccos}_{c(F)}^\#(x + (\mathbf{i} \cdot y)), \text{nearest}_F) \\
&\quad \quad \text{if } x, y \in F \text{ and } (y \neq 0 \text{ or } |x| > 1) \\
&\quad = \text{arccos}_F(x) + \mathbf{i} \cdot (-0) \quad \text{if } x \in F \text{ and } y = 0 \text{ and } |x| \leq 1 \\
&\quad = \text{up}_F(\pi/2) + \mathbf{i} \cdot \text{neg}_F(\text{arcsinh}_F(y)) \\
&\quad \quad \text{if } x = -0 \\
&\quad = \text{conj}_{c(F)}(\text{arccos}_{c(F)}(x + \mathbf{i} \cdot 0)) \\
&\quad \quad \text{if } y = -0 \text{ and } x \neq -0 \\
&\quad = \text{arc}_F(x, y) + \mathbf{i} \cdot (-\infty) \quad \text{if } x \in \{-\infty, +\infty\} \text{ and} \\
&\quad \quad ((y \in F \text{ and } y \geq 0) \text{ or } y = +\infty) \\
&\quad = \text{arc}_F(x, \text{neg}_F(y)) + \mathbf{i} \cdot (+\infty) \\
&\quad \quad \text{if } x \in \{-\infty, +\infty\} \text{ and} \\
&\quad \quad ((y \in F \text{ and } y < 0) \text{ or } y = -\infty) \\
&\quad = \text{no_result}_{c(F)}(x + \mathbf{i} \cdot y) \quad \text{otherwise}
\end{aligned}$$

NOTE – The inverse of $\cos$ is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arccos function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| > 1\}$. Thus $\text{arccos}_{c(F)}(x + \mathbf{i} \cdot 0) \neq \text{arccos}_{c(F)}(x + \mathbf{i} \cdot (-0))$ when $|x| > 1$.

5.3.2.10 Radian arc tangent

The $\arctan_{i(F)}$ operation:

$$\arctan_{i(F)} : i(F) \rightarrow i(F) \cup \{\mathbf{infinitary}, \mathbf{invalid}\}$$

$$\arctan_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \arctanh_F(y)$$

The $\arctan_{i(F) \rightarrow c(F)}$ operation:

$$\arctan_{i(F) \rightarrow c(F)} : i(F) \rightarrow c(F \cup \{-0\}) \cup \{\mathbf{infinitary}\}$$

$$\begin{aligned} \arctan_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= up_F(-\pi/2) + \hat{\mathbf{i}} \cdot arccoth_F(y) && \text{if } (y \in F \text{ and } y < -1) \text{ or } y = -\infty \\ &= mul_F(0, y) + \hat{\mathbf{i}} \cdot \arctanh_F(y) && \text{if } (y \in F \text{ and } |y| \leq 1) \text{ or } y = -0 \\ &= down_F(\pi/2) + \hat{\mathbf{i}} \cdot arccoth_F(y) && \text{if } (y \in F \text{ and } y > 1) \text{ or } y = +\infty \\ &= no_result_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) && \text{otherwise} \end{aligned}$$

The $\arctan_{c(F)}^*$ approximation helper function:

$$\arctan_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\arctan_{c(F)}^*(z)$ returns a close approximation to $\arctan(z)$ in $\mathcal{C}$ with maximum error $max_error_tan_{c(F)}$.

Further requirements on the $\arctan_{c(F)}^*$ approximation helper function are:

$$\begin{aligned} \arctan_{c(F)}^*(conj(z)) &= conj(\arctan_{c(F)}^*(z)) && \text{if } z \in \mathcal{C}_F \\ \arctan_{c(F)}^*(-z) &= -\arctan_{c(F)}^*(z) && \text{if } z \in \mathcal{C}_F \text{ and } (\Re(z) < 1 \text{ or } \Im(z) \neq 0) \end{aligned}$$

The relationships to the $\arctan_F^*$, $\arctanh_F^*$, and $arccoth_F^*$ approximation helper functions for $\arctan_F$, $\arctanh_F$, and $arccoth_F$ operations in an associated library for real-valued operations shall be:

$$\begin{aligned} \arctan_{c(F)}^*(x) &= \arctan_F^*(x) && \text{if } x \in F \\ \arctan_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \tilde{\mathbf{i}} \cdot \arctanh_F^*(y) && \text{if } y \in F \text{ and } |y| < 1 \\ \arctan_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= (\pi/2) + (\tilde{\mathbf{i}} \cdot arccoth_F^*(y)) && \text{if } y \in F \text{ and } |y| > 1 \end{aligned}$$

The $\arctan_{c(F)}^\#$ range limitation helper function:

$$\arctan_{c(F)}^\#(z) = \max\{up_F(-\pi/2), \min\{\Re(\arctan_{c(F)}^*(z)), down_F(\pi/2)\}\} + (\tilde{\mathbf{i}} \cdot \Im(\arctan_{c(F)}^*(z)))$$

The $\arctan_{c(F)}$ operation:

$$\arctan_{c(F)} : c(F) \rightarrow c(F) \cup \{\mathbf{underflow}, \mathbf{infinitary}\}$$

$$\begin{aligned} \arctan_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= result_{c(F)}^*(\arctan_{c(F)}^\#(x + (\tilde{\mathbf{i}} \cdot y)), nearest_F) && \text{if } x, y \in F \text{ and } (x \neq 0 \text{ or } |y| \neq 1) \\ &= \mathbf{infinitary}(0 + \hat{\mathbf{i}} \cdot mul_F(y, +\infty)) && \text{if } x = 0 \text{ and } y \in \{-1, 1\} \\ &= neg_{c(F)}(\arctan_{c(F)}(0 + \hat{\mathbf{i}} \cdot neg_F(y))) && \text{if } x = -0 \\ &= conj_{c(F)}(\arctan_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) && \text{if } y = -0 \text{ and } x \neq -0 \end{aligned}$$

$$\begin{aligned}
&= \text{mul}_F(\text{signum}_F(x), \text{down}_F(\pi/2)) + \hat{\mathbf{i}} \cdot \text{mul}_F(\text{signum}_F(y), 0) \\
&\quad \text{if } (x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\}) \text{ or} \\
&\quad \quad (x \in F \text{ and } y \in \{-\infty, +\infty\}) \\
&= \text{no_result}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \quad \text{otherwise}
\end{aligned}$$

NOTE – The inverse of tan is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ (even any integer multiple of π) added to it, and the result is also in the solution set. The arctan function (returning the principal value for the inverse) branch cuts at $\{\tilde{\mathbf{i}} \cdot y \mid y \in F \text{ and } |y| > 1\}$. Thus $\text{arctan}_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) \neq \text{arctan}_{c(F)}((-0) + \hat{\mathbf{i}} \cdot y)$ when $|y| > 1$.

5.3.2.11 Radian arc cotangent

The $\text{arccot}_{i(F)}$ operation:

$$\begin{aligned}
\text{arccot}_{i(F)} : i(F) &\rightarrow i(F) \cup \{\text{underflow, infinitary, invalid}\} \\
\text{arccot}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{arccoth}_F(\text{neg}_F(y))
\end{aligned}$$

The $\text{arccot}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}
\text{arccot}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \cup \{\text{underflow, infinitary}\} \\
\text{arccot}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= \text{mul}_F(\text{signum}_F(y), 0) + \hat{\mathbf{i}} \cdot \text{arccoth}_F(\text{neg}_F(y)) \\
&\quad \text{if } y \in F \text{ and } |y| \geq 1 \text{ or } y \in \{-\infty, +\infty\} \\
&= \text{up}_F(-\pi/2) + \hat{\mathbf{i}} \cdot \text{arctanh}_F(y) \\
&\quad \text{if } (y \in F \text{ and } -1 < y \text{ and } y < 0) \text{ or } y = -0 \\
&= \text{down}_F(\pi/2) + \hat{\mathbf{i}} \cdot \text{arctanh}_F(y) \\
&\quad \text{if } y \in F \text{ and } 0 \leq y \text{ and } y < 1 \\
&= \text{no_result}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \quad \text{otherwise}
\end{aligned}$$

The $\text{arccot}_{c(F)}^*$ approximation helper function:

$$\text{arccot}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{arccot}_{c(F)}^*(z)$ returns a close approximation to $\text{arccot}(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_{c(F)}$.

Further requirements on the $\text{arccot}_{c(F)}^*$ approximation helper function are:

$$\begin{aligned}
\text{arccot}_{c(F)}^*(\text{conj}(z)) &= \text{conj}(\text{arccot}_{c(F)}^*(z)) & \text{if } z \in \mathcal{C}_F \text{ and } (\Re(z) \neq 0 \text{ or } |\Im(z)| > 1) \\
\text{arccot}_{c(F)}^*(-z) &= -\text{arccot}_{c(F)}^*(z) & \text{if } z \in \mathcal{C}_F \text{ and } (\Re(z) \neq 0 \text{ or } |\Im(z)| > 1) \\
\Re(\text{arccot}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y)) &= \pi/2 & \text{if } y \in F \text{ and } |y| < 1
\end{aligned}$$

The relationships to the arccot_F^* , arccoth_F^* , and arctanh_F^* approximation helper functions for arccot_F , arccoth_F , and arctanh_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned}
\text{arccot}_{c(F)}^*(x) &= \text{arccot}_F^*(x) & \text{if } x \in F \\
\text{arccot}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= \tilde{\mathbf{i}} \cdot \text{arccoth}_F^*(-y) & \text{if } y \in F \text{ and } |y| > 1 \\
\text{arccot}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= (-\pi/2) + (\tilde{\mathbf{i}} \cdot \text{arctanh}_F^*(y)) & \text{if } y \in F \text{ and } |y| < 1
\end{aligned}$$

The $\text{arccot}_{c(F)}^\#$ range limitation helper function:

$$\text{arccot}_{c(F)}^\#(z) = \max\{\text{up}_F(-\pi/2), \min\{\Re(\text{arccot}_{c(F)}^*(z)), \text{down}_F(\pi/2)\}\} + (\tilde{\mathbf{i}} \cdot \Im(\text{arccot}_{c(F)}^*(z)))$$

The $\text{arccot}_{c(F)}$ operation:

$$\begin{aligned}
& \text{arccot}_{c(F)} : c(F) \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow, infinitary}\} \\
& \text{arccot}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\
& \quad = \text{result}_{c(F)}^*(\text{arccot}_{c(F)}^\#(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\
& \quad \quad \text{if } x, y \in F \text{ and } (|y| \neq 1 \text{ or } x \neq 0) \text{ and } y \neq 0 \\
& \quad = \text{arccot}_F(x) + \hat{\mathbf{i}} \cdot (-0) \quad \text{if } x \in F \cup \{-\infty, -0, +\infty\} \text{ and } y = 0 \\
& \quad = \text{neg}_{c(F)}(\text{arccot}_{c(F)}(0 + \hat{\mathbf{i}} \cdot \text{neg}_F(y))) \\
& \quad \quad \text{if } x = -0 \\
& \quad = \text{conj}_{c(F)}(\text{arccot}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0)) \\
& \quad \quad \text{if } y = -0 \text{ and } x \neq -0 \\
& \quad = \text{mul}_F(\text{signum}_F(x), 0) + \hat{\mathbf{i}} \cdot \text{mul}_F(\text{signum}_F(y), 0) \\
& \quad \quad \text{if } (x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\}) \text{ or} \\
& \quad \quad \quad (x \in F \text{ and } y \in \{-\infty, +\infty\}) \\
& \quad = \text{infinitary}(0 + \hat{\mathbf{i}} \cdot \text{mul}_F(y, -\infty)) \\
& \quad \quad \text{if } x = 0 \text{ and } y \in \{-1, 1\} \\
& \quad = \text{no_result}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \quad \text{otherwise}
\end{aligned}$$

NOTE – The inverse of cot is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ (even any integer multiple of π) added to it, and the result is also in the solution set. The arccot function (returning the principal value for the inverse) branch cuts at $\{\hat{\mathbf{i}} \cdot y \mid y \in \mathcal{R} \text{ and } |y| < 1\}$. Thus $\text{arccot}_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) \neq \text{arccot}_{c(F)}((-0) + \hat{\mathbf{i}} \cdot y)$ when $|y| < 1$ or $y = -0$.

5.3.2.12 Radian arc secant

The $\text{arcsec}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \text{arcsec}_{F \rightarrow c(F)} : F \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow, infinitary}\} \\
& \text{arcsec}_{F \rightarrow c(F)}(x) = \text{arcsec}_F(x) + \hat{\mathbf{i}} \cdot \text{mul}_F(\text{signum}_F(x), -0) \\
& \quad \quad \text{if } (x \in F \text{ and } |x| \geq 1) \text{ or } x \in \{-\infty, +\infty\} \\
& \quad = 0 + \hat{\mathbf{i}} \cdot \text{neg}_F(\text{arcsech}_F(x)) \\
& \quad \quad \text{if } (x \in F \text{ and } -1 < x < 0) \text{ or } x = -0 \\
& \quad = \text{down}_F(\pi) + \hat{\mathbf{i}} \cdot \text{arcsech}_F(\text{neg}_F(x)) \\
& \quad \quad \text{if } x \in F \text{ and } 0 \leq x < 1 \\
& \quad = \text{no_result}_{F \rightarrow c(F)}(x) \quad \text{otherwise}
\end{aligned}$$

The $\text{arcsec}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \text{arcsec}_{i(F) \rightarrow c(F)} : i(F) \rightarrow c(F) \cup \{\text{underflow, infinitary}\} \\
& \text{arcsec}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \\
& \quad = \text{down}_F(\pi/2) + \hat{\mathbf{i}} \cdot \text{arccsch}_F(y) \\
& \quad \quad \text{if } (y \in F \text{ and } y < 0) \text{ or } y \in \{-\infty, -0\} \\
& \quad = \text{up}_F(\pi/2) + \hat{\mathbf{i}} \cdot \text{arccsch}_F(y) \\
& \quad \quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
& \quad = \text{no_result}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \quad \text{otherwise}
\end{aligned}$$

The $\text{arcsec}_{c(F)}^*$ approximation helper function:

$$\text{arcsec}_{c(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\text{arcsec}_{c(F)}^*(z)$ returns a close approximation to $\text{arcsec}(z)$ in $\mathcal{C}$ with maximum error $\text{max_error_tan}_{c(F)}$.

Further requirements on the $\text{arcsec}_{c(F)}^*$ approximation helper function are:

$$\text{arcsec}_{c(F)}^*(\text{conj}(z)) = \text{conj}(\text{arcsec}_{c(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F \text{ and } (\Im(z) \neq 0 \text{ or } |\Re(z)| \geq 1)$$

The relationships to the arcsec_F^* , arcsech_F^* , and arccsch_F^* approximation helper functions for arcsec_F , arcsech_F , and arccsch_F operations in an associated library for real-valued operations shall be:

$$\begin{aligned}\text{arcsec}_{c(F)}^*(x) &= \text{arcsec}_F^*(x) && \text{if } x \in F \text{ and } |x| \geq 1 \\ \text{arcsec}_{c(F)}^*(x) &= -\tilde{\mathbf{i}} \cdot \text{arcsech}_F^*(x) && \text{if } x \in F \text{ and } |x| < 1 \text{ and } x \neq 0 \\ \text{arcsec}_{c(F)}^*(\tilde{\mathbf{i}} \cdot y) &= (\pi/2) + (\tilde{\mathbf{i}} \cdot \text{arccsch}_F^*(y)) && \text{if } y \in F\end{aligned}$$

The $\text{arcsec}_{c(F)}^\#$ range limitation helper function:

$$\begin{aligned}\text{arcsec}_{c(F)}^\#(z) &= \min\{\Re(\text{arcsec}_{c(F)}^*(z)), \text{down}_F(\pi/2)\} + (\tilde{\mathbf{i}} \cdot \Im(\text{arcsec}_{c(F)}^*(z))) \\ &\quad \text{if } \Re(z) \geq 0 \\ &= \max\{\text{up}_F(\pi/2), \min\{\Re(\text{arcsec}_{c(F)}^*(z)), \text{down}_F(\pi)\}\} + (\tilde{\mathbf{i}} \cdot \Im(\text{arcsec}_{c(F)}^*(z))) \\ &\quad \text{if } \Re(z) < 0\end{aligned}$$

The $\text{arcsec}_{c(F)}$ operation:

$$\begin{aligned}\text{arcsec}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{infinitary}\} \\ \text{arcsec}_{c(F)}(x + \tilde{\mathbf{i}} \cdot y) &= \text{result}_{c(F)}^*(\text{arcsec}_{c(F)}^\#(x + (\tilde{\mathbf{i}} \cdot y)), \text{nearest}_F) \\ &\quad \text{if } x, y \in F \text{ and } (x \neq 0 \text{ or } y \neq 0) \\ &= \text{infinitary}(\text{down}(\pi/2) + \tilde{\mathbf{i}} \cdot (+\infty)) \\ &\quad \text{if } y = 0 \text{ and } x = 0 \\ &= \text{arcsec}_{c(F)}(0 + \tilde{\mathbf{i}} \cdot \text{neg}_F(y)) \\ &\quad \text{if } x = -0 \\ &= \text{conj}_{c(F)}(\text{arcsec}_{c(F)}(x + \tilde{\mathbf{i}} \cdot 0)) \\ &\quad \text{if } y = -0 \text{ and } x \neq -0 \\ &= \text{mul}_F(\text{signum}_F(x), \text{down}_F(\pi/2)) + \tilde{\mathbf{i}} \cdot \text{mul}_F(\text{signum}_F(y), 0) \\ &\quad \text{if } (x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\}) \text{ or} \\ &\quad \quad (y \in \{-\infty, +\infty\} \text{ and } x \in F) \\ &= \text{no_result}_{c(F)}(x + \tilde{\mathbf{i}} \cdot y) \quad \text{otherwise}\end{aligned}$$

NOTE – The inverse of sec is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arcsec function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| < 1\}$. Thus $\text{arcsec}_{c(F)}(x + \tilde{\mathbf{i}} \cdot 0) \neq \text{arcsec}_{c(F)}(x + \tilde{\mathbf{i}} \cdot (-0))$ when $|x| < 1$ or $x = -0$.

5.3.2.13 Radian arc cosecant

The $\text{arccsc}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}\text{arccsc}_{F \rightarrow c(F)} : F &\rightarrow c(F \cup \{-0\}) \cup \{\text{underflow}, \text{infinitary}\} \\ \text{arccsc}_{F \rightarrow c(F)}(x) &= \text{arccsc}_F(x) + \tilde{\mathbf{i}} \cdot \text{mul}_F(\text{signum}_F(x), 0) \\ &\quad \text{if } (x \in F \text{ and } |x| \geq 1) \text{ or } x \in \{-\infty, +\infty\} \\ &= \text{up}_F(-\pi/2) + \tilde{\mathbf{i}} \cdot \text{neg}_F(\text{arcsech}_F(\text{neg}_F(x))) \\ &\quad \text{if } (x \in F \text{ and } -1 < x < 0) \text{ or } x = -0 \\ &= \text{down}_F(\pi/2) + \tilde{\mathbf{i}} \cdot \text{arcsech}_F(x) \\ &\quad \text{if } x \in F \text{ and } 0 \leq x < 1 \\ &= \text{no_result}_{F \rightarrow c(F)}(x) \quad \text{otherwise}\end{aligned}$$

The $\text{arccsc}_{i(F)}$ operation:

$$\operatorname{arccsc}_{\mathbf{i}(F)} : \mathbf{i}(F) \rightarrow \mathbf{i}(F) \cup \{\mathbf{underflow}, \mathbf{infinitary}\}$$

$$\operatorname{arccsc}_{\mathbf{i}(F)}(\mathbf{i} \cdot y) = \mathbf{i} \cdot \operatorname{arccsch}_F(\operatorname{neg}_F(y))$$

The $\operatorname{arccsc}_{\mathbf{c}(F)}^*$ approximation helper function:

$$\operatorname{arccsc}_{\mathbf{c}(F)}^* : \mathcal{C}_F \rightarrow \mathcal{C}$$

$\operatorname{arccsc}_{\mathbf{c}(F)}^*(z)$ returns a close approximation to $\operatorname{arccsc}(z)$ in $\mathcal{C}$ with maximum error $\max_error_tan_{\mathbf{c}(F)}$.

Further requirements on the $\operatorname{arccsc}_{\mathbf{c}(F)}^*$ approximation helper function are:

$$\operatorname{arccsc}_{\mathbf{c}(F)}^*(\operatorname{conj}(z)) = \operatorname{conj}(\operatorname{arccsc}_{\mathbf{c}(F)}^*(z)) \quad \text{if } z \in \mathcal{C}_F \text{ and } (\Im(z) \neq 0 \text{ or } |\Re(z)| \geq 0)$$

$$\operatorname{arccsc}_{\mathbf{c}(F)}^*(-z) = -\operatorname{arccsc}_{\mathbf{c}(F)}^*(z) \quad \text{if } z \in \mathcal{C}_F \text{ and } (\Im(z) \neq 0 \text{ or } |\Re(z)| \geq 0)$$

The relationships to the $\operatorname{arccsc}_F^*$, $\operatorname{arccsch}_F^*$, and $\operatorname{arcsec}_F^*$ approximation helper functions for arccsc_F , $\operatorname{arccsch}_F$, and arcsec_F operations in an associated library for real-valued operations shall be:

$$\operatorname{arccsc}_{\mathbf{c}(F)}^*(x) = \operatorname{arccsc}_F^*(x) \quad \text{if } x \in F \text{ and } |x| \geq 1$$

$$\operatorname{arccsc}_{\mathbf{c}(F)}^*(x) = (\pi/2) + (\mathbf{i} \cdot \operatorname{arcsec}_F^*(x)) \quad \text{if } x \in F \text{ and } |x| < 1$$

$$\operatorname{arccsc}_{\mathbf{c}(F)}^*(\mathbf{i} \cdot y) = \mathbf{i} \cdot \operatorname{arccsch}_F^*(-y) \quad \text{if } y \in F$$

The $\operatorname{arccsc}_{\mathbf{c}(F)}^\#$ range limitation helper function:

$$\operatorname{arccsc}_{\mathbf{c}(F)}^\#(z) = \max\{up_F(-\pi/2), \min\{\Re(\operatorname{arccsc}_{\mathbf{c}(F)}^*(z)), down_F(\pi/2)\}\} + (\mathbf{i} \cdot \Im(\operatorname{arccsc}_{\mathbf{c}(F)}^*(z)))$$

The $\operatorname{arccsc}_{\mathbf{c}(F)}$ operation:

$$\operatorname{arccsc}_{\mathbf{c}(F)} : \mathbf{c}(F) \rightarrow \mathbf{c}(F \cup \{-\mathbf{0}\}) \cup \{\mathbf{underflow}, \mathbf{infinitary}\}$$

$$\begin{aligned} \operatorname{arccsc}_{\mathbf{c}(F)}(x + \mathbf{i} \cdot y) &= \operatorname{result}_{\mathbf{c}(F)}^*(\operatorname{arccsc}_{\mathbf{c}(F)}^\#(x + (\mathbf{i} \cdot y)), nearest_F) \\ &\quad \text{if } x, y \in F \text{ and } (y \neq 0 \text{ or } 0 < |x| < 1) \\ &= \operatorname{arccsc}_F(x) + \mathbf{i} \cdot (-\mathbf{0}) \quad \text{if } x \in F \text{ and } y = 0 \text{ and } |x| \geq 1 \\ &= \mathbf{infinitary}(0 + \mathbf{i} \cdot (-\infty)) \quad \text{if } x = 0 \text{ and } y = 0 \\ &= \operatorname{neg}_{\mathbf{c}(F)}(\operatorname{arccsc}_{\mathbf{c}(F)}(0 + \mathbf{i} \cdot \operatorname{neg}_F(y))) \quad \text{if } x = -\mathbf{0} \\ &= \operatorname{conj}_{\mathbf{c}(F)}(\operatorname{arccsc}_{\mathbf{c}(F)}(x + \mathbf{i} \cdot 0)) \quad \text{if } y = -\mathbf{0} \text{ and } x \neq -\mathbf{0} \\ &= \operatorname{mul}_F(\operatorname{signum}_F(x), 0) + \mathbf{i} \cdot \operatorname{mul}_F(\operatorname{signum}_F(y), -\mathbf{0}) \quad \text{if } (x \in \{-\infty, +\infty\} \text{ and } y \in F \cup \{-\infty, +\infty\}) \text{ or} \\ &\quad y \in \{-\infty, +\infty\} \text{ and } x \in F \\ &= no_result_{\mathbf{c}(F)}(x + \mathbf{i} \cdot y) \quad \text{otherwise} \end{aligned}$$

NOTE – The inverse of csc is multi-valued, the real part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arccsc function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| < 1\}$. Thus $\operatorname{arccsc}_{\mathbf{c}(F)}(x + \mathbf{i} \cdot 0) \neq \operatorname{arccsc}_{\mathbf{c}(F)}(x + \mathbf{i} \cdot (-\mathbf{0}))$ when $|x| < 1$ or $x = -\mathbf{0}$.

5.3.3 Operations for hyperbolic elementary functions

NOTE – The *correspondences* specified below to other ISO/IEC 10967 operations are exact, not approximate.

5.3.3.1 Hyperbolic normalisation

$$\text{radh}_F : F \rightarrow F$$

$$\text{radh}_F(x) = x$$

$$\text{radh}_{i(F)} : i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\}$$

$$\text{radh}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{rad}_F(y)$$

$$\text{radh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\}$$

$$\begin{aligned} \text{radh}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\ = x \mathrel{+} \hat{\mathbf{i}} \cdot \text{rad}_F(y) \end{aligned}$$

5.3.3.2 Hyperbolic sine

The $\text{sinh}_{i(F)}$ operation:

$$\text{sinh}_{i(F)} : i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{absolute_precision_underflow}\}$$

$$\text{sinh}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{sin}_F(y)$$

The $\text{sinh}_{c(F)}$ operation:

$$\text{sinh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\}$$

$$\begin{aligned} \text{sinh}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\ = \text{itimes}_{c(F)}(\text{sin}_{c(F)}(y \mathrel{+} \hat{\mathbf{i}} \cdot \text{neg}_F(x))) \end{aligned}$$

5.3.3.3 Hyperbolic cosine

The $\text{cosh}_{i(F)}$ operation:

$$\text{cosh}_{i(F)} : i(F) \rightarrow F \cup \{\text{underflow}, \text{absolute_precision_underflow}\}$$

$$\text{cosh}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \text{cos}_F(y)$$

The $\text{cosh}_{c(F)}$ operation:

$$\text{cosh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\}$$

$$\begin{aligned} \text{cosh}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\ = \text{cos}_{c(F)}(y \mathrel{+} \hat{\mathbf{i}} \cdot \text{neg}_F(x)) \end{aligned}$$

5.3.3.4 Hyperbolic tangent

The $\text{tanh}_{i(F)}$ operation:

$$\text{tanh}_{i(F)} : i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\}$$

$$\text{tanh}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{tan}_F(y)$$

The $\text{tanh}_{c(F)}$ operation:

$$\text{tanh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\}$$

$$\begin{aligned} \text{tanh}_{c(F)}(x \mathrel{+} \hat{\mathbf{i}} \cdot y) \\ = \text{itimes}_{c(F)}(\text{tan}_{c(F)}(y \mathrel{+} \hat{\mathbf{i}} \cdot \text{neg}_F(x))) \end{aligned}$$

5.3.3.5 Hyperbolic cotangent

The $\text{coth}_{i(F)}$ operation:

$$\begin{aligned}\text{coth}_{i(F)} : i(F) &\rightarrow i(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{absolute_precision_underflow}\} \\ \text{coth}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{cot}_F(\text{neg}_F(y))\end{aligned}$$

The $\text{coth}_{c(F)}$ operation:

$$\begin{aligned}\text{coth}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{absolute_precision_underflow}\} \\ \text{coth}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{itimes}_{c(F)}(\text{cot}_{c(F)}(\text{neg}_F(y) + \hat{\mathbf{i}} \cdot x))\end{aligned}$$

5.3.3.6 Hyperbolic secant

The $\text{sech}_{i(F)}$ operation:

$$\begin{aligned}\text{sech}_{i(F)} : i(F) &\rightarrow F \cup \{\text{overflow}, \text{absolute_precision_underflow}\} \\ \text{sech}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \text{sec}_F(\text{neg}_F(y))\end{aligned}$$

The $\text{sech}_{c(F)}$ operation:

$$\begin{aligned}\text{sech}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{absolute_precision_underflow}\} \\ \text{sech}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{sec}_{c(F)}(\text{neg}_F(y) + \hat{\mathbf{i}} \cdot x)\end{aligned}$$

5.3.3.7 Hyperbolic cosecant

The $\text{csch}_{i(F)}$ operation:

$$\begin{aligned}\text{csch}_{i(F)} : i(F) &\rightarrow i(F) \cup \{\text{overflow}, \text{infinitary}, \text{absolute_precision_underflow}\} \\ \text{csch}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{csc}_F(\text{neg}_F(y))\end{aligned}$$

The $\text{csch}_{c(F)}$ operation:

$$\begin{aligned}\text{csch}_{c(F)} : c(F) &\rightarrow c(F) \cup \{\text{underflow}, \text{overflow}, \text{infinitary}, \text{absolute_precision_underflow}\} \\ \text{csch}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) &= \text{itimes}_{c(F)}(\text{csc}_{c(F)}(\text{neg}_F(y) + \hat{\mathbf{i}} \cdot x))\end{aligned}$$

5.3.3.8 Inverse hyperbolic sine

The $\text{arcsinh}_{i(F)}$ operation:

$$\begin{aligned}\text{arcsinh}_{i(F)} : i(F) &\rightarrow i(F) \cup \{\text{invalid}\} \\ \text{arcsinh}_{i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot \text{arcsin}_F(y)\end{aligned}$$

The $\text{arcsinh}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}\text{arcsinh}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F) \\ \text{arcsinh}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) &= \text{itimes}_{i(F)}(\text{arcsin}_{F \rightarrow c(F)}(y))\end{aligned}$$

The $\text{arcsinh}_{c(F)}$ operation:

$$\begin{aligned}
& \operatorname{arcsinh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\mathbf{underflow}\} \\
& \operatorname{arcsinh}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\
& \quad = \operatorname{itimes}_{c(F)}(\operatorname{arcsin}_{c(F)}(y + \hat{\mathbf{i}} \cdot \operatorname{neg}_F(x)))
\end{aligned}$$

NOTE – The inverse of sinh is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arcsinh function (returning the principal value for the inverse) branch cuts at $\{\hat{\mathbf{i}} \cdot y \mid y \in \mathcal{R} \text{ and } |y| > 1\}$. Thus $\operatorname{arcsinh}_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) \neq \operatorname{arcsinh}_{c(F)}(-0 + \hat{\mathbf{i}} \cdot y)$ when $|y| > 1$.

5.3.3.9 Inverse hyperbolic cosine

The $\operatorname{arccosh}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \operatorname{arccosh}_{F \rightarrow c(F)} : F \rightarrow c(F) \cup \{-0\} \\
& \operatorname{arccosh}_{F \rightarrow c(F)}(x) \\
& \quad = \operatorname{itimes}_{c(F)}(\operatorname{arccos}_{F \rightarrow c(F)}(x)) \\
& \quad \quad \text{if } (x \in F \text{ and } x < 0) \text{ or } x \in \{-\infty, -0\} \\
& \quad = \operatorname{neg}_{c(F)}(\operatorname{itimes}_{c(F)}(\operatorname{arccos}_{F \rightarrow c(F)}(x))) \\
& \quad \quad \text{otherwise}
\end{aligned}$$

The $\operatorname{arccosh}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \operatorname{arccosh}_{i(F) \rightarrow c(F)} : i(F) \rightarrow c(F) \\
& \operatorname{arccosh}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \\
& \quad = \operatorname{itimes}_{c(F)}(\operatorname{arccos}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y)) \\
& \quad \quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
& \quad = \operatorname{neg}_{c(F)}(\operatorname{itimes}_{c(F)}(\operatorname{arccos}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y))) \\
& \quad \quad \text{otherwise}
\end{aligned}$$

The $\operatorname{arccosh}_{c(F)}$ operation:

$$\begin{aligned}
& \operatorname{arccosh}_{c(F)} : c(F) \rightarrow c(F) \\
& \operatorname{arccosh}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\
& \quad = \operatorname{itimes}_{c(F)}(\operatorname{arccos}_{c(F)}(x + \hat{\mathbf{i}} \cdot y)) \\
& \quad \quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
& \quad = \operatorname{neg}_{c(F)}(\operatorname{itimes}_{c(F)}(\operatorname{arccos}_{c(F)}(x + \hat{\mathbf{i}} \cdot y))) \\
& \quad \quad \text{otherwise}
\end{aligned}$$

NOTE – The inverse of cosh is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arccosh function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } x < 1\}$. Thus $\operatorname{arccosh}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq \operatorname{arccosh}_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $x < 1$ or $x = -0$.

5.3.3.10 Inverse hyperbolic tangent

The $\operatorname{arctanh}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \operatorname{arctanh}_{F \rightarrow c(F)} : F \rightarrow c(F) \cup \{-0\} \cup \{\mathbf{underflow}, \mathbf{infinitary}\} \\
& \operatorname{arctanh}_{F \rightarrow c(F)}(x) \\
& \quad = \hat{\mathbf{i}} \cdot \operatorname{arctan}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot \operatorname{neg}_F(x))
\end{aligned}$$

The $\operatorname{arctanh}_{i(F)}$ operation:

$$\operatorname{arctanh}_{i(F)} : i(F) \rightarrow i(F)$$

$$\operatorname{arctanh}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \operatorname{arctan}_F(y)$$

The $\operatorname{arctanh}_{c(F)}$ operation:

$$\operatorname{arctanh}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow, infinitary}\}$$

$$\begin{aligned} \operatorname{arctanh}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\ = \operatorname{itimes}_{c(F)}(\operatorname{arctan}_{c(F)}(y + \hat{\mathbf{i}} \cdot \operatorname{neg}_F(x))) \end{aligned}$$

NOTE – The inverse of $\tanh$ is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ (even any integer multiple of π) added to it, and the result is also in the solution set. The $\operatorname{arctanh}$ function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| > 1\}$. Thus $\operatorname{arctanh}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq \operatorname{arctanh}_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $|x| > 1$.

5.3.3.11 Inverse hyperbolic cotangent

The $\operatorname{arccoth}_{F \rightarrow c(F)}$ operation:

$$\operatorname{arccoth}_{F \rightarrow c(F)} : F \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow, infinitary}\}$$

$$\begin{aligned} \operatorname{arccoth}_{F \rightarrow c(F)}(x) \\ = \hat{\mathbf{i}} \cdot \operatorname{arccot}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot x) \end{aligned}$$

The $\operatorname{arccoth}_{i(F)}$ operation:

$$\operatorname{arccoth}_{i(F)} : i(F) \rightarrow i(F) \cup \{\text{underflow}\}$$

$$\operatorname{arccoth}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \operatorname{arccot}_F(\operatorname{neg}_F(y))$$

The $\operatorname{arccoth}_{c(F)}$ operation:

$$\operatorname{arccoth}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow, infinitary}\}$$

$$\begin{aligned} \operatorname{arccoth}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\ = \operatorname{itimes}_{c(F)}(\operatorname{arccot}_{c(F)}(\operatorname{neg}_F(y) + \hat{\mathbf{i}} \cdot x)) \end{aligned}$$

NOTE – The inverse of $\coth$ is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ (even any integer multiple of π) added to it, and the result is also in the solution set. The $\operatorname{arccoth}$ function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } |x| < 1\}$. Thus $\operatorname{arccoth}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq \operatorname{arccoth}_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $|x| < 1$ or $x = -0$.

5.3.3.12 Inverse hyperbolic secant

The $\operatorname{arcsech}_{F \rightarrow c(F)}$ operation:

$$\operatorname{arcsech}_{F \rightarrow c(F)} : F \rightarrow c(F \cup \{-0\}) \cup \{\text{underflow, infinitary}\}$$

$$\begin{aligned} \operatorname{arcsech}_{F \rightarrow c(F)}(x) \\ = \operatorname{itimes}_{c(F)}(\operatorname{arcsec}_{F \rightarrow c(F)}(x)) \\ \quad \text{if } (x \in F \text{ and } x < 0) \text{ or } x \in \{-\infty, -0\} \\ = \operatorname{neg}_{c(F)}(\operatorname{itimes}_{c(F)}(\operatorname{arcsec}_{F \rightarrow c(F)}(x))) \\ \quad \text{otherwise} \end{aligned}$$

The $\operatorname{arcsech}_{i(F) \rightarrow c(F)}$ operation:

$$\operatorname{arcsech}_{i(F) \rightarrow c(F)} : i(F) \rightarrow c(F) \cup \{\text{underflow, infinitary}\}$$

$$\begin{aligned}
& \text{arcsech}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \\
&= \text{itimes}_{c(F)}(\text{arcsec}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y)) \\
&\quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
&= \text{neg}_{c(F)}(\text{itimes}_{c(F)}(\text{arcsec}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y))) \\
&\quad \text{otherwise}
\end{aligned}$$

The $\text{arcsech}_{c(F)}$ operation:

$$\begin{aligned}
& \text{arcsech}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{infinitary}\} \\
& \text{arcsech}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\
&= \text{itimes}_{c(F)}(\text{arcsec}_{c(F)}(x + \hat{\mathbf{i}} \cdot y)) \\
&\quad \text{if } (y \in F \text{ and } y \geq 0) \text{ or } y = +\infty \\
&= \text{neg}_{c(F)}(\text{itimes}_{c(F)}(\text{arcsec}_{c(F)}(x + \hat{\mathbf{i}} \cdot y))) \\
&\quad \text{otherwise}
\end{aligned}$$

NOTE – The inverse of sech is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arcsech function (returning the principal value for the inverse) branch cuts at $\{x \mid x \in \mathcal{R} \text{ and } x \leq 0 \text{ or } x > 1\}$. Thus $\text{arcsech}_{c(F)}(x + \hat{\mathbf{i}} \cdot 0) \neq \text{arcsech}_{c(F)}(x + \hat{\mathbf{i}} \cdot (-0))$ when $x \leq 0$ or $x = -0$ or $x > 1$.

5.3.3.13 Inverse hyperbolic cosecant

The $\text{arccsch}_{i(F)}$ operation:

$$\begin{aligned}
& \text{arccsch}_{i(F)} : i(F) \rightarrow i(F) \cup \{\text{underflow}, \text{invalid}\} \\
& \text{arccsch}_{i(F)}(\hat{\mathbf{i}} \cdot y) = \hat{\mathbf{i}} \cdot \text{arccsc}_F(\text{neg}_F(y))
\end{aligned}$$

The $\text{arccsch}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned}
& \text{arccsch}_{i(F) \rightarrow c(F)} : i(F) \rightarrow c(F) \cup \{\text{underflow}, \text{infinitary}\} \\
& \text{arccsch}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \\
&= \text{itimes}_{c(F)}(\text{arccsc}_{F \rightarrow c(F)}(y))
\end{aligned}$$

The $\text{arccsch}_{c(F)}$ operation:

$$\begin{aligned}
& \text{arccsch}_{c(F)} : c(F) \rightarrow c(F) \cup \{\text{underflow}, \text{infinitary}\} \\
& \text{arccsch}_{c(F)}(x + \hat{\mathbf{i}} \cdot y) \\
&= \text{itimes}_{c(F)}(\text{arccsc}_{c(F)}(\text{neg}_F(y) + \hat{\mathbf{i}} \cdot x))
\end{aligned}$$

NOTE – The inverse of csch is multi-valued, the imaginary part may have any integer multiple of $2 \cdot \pi$ added to it, and the result is also in the solution set. The arccsch function (returning the principal value for the inverse) branch cuts at $\{\hat{\mathbf{i}} \cdot y \mid y \in \mathcal{R} \text{ and } |y| < 1\}$. Thus $\text{arccsch}_{c(F)}(0 + \hat{\mathbf{i}} \cdot y) \neq \text{arccsch}_{c(F)}(-0 + \hat{\mathbf{i}} \cdot y)$ when $|y| < 1$ or $y = -0$.

5.4 Operations for conversion between imaginary and complex numeric datatypes

5.4.1 Integer to complex integer conversions

Let I and I' be the non-special value sets for two integer datatypes, at least one of which of those occurring in each operation's signature expression conforms to part 1 of ISO/IEC 10967.

$$\text{convert}_{I \rightarrow c(I)} : I \rightarrow c(I)$$

$$\begin{aligned} \text{convert}_{I \rightarrow c(I)}(x) \\ = x + \hat{\mathbf{i}} \cdot 0 \end{aligned}$$

$$\begin{aligned} \text{convert}_{i(I) \rightarrow c(I)} : i(I) &\rightarrow c(I) \\ \text{convert}_{i(I) \rightarrow c(I)}(\hat{\mathbf{i}} \cdot y) \\ &= 0 + \hat{\mathbf{i}} \cdot y \end{aligned}$$

$$\begin{aligned} \text{convert}_{i(I) \rightarrow i(I')} : i(I) &\rightarrow i(I') \cup \{\mathbf{overflow}\} \\ \text{convert}_{i(I) \rightarrow i(I')}(\hat{\mathbf{i}} \cdot y) \\ &= \hat{\mathbf{i}} \cdot \text{convert}_{I \rightarrow I'}(y) \end{aligned}$$

$$\begin{aligned} \text{convert}_{c(I) \rightarrow c(I')} : c(I) &\rightarrow c(I') \cup \{\mathbf{overflow}\} \\ \text{convert}_{c(I) \rightarrow c(I')}(x + \hat{\mathbf{i}} \cdot y) \\ &= \text{convert}_{I \rightarrow I'}(x) + \hat{\mathbf{i}} \cdot \text{convert}_{I \rightarrow I'}(y) \end{aligned}$$

5.4.2 Floating point to complex floating point conversions

Let F and F' be the non-special value sets for two floating point datatypes, at least one of which of those occurring in each operation's signature expression conforms to part 1. Let D be the non-special value set for a fixed point datatype (see Clause 5.4.5 in part 2).

The $\text{convert}_{F \rightarrow c(F)}$ operation:

$$\begin{aligned} \text{convert}_{F \rightarrow c(F)} : F &\rightarrow c(F \cup \{-0\}) \\ \text{convert}_{F \rightarrow c(F)}(x) \\ &= x + \hat{\mathbf{i}} \cdot \text{im}_F(x) && \text{if } x \in F \cup \{-\infty, -0, +\infty\} \\ &= \text{no_result}_{F \rightarrow c(F)}(x) && \text{otherwise} \end{aligned}$$

The $\text{convert}_{i(F) \rightarrow c(F)}$ operation:

$$\begin{aligned} \text{convert}_{i(F) \rightarrow c(F)} : i(F) &\rightarrow c(F \cup \{-0\}) \\ \text{convert}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) \\ &= \text{re}_{i(F)}(\hat{\mathbf{i}} \cdot y) + \hat{\mathbf{i}} \cdot y && \text{if } y \in F \cup \{-\infty, -0, +\infty\} \\ &= \text{no_result}_{i(F) \rightarrow c(F)}(\hat{\mathbf{i}} \cdot y) && \text{otherwise} \end{aligned}$$

The $\text{convert}_{i(F) \rightarrow i(F')}$ operation:

$$\begin{aligned} \text{convert}_{i(F) \rightarrow i(F')} : i(F) &\rightarrow i(F') \cup \{\mathbf{underflow}, \mathbf{overflow}\} \\ \text{convert}_{i(F) \rightarrow i(F')}(\hat{\mathbf{i}} \cdot y) \\ &= \hat{\mathbf{i}} \cdot \text{convert}_{F \rightarrow F'}(y) \end{aligned}$$

The $\text{convert}_{c(F) \rightarrow c(F')}$ operation:

$$\begin{aligned} \text{convert}_{c(F) \rightarrow c(F')} : c(F) &\rightarrow c(F') \cup \{\mathbf{underflow}, \mathbf{overflow}\} \\ \text{convert}_{c(F) \rightarrow c(F')}(x + \hat{\mathbf{i}} \cdot y) \\ &= \text{convert}_{F \rightarrow F'}(x) + \hat{\mathbf{i}} \cdot \text{convert}_{F \rightarrow F'}(y) \end{aligned}$$

The $convert_{i(F) \rightarrow i(D)}$ operation:

$$\begin{aligned} convert_{i(F) \rightarrow i(D)} : i(F) &\rightarrow i(D) \cup \{\mathbf{overflow}\} \\ convert_{i(F) \rightarrow i(D)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot convert_{F \rightarrow D}(y) \end{aligned}$$

The $convert_{c(F) \rightarrow c(D)}$ operation:

$$\begin{aligned} convert_{c(F) \rightarrow c(D)} : c(F) &\rightarrow c(D) \cup \{\mathbf{overflow}\} \\ convert_{c(F) \rightarrow c(D)}(x + \hat{\mathbf{i}} \cdot y) &= convert_{F \rightarrow D}(x) + \hat{\mathbf{i}} \cdot convert_{F \rightarrow D}(y) \end{aligned}$$

The $convert_{i(D) \rightarrow i(F)}$ operation:

$$\begin{aligned} convert_{i(D) \rightarrow i(F)} : i(D) &\rightarrow i(F) \cup \{\mathbf{underflow}, \mathbf{overflow}\} \\ convert_{i(D) \rightarrow i(F)}(\hat{\mathbf{i}} \cdot y) &= \hat{\mathbf{i}} \cdot convert_{D \rightarrow F}(y) \end{aligned}$$

The $convert_{c(D) \rightarrow c(F)}$ operation:

$$\begin{aligned} convert_{c(D) \rightarrow c(F)} : c(D) &\rightarrow c(F) \cup \{\mathbf{underflow}, \mathbf{overflow}\} \\ convert_{c(D) \rightarrow c(F)}(x + \hat{\mathbf{i}} \cdot y) &= convert_{D \rightarrow F}(x) + \hat{\mathbf{i}} \cdot convert_{D \rightarrow F}(y) \end{aligned}$$

5.5 Support for imaginary and complex numerals

Rather than specifying special numerals for imaginary and complex values, imaginary and complex units are specified in this document. These imaginary and complex units can be used as a basis for expressions that transform ordinary numerals to imaginary or complex numerals. Many programming languages also specify special numeral forms for imaginary or complex numerals, either for the programming language syntax or for input/output formats. These can be expressed using the units specified here.

The imaginary and complex unit operations are as follows:

$$\begin{aligned} imaginary_unit_{i(F)} &= \hat{\mathbf{i}} \cdot 1 \end{aligned}$$

$$\begin{aligned} imaginary_unit_{c(I)} &= 0 + \hat{\mathbf{i}} \cdot 1 \end{aligned}$$

$$\begin{aligned} imaginary_unit_{i(F)} &= \hat{\mathbf{i}} \cdot 1 \end{aligned}$$

$$\begin{aligned} imaginary_unit_{c(F)} &= 0 + \hat{\mathbf{i}} \cdot 1 \end{aligned}$$

6 Notification

Notification is the process by which a user or program is informed that an arithmetic operation cannot return a suitable numeric result. Specifically, a notification shall occur when any arithmetic operation returns an exceptional value. Notification shall be performed according to the requirements of Clause 6 of part 1 of ISO/IEC 10967.

An implementation shall not give notifications for operations conforming to this document, unless the specification requires that an exceptional value results for the given arguments.

The default method of notification should be recording of indicators.

6.1 Continuation values

If notifications are handled by a recording of indicators, in the event of notification the implementation shall provide a *continuation value* to be used in subsequent arithmetic operations. Continuation values may be in I , $i(I)$, $c(I)$, F , $i(F)$ or $c(F)$ (as appropriate), or be special values (where the real or imaginary component is -0 , $-\infty$, $+\infty$, or a **qNaN**).

Floating point datatypes that satisfy the requirements of IEC 60559 have special values in addition to the values in F . These are: -0 , $+\infty$, $-\infty$, *signaling NaNs* (**sNaN**), and *quiet NaNs* (**qNaN**). Such values may be components of complex floating point datatypes, and may be included in values passed as arguments to operations, and used in results or continuation values. Floating point types that do not fully conform to IEC 60559 can also have values corresponding to -0 , $+\infty$, $-\infty$, or **NaN**.

7 Relationship with language standards

A computing system often provides some of the operations specified in this document within the context of a programming language. The requirements of this document shall be in addition to those imposed by the relevant programming language standards.

This document does not define the syntax of arithmetic expressions. However, programmers need to know how to reliably access the operations specified in this document.

NOTE 1 – Providing the information required in this clause is properly the responsibility of programming language standards (or a binding standard between this document and a programming language). An individual implementation would only need to provide details if it could not cite an appropriate clause of the language or binding standard.

An implementation shall document the notation that should be used to invoke an operation or to access a parameter specified in this document and made available. An implementation should document the notation that should be used to invoke an operation or to access a parameter specified in this document and that could be made available.

NOTES

- 2 For example, the complex radian arc sine operation for an argument x ($\arcsin_{c(F)}(x)$) might be invoked as

<code>arcsin(x)</code>	in Ada [2]
<code>casin(x)</code>	in C [4]
<code>asin(x)</code>	in Fortran [6] and C++ [5]
<code>(asin x)</code>	in Common Lisp [10]

with a suitable expression of the argument (x).

- 3 The requirement that the notation that should be used for accessing operations or parameters that are not made available should also be documented is a recommendation to reserve that notation for future use (or present use by some implementations). Thus, doing so would be a way of saying that that notation should not be used for something else, nor, e.g., for a less accurate but otherwise similar operation.

An implementation shall document the semantics of arithmetic expressions in terms of compositions of the operations specified in Clause 5 of this document and in Clause 5 of part 1 and Clause 5 of part 2 of ISO/IEC 10967.

Compilers often “optimize” code as part of compilation. Thus, an arithmetic expression might not be executed as written. An implementation shall document the possible transformations of arithmetic expressions (or groups of expressions) that it permits. Typical transformations include:

- a) Insertion of operations, such as datatype conversions or changes in precision.
- b) Replacing operations (or entire subexpressions) with others, such as “`cos(-x)`” → “`cos(x)`” (exactly the same result) or “`pi - arccos(x)`” → “`arccos(-x)`” (more accurate result).
- c) Evaluating constant subexpressions.
- d) Eliminating unneeded subexpressions.

Only transformations which alter the semantics of an expression (the values produced, and the notifications generated) need be documented. Only the kinds of permitted transformations need be documented. It is not necessary to describe the specific choice of transformations that will be applied to a particular expression.

The textual scope of such transformations shall be documented, and any mechanisms that provide programmer control over this process should be documented as well.

NOTE 4 – It is highly desirable that programming languages intended for numerical use provide means for limiting the transformations applied to particular arithmetic expressions. Control over changes of precision is particularly useful.

8 Documentation requirements

In order to conform to this part, an implementation shall include documentation providing the following information to programmers.

NOTE 1 – Much of the documentation required in this clause is properly the responsibility of programming language or binding standards. An individual implementation would only need to provide details if it could not cite an appropriate clause of the language or binding standard.

- a) A list of the provided operations that conform to this document.
- b) For each box error mode parameter, the value of that parameter. Only box error mode parameters that are relevant to the provided operations need be given.
- c) For each maximum error parameter, the value of that parameter or definition of that parameter function. Only maximum error parameters that are relevant to the provided operations need be given.
- d) The value of the parameters *big-angle-r_F* and *big-angle-u_F*. Only big angle parameters that are relevant to the provided operations need be given.

- e) For the *nearest_F* function, the rule used for rounding halfway cases, unless *iec_559_F* is fixed to **true**.
- f) For each conforming operation, the continuation value provided for each notification condition. Specific continuation values that are required by this document need not be documented. If the notification mechanism does not make use of continuation values (see Clause 6), continuation values need not be documented.

NOTE 2 – Implementations that do not provide infinities or NaNs will have to document any continuation values used in place of such values.

- g) For each conforming operation, how the results depend on the rounding mode, if rounding modes are provided. Operations may be insensitive to the rounding mode, or sensitive to it, but even then need not heed the rounding mode.
- h) For each conforming operation, the notation to be used for invoking that operation.
- i) For each box error mode parameter, the notation to be used to access that parameter.
- j) For each maximum error parameter, the notation to be used to access that parameter.
- k) The notation to be used to access the parameters *big_angle_r_F* and *big_angle_u_F*.
- l) For each of the provided operations where this document specifies a relation to another operation specified in this International Standard, the binding for that other operation.
- m) For numerals conforming to this International Standard, which available string conversion operations, including reading from input, give exactly the same conversion results, even if the string syntaxes for ‘internal’ and ‘external’ numerals are different.

Since the integer and floating point datatypes used in conforming operations shall satisfy the requirements of part 1 of ISO/IEC 10967, the following information shall also be provided by any conforming implementation.

- n) The means for selecting the modes of operation that ensure conformity.
- o) The translation of arithmetic expressions into combinations of the operations provided by any part of ISO/IEC 10967, including any use made of higher precision. (See Clause 7 of part 1.)
- p) The methods used for notification, and the information made available about the notification. (See Clause 6 of part 1.)
- q) The means for selecting among the notification methods, and the notification method used in the absence of a user selection. (See Clause 6.3 of part 1.)
- r) When “recording of indicators” is the method of notification, the datatype used to represent *Ind* (see Clause 6.1.2 of part 1), the method for denoting the values of *Ind*, and the notation for invoking each of the “indicator” operations. *E* is the set of notification indicators. The association of values in *Ind* with subsets of *E* must be clear. In interpreting Clause 6.1.2 of part 1, the set of indicators *E* shall be interpreted as including all exceptional values listed in the signatures of conforming operations. In particular, *E* may need to contain **infinitary** and **absolute_precision_underflow**.

Annex A (normative)

Partial conformity

If an implementation of an operation fulfills all relevant requirements according to the main normative text in this part, except the ones relaxed in this annex, the implementation of that operation is said to *partially conform* to this part.

Conformity to this part shall not be claimed for operations that only fulfill partial conformity.

Partial conformity shall not be claimed for operations that relax other requirements than those relaxed in this annex.

A.1 Maximum error relaxation

This part has the following maximum error requirements for conformity.

$$\max_error_mul_{c(F)} \in [0.5, 5]$$

$$\max_error_div_{c(F)} \in [0.5, 13]$$

$$\max_error_exp_{c(F)} \in [0.5, 7]$$

$$\max_error_power_{c(F)} \in [0.5, 15]$$

$$\max_error_sin_{c(F)} \in [0.5, 11]$$

$$\max_error_tan_{c(F)} \in [0.5, 14]$$

In a partially conforming implementation the maximum error parameters may be greater than what is specified by this part. The maximum error parameter values given by an implementation shall still adequately reflect the accuracy of the relevant operations, if a claim of partial conformity is made.

A partially conforming implementation shall document which maximum error parameters have greater values than specified by this part, and their values.

A.2 Extra accuracy requirements relaxation

This part has a number of extra accuracy requirements. These are detailed in the paragraphs beginning “Further requirements on the op_F^* approximation helper function are:”.

In a partially conforming implementation these further requirements need not be fulfilled. The values returned must still be within the maximum error bounds that are given by the maximum error parameters, if a claim of partial conformity is made.

A partially conforming implementation shall document which extra accuracy requirements are not fulfilled by the implementation.

A.3 Relationships to other operations relaxation

This part has a number of requirements giving relations to other operations. These are detailed in the paragraphs beginning with wordings like “Relationship to the op_F^* approximation helper function for operations in an associated library shall be:”.

In a partially conforming implementation these relationships need not be fulfilled. The values returned must still be within the maximum error bounds that are given by the values provided for the maximum error parameters, if a claim of partial conformity is made.

A partially conforming implementation shall document which operation relationships are not fulfilled by the implementation.

A.4 Part 1 and part 2 requirements relaxation

Part 3 depends on the datatypes and operations specified in part 1, as well as the operations (and approximation helper functions) specified in part 2. Part 1 and part 2 allow for partial conformity. Part 3 operations may thus be only partially conforming if a relevant datatype or operation is only partially conforming to part 1 or part 2.

IECNORM.COM : Click to view the full PDF of ISO/IEC 10967-3:2006

Annex B (informative)

Rationale

In the informative annexes of this document, LIA-1 refers to part 1 of ISO/IEC 10967, LIA-2 refers to part 2, and LIA-3 refers to this document.

This annex explains and clarifies some of the ideas behind *Information technology – Language independent arithmetic – Part 3: Complex floating point arithmetic and complex elementary numerical functions*. The clause numbering matches that of the standard.

B.1 Scope

B.1.1 Inclusions

Integer complex datatypes are included in Common Lisp [10]. Integer complex numbers are sometimes referred to as Gaussian numbers.

Imaginary datatypes are included informatively in the current version of the C programming language standard (annex G of [4]) and normatively in the current version of the Ada language standard (annex G of [2]).

B.1.2 Exclusions

LIA-3 only covers Cartesian complex datatypes. Polar complex datatypes are not included since no widely known programming language provides them. However, operations for conversion from polar representation to Cartesian representation are included both in some programming languages and in LIA-3.

Neither rational nor fixed point complex datatypes are covered by LIA-3 since (as yet) no part of LIA covers rational or fixed point datatypes. (Note that Common Lisp includes complex datatypes that have rational datatype values in the parts.)

B.2 Conformity

Conformity to this standard is dependent on the existence of language binding standards. Each programming language committee (or other organization responsible for a programming language or other specification to which LIA-1, LIA-2, and LIA-3 may apply) is encouraged to produce a binding covering at least those operations already required by the programming language (or similar) and also specified in LIA-3.

The term “programming language” is here used in a generalised sense to include other computing entities such as calculators, spread sheets, page description languages, web-script languages, and database query languages to the extent that they provide the operations covered by LIA-3.

Suggestions for bindings are provided in Annex C. Annex C has partial binding examples for a number of existing programming languages and LIA-3. In addition to the bindings for the operations in LIA-3, it is also necessary to provide bindings for the maximum error parameters

specified by LIA-3. Annex C contains suggestions for these bindings. To conform to this standard, in the absence of a binding standard, an implementation should create a binding, following the suggestions in annex C.

B.3 Normative references

The referenced IEC 60559 standard is identical to the IEEE 754 standard and the former IEC 559 standard.

The origin of IEC 60559:1989, IEEE 754:1984, was undergoing revision during the development of this part.

See Annex E of ISO/IEC 10967-2:2001 regarding the first edition (published in 1994) of ISO/IEC 10967-1.

See Annex F of this part regarding the first edition (published in 2001) of ISO/IEC 10967-2.

B.4 Symbols and definitions

B.4.1 Symbols

B.4.1.1 Sets and intervals

The interval notation is in common use. It has been chosen over the other commonly used interval notation because the chosen notation has no risk of confusion with the pair notation.

B.4.1.2 Operators and relations

Note that all operators, relations, and other mathematical notation used in LIA-3 are used in their conventional exact mathematical sense. They are not used to stand for operations specified by IEC 60559, LIA-1, LIA-2, LIA-3, or, with the exception of program excerpts which are clearly marked, any programming language. E.g. x/u stands for the mathematically exact result of dividing x by u , whether that value is representable in any floating point datatype or not, and $x/u \neq \text{div}_F(x, u)$ is often the case. Likewise, $=$ is the mathematical equality, not the eq_F operation: $0 \neq -0$, while $eq_F(0, -0) = \text{true}$.

For mathematical complex values, conventional notation with $+$ and $\cdot$ is used. For the imaginary unit, the symbol i is used.

It is important to distinguish this mathematical notation for values in $\mathcal{C}$ and the notation used to express values in $c(X)$ or $i(X)$. For $c(X)$ the binary operator $\mathbf{+i\cdot}$ is used. Its only function is to create a pair of values corresponding to a value in $\mathcal{C}$. Similarly, for $i(X)$, the unary operator $\mathbf{i\cdot}$ is used. It creates a 1-tuple corresponding to a value in $\mathcal{C}$ where the real part is 0.

B.4.1.3 Mathematical functions

The elementary functions named $\sin$, $\cos$, etc., used in LIA-3 are the exact mathematical functions, not any approximation. The approximations to these mathematical functions are introduced in clauses 5.3 and 5.4 and are written in a way clearly distinct from the exact mathematical functions. E.g., $\sin_{c(F)}^*$, $\cos_{c(F)}^*$, etc., which are unspecified mathematical functions approximating the targeted exact mathematical functions to a specified degree; $\sin_{c(F)}$, $\cos_{c(F)}$, etc., which are

the operations specified by LIA-3 based on the respective approximating function; `sin`, `cos`, etc., which are programming language names bound to LIA-3 operations. `sin` is thus very different from `sin`.

B.4.1.4 Exceptional values

LIA-3 uses the same set of exceptional values as LIA-2.

B.4.1.5 Datatypes and special values

LIA allows for three methods for handling notifications: recording of indicators, change of control flow (returnable or not), and termination of program. The preferred method is recording of indicators. This allows the computation to continue using the continuation values. For **underflow** and **infinitary** notifications this course of action is strongly preferred, provided that a suitable continuation value can be represented in the result datatype.

Not all occurrences of the same exceptional value need be handled the same. There may be explicit mode changes in how notifications are handled, and there may be implicit changes. E.g., **invalid** without a specified (by LIA-3 or binding) continuation value to cause change of control flow (like an Ada [2] exception), while **invalid** with a specified continuation value to use recording of indicators. This should be specified by bindings or by implementations.

The operations may return any of the exceptional values **overflow**, **underflow**, **invalid**, **infinitary**, or **absolute_precision_underflow**. This does *not* imply that the implemented operations are to actually *return* any of these values. When these values are returned according to the LIA specification, that means that the implementation is to perform a notification handling for that exceptional value. If the notification handling is by recording of indicators, then what is actually returned by the implemented operation is the continuation value.

B.4.1.6 Complex value constructors and complex datatype constructors

For integer and floating point datatypes, the set of non-special values are subsets of $\mathcal{Z}$ and $\mathcal{R}$ respectively (LIA-1). However, if we do the parallel here, defining the set of non-special values as subsets of $\mathcal{G}$ and $\mathcal{C}$, the handling of special values becomes awkward. The latter would need a pair representation for the full complex case. Instead, we use a pair representation for all complex values. In addition, we then also get an easy notational difference between imaginary and full complex datatypes. Unfortunately, we also get a distinction between this pair formation, and the arithmetic expression resulting in a complex value in $\mathcal{G}$ or $\mathcal{C}$. This distinction has to be carefully maintained in the specifications in the rest of the document.

B.4.2 Definitions of terms

Note the LIA distinction between exceptional values, exceptions, and exception handling (handling of notification by non-returnable change of control flow; as in e.g. Ada). LIA exceptional values are not the same as Ada exceptions, nor are they the same as IEC 60559 special values.

B.5 Specifications for the imaginary and complex datatypes and operations

B.5.1 Imaginary and complex integer datatypes and operations

Of the programming languages covered in annex C, only Common Lisp has support for complex integers, but not imaginary integers as separate from complex integers.

Since LIA-3 distinguishes between complex and imaginary floating point, that same distinction is also made for the integer case.

“Real” integers and imaginary integers can be compared for order, while complex integers only can be compared for equality or inequality.

The re_I operation is specified for completeness. It is the identity function on integers; except possibly for NaN values (which are rare in integer datatypes). Likewise is the im_I operation specified for completeness. It always returns zero. Note, however, that the imaginary part of a “real” floating point value is not the same for all arguments, even though it is always a zero. Similarly, im_I could return signed zero as im_F does, if the integer datatype supports signed zero. This is, however, rare, so is not specified explicitly by LIA-3. Also the $conj_I$ and $conj_{i(I)}$ operations are specified for completeness.

The binary operations, except comparisons for order, allow for combining “real” integer, imaginary integer and complex integer of the same base integer datatype. Bindings may generalise this to also allow different base integer datatypes, converting the smaller ones to the largest one occurring as arguments to the operation.

$signum_I$ is specified in this part of LIA, since the $sign_I$ operation specified in the first edition of LIA-1 is not consistent with the $signum_F$ operation, also specified in this part. The $sign_I$ and $sign_F$ of the first edition of LIA-1 return zero for a zero argument. This is not consistent with the branch cuts specified for elementary operations in this part. For $signum_I$ and $signum_F$ the branch cut view is used, and they only return -1 or 1 , never zero (of any sign).

There is no $signum_{c(I)}$ operation. The mathematical function would normalise the value to the unit circle. But since only four values in $c(I)$ is on, or even close, to the unit circle, the approximation for returning a $c(I)$ value is too great. If one wants to compute an approximation to the signum of a $c(I)$ value, first convert the value to $c(F)$, and then compute $signum_{c(F)}$ of that value.

$divides_{c(I)}$ is the operation version of the divides relation specified in clause 4.1.

Max and min operations are defined only for imaginary integer values, since complex integers are not ordered in mathematics. A lexicographic order is easy to define, but that is another kind of ordering that is out of scope for LIA.

B.5.2 Imaginary and complex floating point datatypes and operations

B.5.2.1 Maximum error requirements

B.5.2.2 Sign requirements

B.5.2.3 Maximum error requirements

The values for the error parameters given later in the normative part of this document are largely taken from the Ada standard [2].

B.5.2.4 Basic arithmetic for complex floating point

F must be a subset of $\mathcal{R}$. Floating point datatypes can have infinity values as well as NaN values, and also may have a -0 . These values are not in F . The special values are, however, commonly available in floating point datatypes today, thanks to the wide adoption of IEC 60599.

Note that for some operations the exceptional value **invalid** is produced only for argument values involving -0 , $+\infty$, $-\infty$, or **sNaN**. For these operations the signature given in LIA-3 does not contain **invalid**.

The report *Floating-Point C Extensions* [16], issued by the ANSI X3J11 committee, discusses possible ways of exploiting the IEC 60559 special values. The report identifies some of its suggestions as controversial and cites *Branch Cuts for Complex Elementary Functions, or Much Ado about Nothing's Sign Bit* [14] as justification. The report *Complex Floating-Point C Extensions* [17], also issued by the ANSI X3J11 committee, extends these recommendations to imaginary and complex datatypes. These reports have strongly influenced the C99 [4] programming language.

The implicit zero imaginary part of a value in F and the implicit zero real part of a value in $i(F)$ is a very strong zero, here written **00**. Note that **00** is *not* represented in the datatype corresponding to F , nor in $i(F)$ or $c(F)$, not even as a component. **00** is strong in the sense that **00** overtrumps NaN and infinity values, which neither 0 nor -0 does. I.e. the implicit zeroes (**00**) work *as if* the rules of the following table applies (plus commutativity):

$mul_F(0, +\infty) = \text{invalid}(\text{qNaN})$	$mul_F(-0, +\infty) = \text{invalid}(\text{qNaN})$	$mul_F(\text{00}, +\infty) = \text{00}$
$mul_F(0, -\infty) = \text{invalid}(\text{qNaN})$	$mul_F(-0, -\infty) = \text{invalid}(\text{qNaN})$	$mul_F(\text{00}, -\infty) = \text{00}$
$mul_F(0, \text{qNaN}) = \text{qNaN}$	$mul_F(-0, \text{qNaN}) = \text{qNaN}$	$mul_F(\text{00}, \text{qNaN}) = \text{00}$
$mul_F(0, \text{sNaN}) = \text{invalid}(\text{qNaN})$	$mul_F(-0, \text{sNaN}) = \text{invalid}(\text{qNaN})$	$mul_F(\text{00}, \text{sNaN}) = \text{00}$

and further as if $add_F(\text{00}, x) = x$, even for signalling NaNs. The basic idea here is that **00** is not only mostly treated as exact even though it *may* be the result of an approximation (as other values are handled), but that **00** really is exact, never the result of an approximation. Instead of representing **00** in any floating point datatype, which would necessitate new floating point datatypes, several programming languages have chosen to use triplets of datatypes (corresponding to F , $i(F)$ and $c(F)$), but adding special values. LIA-3 adheres to this convention.

When the implicit **00** is made explicit, one has to choose between 0 and -0 (if the latter is representable). In doing so, sign symmetry is observed by the specifications in LIA-3. This is done in such a way that it is coherent with the sign symmetry of some trigonometric (and hyperbolic) operations that accept values in F (or $i(F)$) but return a value in $c(F)$. This is in order to maintain also the sign symmetry, even when -0 is not representable. This also adheres to the sign symmetry conventions of Common Lisp and C99. Further, when negative zero (-0) is not representable, a 0 as the real part or as the imaginary part of the argument is to be considered to be -0 consistently with the above.

Another approach would be not to regard implicit zeroes as **00**, but as ordinary zero. This renders the $i(F)$ datatypes superfluous, and they are usually omitted in this approach. This approach is both less efficient and loses information needlessly. In this approach one also uses just one complex NaN, only one (unsigned) complex infinity, and only one (unsigned) complex 0. Again this loses information needlessly, in particular if the infinities or zeroes are the result of approximations. Again that speaks against this approach. In this approach one sometimes also do not distinguish the signs of zero, forcing programmers to write expressions like `max(posvalue, FMIN)` and `min(negvalue, -FMIN)`, to avoid zero but (with error prone programming) nearly mimic the signed zeroes. In this approach one does also not always follow conjugate symmetry

and sign symmetry rules when part of the argument is zero, which for this approach leads to surprising changes of results when doing algebraic manipulations based on those rules. For the reasons given here, this approach is not taken for LIA-3.

B.5.3 Elementary transcendental imaginary and complex floating point operations

B.5.3.1 Operations for exponentiations and logarithms

The LIA-3 exponentiation and inverse exponentiation operations that take imaginary or complex arguments interpret the imaginary part (of the argument or result respectively) as in radians. Operations that take an extra floating point argument, giving the non-radian cyclic unit, can be made, paralleling such operations (for trigonometric operations) in LIA-2. Such operations are not included in LIA-3, because no programming language requires them for complex arguments.

However, the complex result power operation uses $|z/2| \leq \text{big_angle_}u_F$, as the “angle” cutoff, since that is more logical in connection with the power operation, as well as being easier to implement (no multiplication by π).

B.5.3.2 Operations for radian trigonometric elementary functions

The LIA-3 trigonometric and inverse trigonometric operations that take imaginary or complex arguments interpret the real part (of the argument or result respectively) as in radians. Operations that take an extra floating point argument, giving the non-radian cyclic unit, can be made, paralleling such operations in LIA-2. Such operations are not included in LIA-3, because no programming language require them (with the exception of *phaseu* and *polaru* which are required by Ada).

The notation here is taken from *Handbook of mathematical functions with graphs, formulas, and mathematical tables*. The equalities that involve an uppercased elementary function name refer the multi-valued inverses. The equalities then say the values on the left side (seen as a set) is equal to the values on the right side (seen as a set). This could be expressed more conventionally, using set notation, but this notation is borrowed since it is convenient. The lowercase names for inverses are the corresponding functions, denoting the principal value for the multi-valued inverse.

B.5.3.2.1 Radian angle normalisation

The radian angle normalisation operations normalise the angle part (the real part) of the argument to be within $[\text{down}_F(-\pi), \text{up}_F(\pi)]$ (or, if the correction given in annex F is used: $[\text{up}_F(-\pi), \text{down}_F(\pi)]$). While useful for programs where one wishes to keep angle information at a high accuracy, it is not, and cannot be, sufficiently accurate for implementing the trigonometric operations specified in LIA-3. A normalisation function that is almost sufficient for that is specified in LIA-2: *axis_rad_F*. Some extra accuracy bits for the normalisation are still needed for implementing the trigonometric operations with respect to sign and accuracy requirements. *axis_rad_F* is not generalised to the complex case, since there is little or no convenience to be gained compared to using *axis_rad_F* directly on the relevant part of a complex value.

B.5.3.2.2 Radian sine

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}
 \sin(-z) &= -\sin(z) \\
 \sin(\text{conj}(z)) &= \text{conj}(\sin(z)) \\
 \sin(z + k \cdot 2 \cdot \pi) &= \sin(z) \text{ if } k \in \mathcal{Z} \\
 \sin(z) &= \cos(\pi/2 - z) \\
 \\
 \sin(x) &= -\tilde{\mathbf{i}} \cdot \sinh(\tilde{\mathbf{i}} \cdot x) = \tilde{\mathbf{i}} \cdot \sinh(-\tilde{\mathbf{i}} \cdot x) \\
 \sin(\tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \sinh(-y) = \tilde{\mathbf{i}} \cdot \sinh(y) \\
 \sin(x + \tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \sinh(-y + \tilde{\mathbf{i}} \cdot x) = \tilde{\mathbf{i}} \cdot \sinh(y - \tilde{\mathbf{i}} \cdot x) \\
 \\
 \sin(x + \tilde{\mathbf{i}} \cdot y) &= \sin(x) \cdot \cosh(y) + \tilde{\mathbf{i}} \cdot \cos(x) \cdot \sinh(y)
 \end{aligned}$$

B.5.3.2.3 Radian cosine

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}
 \cos(-z) &= \cos(z) \\
 \cos(\text{conj}(z)) &= \text{conj}(\cos(z)) \\
 \cos(z + k \cdot 2 \cdot \pi) &= \cos(z) \text{ if } k \in \mathcal{Z} \\
 \cos(z) &= \sin(\pi/2 - z) \\
 \\
 \cos(x) &= \cosh(\tilde{\mathbf{i}} \cdot x) = \cosh(-\tilde{\mathbf{i}} \cdot x) \\
 \cos(\tilde{\mathbf{i}} \cdot y) &= \cosh(y) \\
 \cos(x + \tilde{\mathbf{i}} \cdot y) &= \cosh(-y + \tilde{\mathbf{i}} \cdot x) = \cosh(y - \tilde{\mathbf{i}} \cdot x) \\
 \\
 \cos(x + \tilde{\mathbf{i}} \cdot y) &= \cos(x) \cdot \cosh(y) - \tilde{\mathbf{i}} \cdot \sin(x) \cdot \sinh(y)
 \end{aligned}$$

B.5.3.2.4 Radian tangent

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}
 \tan(-z) &= -\tan(z) \\
 \tan(\text{conj}(z)) &= \text{conj}(\tan(z)) \\
 \tan(z + k \cdot 2 \cdot \pi) &= \tan(z) \text{ if } k \in \mathcal{Z} \\
 \tan(z) &= \cot(\pi/2 - z) \\
 \\
 \tan(x) &= -\tilde{\mathbf{i}} \cdot \tanh(\tilde{\mathbf{i}} \cdot x) = \tilde{\mathbf{i}} \cdot \tanh(-\tilde{\mathbf{i}} \cdot x) \\
 \tan(\tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \tanh(-y) = \tilde{\mathbf{i}} \cdot \tanh(y) \\
 \tan(x + \tilde{\mathbf{i}} \cdot y) &= -\tilde{\mathbf{i}} \cdot \tanh(-y + \tilde{\mathbf{i}} \cdot x) = \tilde{\mathbf{i}} \cdot \tanh(y - \tilde{\mathbf{i}} \cdot x)
 \end{aligned}$$

Note that every non-zero integer multiple of π is a cycle for tangent. Since $2 \cdot \pi$ is the smallest positive cycle that is common to all of the trigonometric functions, that is the cycle concentrated on for all of them, even though some have π as cycle too.

B.5.3.2.5 Radian cotangent

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\cot(-z) &= -\cot(z) \\ \cot(\operatorname{conj}(z)) &= \operatorname{conj}(\cot(z)) \\ \cot(z + k \cdot 2 \cdot \pi) &= \cot(z) \text{ if } k \in \mathbb{Z} \\ \cot(z) &= \tan(\pi/2 - z) \\ \cot(z) &= 1/\tan(z)\end{aligned}$$

$$\begin{aligned}\cot(x) &= \tilde{\imath} \cdot \coth(\tilde{\imath} \cdot x) = -\tilde{\imath} \cdot \coth(-\tilde{\imath} \cdot x) \\ \cot(\tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \coth(-y) = -\tilde{\imath} \cdot \coth(y) \\ \cot(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \coth(-y + \tilde{\imath} \cdot x) = -\tilde{\imath} \cdot \coth(y - \tilde{\imath} \cdot x)\end{aligned}$$

B.5.3.2.6 Radian secant

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\sec(-z) &= \sec(z) \\ \sec(\operatorname{conj}(z)) &= \operatorname{conj}(\sec(z)) \\ \sec(z + k \cdot 2 \cdot \pi) &= \sec(z) \text{ if } k \in \mathbb{Z} \\ \sec(z) &= \csc(\pi/2 - z) \\ \sec(z) &= 1/\cos(z)\end{aligned}$$

$$\begin{aligned}\sec(x) &= \operatorname{sech}(\tilde{\imath} \cdot x) = \operatorname{sech}(-\tilde{\imath} \cdot x) \\ \sec(\tilde{\imath} \cdot y) &= \operatorname{sech}(-y) = \operatorname{sech}(y) \\ \sec(x + \tilde{\imath} \cdot y) &= \operatorname{sech}(-y + \tilde{\imath} \cdot x) = \operatorname{sech}(y - \tilde{\imath} \cdot x)\end{aligned}$$

B.5.3.2.7 Radian cosecant

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\csc(-z) &= -\csc(z) \\ \csc(\operatorname{conj}(z)) &= \operatorname{conj}(\csc(z)) \\ \csc(z + k \cdot 2 \cdot \pi) &= \csc(z) \text{ if } k \in \mathbb{Z} \\ \csc(z) &= \sec(\pi/2 - z) \\ \csc(z) &= 1/\sin(z)\end{aligned}$$

$$\begin{aligned}\csc(x) &= \tilde{\imath} \cdot \operatorname{csch}(\tilde{\imath} \cdot x) = -\tilde{\imath} \cdot \operatorname{csch}(-\tilde{\imath} \cdot x) \\ \csc(\tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \operatorname{csch}(-y) = -\tilde{\imath} \cdot \operatorname{csch}(y) \\ \csc(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \operatorname{csch}(-y + \tilde{\imath} \cdot x) = -\tilde{\imath} \cdot \operatorname{csch}(y - \tilde{\imath} \cdot x)\end{aligned}$$

B.5.3.2.8 Radian arc sine

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\arcsin(-z) &= -\arcsin(z) \\ \arcsin(z) &= \pi/2 - \arccos(z) \\ \arcsin(\operatorname{conj}(z)) &= \operatorname{conj}(\arcsin(z)) \text{ if } \Im(z) \neq 0 \text{ or } |\Re(z)| \leq 1 \\ \operatorname{Arcsin}(x + \tilde{\imath} \cdot y) &= -\tilde{\imath} \cdot \operatorname{Arcsinh}(-y + \tilde{\imath} \cdot x) = \tilde{\imath} \cdot \operatorname{Arcsinh}(y - \tilde{\imath} \cdot x)\end{aligned}$$

B.5.3.2.9 Radian arc cosine

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\arccos(-z) &= \tilde{\imath} \cdot \pi - \arccos(z) \\ \arccos(\text{conj}(z)) &= \text{conj}(\arccos(z)) \text{ if } \Im(z) \neq 0 \text{ or } |\Re(z)| \leq 1 \\ \arccos(z) &= \pi/2 - \arcsin(z) \\ \text{Arccos}(x + \tilde{\imath} \cdot y) &= \pm \tilde{\imath} \cdot \text{Arccosh}(x + \tilde{\imath} \cdot y)\end{aligned}$$

B.5.3.2.10 Radian arc tangent

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\arctan(-z) &= -\arctan(z) \\ \arctan(\text{conj}(z)) &= \text{conj}(\arctan(z)) \text{ if } \Re(z) \neq 0 \text{ or } |\Im(z)| \leq 1 \\ \arctan(z) &= \pm \pi/2 - \text{arccot}(z) \\ \text{Arctan}(x + \tilde{\imath} \cdot y) &= -\tilde{\imath} \cdot \text{Arctanh}(-y + \tilde{\imath} \cdot x) = \tilde{\imath} \cdot \text{Arctanh}(y - \tilde{\imath} \cdot x)\end{aligned}$$

B.5.3.2.11 Radian arc cotangent

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\text{arccot}(-z) &= -\text{arccot}(z) \\ \text{arccot}(\text{conj}(z)) &= \text{conj}(\text{arccot}(z)) \text{ if } \Re(z) \neq 0 \text{ or } |\Im(z)| \geq 1 \\ \text{arccot}(z) &= \pm \pi/2 - \arctan(z) \\ \text{Arccot}(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \text{Arccoth}(-y + \tilde{\imath} \cdot x) \\ \text{arccot}(z) &= \arctan(1/z)\end{aligned}$$

B.5.3.2.12 Radian arc secant

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\text{arcsec}(-z) &= \pi - \text{arcsec}(z) \\ \text{arcsec}(\text{conj}(z)) &= \text{conj}(\text{arcsec}(z)) \text{ if } \Im(z) \neq 0 \text{ or } |\Re(z)| \geq 1 \\ \text{arcsec}(z) &= \pi/2 - \text{arccsc}(z) \\ \text{Arcsec}(x + \tilde{\imath} \cdot y) &= \pm \tilde{\imath} \cdot \text{Arcsech}(x + \tilde{\imath} \cdot y) \\ \text{arcsec}(z) &= \arccos(1/z)\end{aligned}$$

B.5.3.2.13 Radian arc cosecant

The following equalities have been used to derive the LIA-3 requirements:

$$\begin{aligned}\text{arccsc}(-z) &= -\text{arccsc}(z) \\ \text{arccsc}(\text{conj}(z)) &= \text{conj}(\text{arccsc}(z)) \text{ if } \Im(z) \neq 0 \text{ or } |\Re(z)| \geq 1 \\ \text{arccsc}(z) &= \pi/2 - \text{arcsec}(z) \\ \text{Arccsc}(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \text{Arccsch}(-y + \tilde{\imath} \cdot x) \\ \text{arccsc}(z) &= \arcsin(1/z)\end{aligned}$$

B.5.3.3 Operations for hyperbolic elementary functions

The operations for hyperbolic elementary functions are all defined directly in terms of “turned” imaginary and complex trigonometric operations of a “turned” argument. Their specifications are therefore rather short. It’s done this way for several reasons:

- a) The hyperbolic operations are less commonly used than the trigonometric ones.
- b) The connections with the corresponding trigonometric operations are exact rather than approximate.
- c) The hyperbolic functions have some ‘irregularities’ (expressed as conditionals in the specifications for the inverse hyperbolic operations for arccosh and arcsech) that the trigonometric operations don’t have. It is therefore slightly easier to specify the hyperbolic operations in terms of the trigonometric ones rather than the other way around.
- d) C also specifies exact correspondences between the complex hyperbolic operations and the complex trigonometric operations (though the other way around, and skipping the “irregular” operations, compared to the specification in LIA-3).

The LIA-3 hyperbolic and inverse hyperbolic operations that take imaginary or complex arguments interpret the imaginary part (of the argument or result respectively) as in radians. Operations that take an extra floating point argument, giving the non-radian cyclic unit, can be made, paralleling such operations (for trigonometric operations) in LIA-2. Such operations are not included in LIA-3, because no programming language require them.

B.5.3.3.1 Hyperbolic normalisation

B.5.3.3.2 Hyperbolic sine

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\sinh(x) &= -\tilde{i} \cdot \sin(\tilde{i} \cdot x) = \tilde{i} \cdot \sin(-\tilde{i} \cdot x) \\ \sinh(\tilde{i} \cdot y) &= -\tilde{i} \cdot \sin(-y) = \tilde{i} \cdot \sin(y) \\ \sinh(x + \tilde{i} \cdot y) &= -\tilde{i} \cdot \sin(-y + \tilde{i} \cdot x) = \tilde{i} \cdot \sin(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.3 Hyperbolic cosine

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\cosh(x) &= \cos(\tilde{i} \cdot x) = \cos(-\tilde{i} \cdot x) \\ \cosh(\tilde{i} \cdot y) &= \cos(y) = \cos(-y) \\ \cosh(x + \tilde{i} \cdot y) &= \cos(-y + \tilde{i} \cdot x) = \cos(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.4 Hyperbolic tangent

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\tanh(x) &= -\tilde{i} \cdot \tan(\tilde{i} \cdot x) = \tilde{i} \cdot \tan(-\tilde{i} \cdot x) \\ \tanh(\tilde{i} \cdot y) &= -\tilde{i} \cdot \tan(-y) = \tilde{i} \cdot \tan(y) \\ \tanh(x + \tilde{i} \cdot y) &= -\tilde{i} \cdot \tan(-y + \tilde{i} \cdot x) = \tilde{i} \cdot \tan(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.5 Hyperbolic cotangent

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\coth(z) &= \tanh(\tilde{i} \cdot \pi/2 - z) \\ \coth(x) &= \tilde{i} \cdot \cot(\tilde{i} \cdot x) = -\tilde{i} \cdot \cot(-\tilde{i} \cdot x) \\ \coth(\tilde{i} \cdot y) &= \tilde{i} \cdot \cot(-y) = -\tilde{i} \cdot \cot(y) \\ \coth(x + \tilde{i} \cdot y) &= \tilde{i} \cdot \cot(-y + \tilde{i} \cdot x) = -\tilde{i} \cdot \cot(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.6 Hyperbolic secant

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\operatorname{sech}(x) &= \sec(\tilde{i} \cdot x) = \sec(-\tilde{i} \cdot x) \\ \operatorname{sech}(\tilde{i} \cdot y) &= \sec(-y) = \sec(y) \\ \operatorname{sech}(x + \tilde{i} \cdot y) &= \sec(-y + \tilde{i} \cdot x) = \sec(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.7 Hyperbolic cosecant

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\operatorname{csch}(z) &= \operatorname{sech}(\tilde{i} \cdot \pi/2 - z) \\ \operatorname{csch}(x) &= \tilde{i} \cdot \csc(\tilde{i} \cdot x) = -\tilde{i} \cdot \csc(-\tilde{i} \cdot x) \\ \operatorname{csch}(\tilde{i} \cdot y) &= \tilde{i} \cdot \csc(-y) = -\tilde{i} \cdot \csc(y) \\ \operatorname{csch}(x + \tilde{i} \cdot y) &= \tilde{i} \cdot \csc(-y + \tilde{i} \cdot x) = -\tilde{i} \cdot \csc(y - \tilde{i} \cdot x)\end{aligned}$$

B.5.3.3.8 Inverse hyperbolic sine

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\operatorname{Arcsinh}(x + \tilde{i} \cdot y) = -\tilde{i} \cdot \operatorname{Arcsin}(-y + \tilde{i} \cdot x) = \tilde{i} \cdot \operatorname{Arcsin}(y - \tilde{i} \cdot x)$$

B.5.3.3.9 Inverse hyperbolic cosine

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\operatorname{Arccosh}(x + \tilde{i} \cdot y) = \pm \tilde{i} \cdot \operatorname{Arccos}(x + \tilde{i} \cdot y)$$

B.5.3.3.10 Inverse hyperbolic tangent

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\operatorname{Arctanh}(x + \tilde{\imath} \cdot y) = -\tilde{\imath} \cdot \operatorname{Arctan}(-y + \tilde{\imath} \cdot x) = \tilde{\imath} \cdot \operatorname{Arctan}(y - \tilde{\imath} \cdot x)$$

B.5.3.3.11 Inverse hyperbolic cotangent

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\operatorname{Arccoth}(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \operatorname{Arccot}(-y + \tilde{\imath} \cdot x) \\ \operatorname{arccoth}(z) &= \operatorname{arctanh}(1/z)\end{aligned}$$

B.5.3.3.12 Inverse hyperbolic secant

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\operatorname{arcsech}(z) &= \tilde{\imath} \cdot \pi/2 - \operatorname{arccsch}(z) \text{ if } \Re(z) > 0 \\ \operatorname{Arcsech}(x + \tilde{\imath} \cdot y) &= \pm \tilde{\imath} \cdot \operatorname{Arcsec}(x + \tilde{\imath} \cdot y) \\ \operatorname{arcsech}(z) &= \operatorname{arccosh}(1/z)\end{aligned}$$

B.5.3.3.13 Inverse hyperbolic cosecant

Only the relationship to the corresponding trigonometric function is used in deriving the LIA-3 specification.

$$\begin{aligned}\operatorname{arccsch}(z) &= \tilde{\imath} \cdot \pi/2 - \operatorname{arcsech}(z) \text{ if } \Re(z) > 0 \\ \operatorname{Arccsch}(x + \tilde{\imath} \cdot y) &= \tilde{\imath} \cdot \operatorname{Arccsc}(-y + \tilde{\imath} \cdot x) \\ \operatorname{arccsch}(z) &= \operatorname{arcsinh}(1/z)\end{aligned}$$

B.5.4 Operations for conversion between imaginary and complex numeric datatypes**B.5.5 Support for imaginary and complex numerals****B.6 Notification****B.6.1 Continuation values**

Some operations are specified applications of one or more other operations. In these cases the notifications of the “suboperations” are recorded, if recording of indicators are used, and a result (continuation value) is to be computed from the continuation values of the “suboperations”.

B.7 Relationship with language standards**B.8 Documentation requirements**

Annex C

(informative)

Example bindings for specific languages

This annex describes how a computing system can simultaneously conform to a language standard (or publicly available specification) and to LIA-3. It contains suggestions for binding the “abstract” operations specified in LIA-3 to concrete language syntax. The format used for these example bindings in this annex is a short form version, suitable for the purposes of this annex. An actual binding is under no obligation to follow this format. An actual binding should, however, as in the bindings examples, give the LIA-3 operation name, or parameter name, bound to an identifier (or expression) by the binding.

Portability of programs can be improved if two conforming LIA-3 systems using the same programming language agree in the manner with which they adhere to LIA-3. For instance, LIA-3 requires that the parameter *big_angle_rF* be provided (if any conforming complex radian trigonometric operations are provided), but if one system provides it by means of the identifier **BigAngle** and another by the identifier **MaxAngle** for the same programming language, portability is impaired. Clearly, it would be best if such names were defined in the relevant language standards or binding standards, but in the meantime, suggestions are given here to aid portability. Name consistency cannot, however, be fully maintained between different programming languages, due to already existing differences in naming conventions, and LIA does *not* require wholesale naming changes, nor expression syntax changes.

The following clauses are suggestions rather than requirements because the areas covered are the responsibility of the various programming language standards committees. Until binding standards are in place, implementors can promote “de facto” portability by following these suggestions on their own.

The languages covered in this annex are:

Ada
C
C++
Fortran
Common Lisp

This list is not exhaustive. Other languages and other computing devices (like ‘scientific’ calculators, ‘web script’ languages, and database ‘query languages’) may be suitable for conformity to LIA-3.

In this annex, the parameters, operations, and exception behaviour of each language are examined to see how closely they fit the requirements of LIA-3.

- a) For parameters and operations provided (semanticswise, or with a very close semantics) via the specification of a programming language, their syntax is marked with a ★ below, occasionally noting minor deviations (which bindings are allowed to have).
- b) Where parameters or operations are not provided by the language, suggested names and syntax for them are given but marked with a †.

- c) Some provided operations are closely related to LIA-3 operations, but not specified by this version of LIA-3; the syntax for such operations are in addition to the † noted as “Not LIA-3”.

This annex describes only the language-level support for LIA-3. An implementation that wishes to conform must ensure that the underlying hardware and software is also configured to conform to LIA-3 requirements.

A complete binding for LIA-3 will include, or refer to, a binding for LIA-2 and LIA-1. In turn, a complete binding for LIA-2/LIA-1 may include, or refer to, a binding for IEC 60559. Such a joint LIA-3/LIA-2/LIA-1/IEC 60559 binding should be developed as a single binding standard. To avoid conflict with ongoing development, only the LIA-3 specific portions of such a binding are exemplified in this annex.

Most language standards permit an implementation to provide, by some means, the parameters and operations required by LIA-3 that are not already part of the language. The method for accessing these additional parameters and operations depends on the implementation and language, and is not specified in LIA-3 nor exemplified in this annex. It could include external subroutine libraries; new intrinsic functions supported by the compiler; constants and functions provided as global “macros”; and so on. The actual method of access through libraries, macros, etc. should of course be given in a real binding.

Most language standards do not constrain the accuracy of elementary numerical functions, or specify the subsequent behaviour after an arithmetic notification occurs.

In the event that there is a conflict between the requirements of the language standard and the requirements of LIA-3, the language binding standard should clearly identify the conflict and state its resolution of the conflict.

C.1 Ada

The programming language Ada is defined by ISO/IEC 8652:1995, *Information Technology – Programming Languages – Ada* [2].

An implementation should follow all the requirements of LIA-3 unless otherwise specified by this language binding.

The operations or parameters marked “†” are not part of the language and should be provided by an implementation that wishes to conform to LIA-3 for that operation or parameter. For each of the marked items a suggested identifier is provided.

The Ada datatype **Boolean** corresponds to the LIA datatype **Boolean**.

Every implementation of Ada has at least one integer datatype, and at least one floating point datatype. The notations *INT* and *FLT* are used to stand for the names of one of these datatypes (respectively) in what follows.

Ada has imaginary and complex floating point datatypes. Standard Ada does not have any imaginary and complex integer datatypes, but such can be added as libraries (including infix syntax for certain operators). The notations *IINT* and *IFLT* are used to stand for the names of one of the imaginary datatypes (for integers and floating point respectively), and the notations *CINT* and *CFLT* are used to stand for the names of one of the complex datatypes (for integers and floating point respectively).

Ada has an overloading system, so that the same name can be used for different types of arguments to the operations. However, the constants *i* and *j* are defined in a generic package, but do not have any arguments whose types can determine which package instance is used (unless there is only one instance). The constants *i* and *j* may therefore need to be qualified with which package instance is used, even if that package is ‘open’.

The LIA-3 integer operations are listed below, along with the syntax used to invoke them:

$eq_{i(I)}(x, y)$	$x = y$	†
$eq_{I,i(I)}(x, y)$	$x = y$	†
$eq_{i(I),I}(x, y)$	$x = y$	†
$eq_{I,c(I)}(x, y)$	$x = y$	†
$eq_{c(I),I}(x, y)$	$x = y$	†
$eq_{i(I),c(I)}(x, y)$	$x = y$	†
$eq_{c(I),i(I)}(x, y)$	$x = y$	†
$eq_{c(I)}(x, y)$	$x = y$	†
$neq_{i(I)}(x, y)$	$x \neq y$	†
$neq_{I,i(I)}(x, y)$	$x \neq y$	†
$neq_{i(I),I}(x, y)$	$x \neq y$	†
$neq_{I,c(I)}(x, y)$	$x \neq y$	†
$neq_{c(I),I}(x, y)$	$x \neq y$	†
$neq_{i(I),c(I)}(x, y)$	$x \neq y$	†
$neq_{c(I),i(I)}(x, y)$	$x \neq y$	†
$neq_{c(I)}(x, y)$	$x \neq y$	†
$lss_{i(I)}(x, y)$	$x < y$	†
$leq_{i(I)}(x, y)$	$x \leq y$	†
$gtr_{i(I)}(x, y)$	$x > y$	†
$geq_{i(I)}(x, y)$	$x \geq y$	†
$itimes_{I \rightarrow i(I)}(x)$	$ii * x$	†
$itimes_{i(I) \rightarrow I}(x)$	$ii * x$	†
$itimes_{c(I)}(x)$	$ii * x$	†
$re_I(x)$	$Re(x)$	†
$re_{i(I)}(x)$	$Re(x)$	†
$re_{c(I)}(x)$	$Re(x)$	†
$im_I(x)$	$Im(x)$	†
$im_{i(I)}(x)$	$Im(x)$	†
$im_{c(I)}(x)$	$Im(x)$	†
$plusitimes_I(x, y)$	$Compose_From_Cartesian(x, y)$ or $x + ii * y$	†
$neg_{i(I)}(x)$	$-x$	†
$neg_{c(I)}(x)$	$-x$	†
$conj_I(x)$	$Conjugate(x)$	†
$conj_{i(I)}(x)$	$Conjugate(x)$	†
$conj_{c(I)}(x)$	$Conjugate(x)$	†

$add_{i(I)}(x, y)$	$x + y$	†
$add_{I,i(I)}(x, y)$	$x + y$	†
$add_{i(I),I}(x, y)$	$x + y$	†
$add_{I,c(I)}(x, y)$	$x + y$	†
$add_{c(I),I}(x, y)$	$x + y$	†
$add_{i(I),c(I)}(x, y)$	$x + y$	†
$add_{c(I),i(I)}(x, y)$	$x + y$	†
$add_{c(I)}(x, y)$	$x + y$	†
$sub_{i(I)}(x, y)$	$x - y$	†
$sub_{I,i(I)}(x, y)$	$x - y$	†
$sub_{i(I),I}(x, y)$	$x - y$	†
$sub_{I,c(I)}(x, y)$	$x - y$	†
$sub_{c(I),I}(x, y)$	$x - y$	†
$sub_{i(I),c(I)}(x, y)$	$x - y$	†
$sub_{c(I),i(I)}(x, y)$	$x - y$	†
$sub_{c(I)}(x, y)$	$x - y$	†
$mul_{i(I)}(x, y)$	$x * y$	†
$mul_{I,i(I)}(x, y)$	$x * y$	†
$mul_{i(I),I}(x, y)$	$x * y$	†
$mul_{I,c(I)}(x, y)$	$x * y$	†
$mul_{c(I),I}(x, y)$	$x * y$	†
$mul_{i(I),c(I)}(x, y)$	$x * y$	†
$mul_{c(I),i(I)}(x, y)$	$x * y$	†
$mul_{c(I)}(x, y)$	$x * y$	†
$abs_{i(I)}(x)$	$abs\ x$	†
$signum_I(x)$	$Signum(x)$	†
$signum_{i(I)}(x)$	$Signum(x)$	†
$divides_{i(I)}(x, y)$	$Divides(x, y)$	†
$divides_{I,i(I)}(x, y)$	$Divides(x, y)$	†
$divides_{i(I),I}(x, y)$	$Divides(x, y)$	†
$divides_{I,c(I)}(x, y)$	$Divides(x, y)$	†
$divides_{c(I),I}(x, y)$	$Divides(x, y)$	†
$divides_{i(I),c(I)}(x, y)$	$Divides(x, y)$	†
$divides_{c(I),i(I)}(x, y)$	$Divides(x, y)$	†
$divides_{c(I)}(x, y)$	$Divides(x, y)$	†
$quot_{i(I)}(x, y)$	$Quotient(x, y)$	†
$quot_{I,i(I)}(x, y)$	$Quotient(x, y)$	†
$quot_{i(I),I}(x, y)$	$Quotient(x, y)$	†
$quot_{I,c(I)}(x, y)$	$Quotient(x, y)$	†
$quot_{c(I),I}(x, y)$	$Quotient(x, y)$	†
$quot_{i(I),c(I)}(x, y)$	$Quotient(x, y)$	†
$quot_{c(I),i(I)}(x, y)$	$Quotient(x, y)$	†

$quot_{c(I)}(x, y)$	Quotient(x, y)	†
$mod_{i(I)}(x, y)$	$x \bmod y$	†
$mod_{I,i(I)}(x, y)$	$x \bmod y$	†
$mod_{i(I),I}(x, y)$	$x \bmod y$	†
$mod_{I,c(I)}(x, y)$	$x \bmod y$	†
$mod_{c(I),I}(x, y)$	$x \bmod y$	†
$mod_{i(I),c(I)}(x, y)$	$x \bmod y$	†
$mod_{c(I),i(I)}(x, y)$	$x \bmod y$	†
$mod_{c(I)}(x, y)$	$x \bmod y$	†
$ratio_{i(I)}(x, y)$	Ratio(x, y)	†
$ratio_{I,i(I)}(x, y)$	Ratio(x, y)	†
$ratio_{i(I),I}(x, y)$	Ratio(x, y)	†
$ratio_{I,c(I)}(x, y)$	Ratio(x, y)	†
$ratio_{c(I),I}(x, y)$	Ratio(x, y)	†
$ratio_{i(I),c(I)}(x, y)$	Ratio(x, y)	†
$ratio_{c(I),i(I)}(x, y)$	Ratio(x, y)	†
$ratio_{c(I)}(x, y)$	Ratio(x, y)	†
$residue_{i(I)}(x, y)$	Residue(x, y)	†
$residue_{I,i(I)}(x, y)$	Residue(x, y)	†
$residue_{i(I),I}(x, y)$	Residue(x, y)	†
$residue_{I,c(I)}(x, y)$	Residue(x, y)	†
$residue_{c(I),I}(x, y)$	Residue(x, y)	†
$residue_{i(I),c(I)}(x, y)$	Residue(x, y)	†
$residue_{c(I),i(I)}(x, y)$	Residue(x, y)	†
$residue_{c(I)}(x, y)$	Residue(x, y)	†
$group_{i(I)}(x, y)$	Group(x, y)	†
$group_{I,i(I)}(x, y)$	Group(x, y)	†
$group_{i(I),I}(x, y)$	Group(x, y)	†
$group_{I,c(I)}(x, y)$	Group(x, y)	†
$group_{c(I),I}(x, y)$	Group(x, y)	†
$group_{i(I),c(I)}(x, y)$	Group(x, y)	†
$group_{c(I),i(I)}(x, y)$	Group(x, y)	†
$group_{c(I)}(x, y)$	Group(x, y)	†
$pad_{i(I)}(x, y)$	Pad(x, y)	†
$pad_{I,i(I)}(x, y)$	Pad(x, y)	†
$pad_{i(I),I}(x, y)$	Pad(x, y)	†
$pad_{I,c(I)}(x, y)$	Pad(x, y)	†
$pad_{c(I),I}(x, y)$	Pad(x, y)	†
$pad_{i(I),c(I)}(x, y)$	Pad(x, y)	†
$pad_{c(I),i(I)}(x, y)$	Pad(x, y)	†
$pad_{c(I)}(x, y)$	Pad(x, y)	†

$max_{i(I)}(x, y)$	$IINT'Max(x, y)$	†
$min_{i(I)}(x, y)$	$IINT'Min(x, y)$	†
$max_seq_{i(I)}(xs)$	$Max(xs)$	†
$min_seq_{i(I)}(xs)$	$Min(xs)$	†

where x and y are expressions of type INT , $IINT$, or $CINT$ as appropriate, and xs is an expression of type **array** (Integer range <>) of $IINT$.

The parameters for LIA-3 operations approximating real complex valued functions can be accessed by the following syntax:

$box_error_mode_mul_{c(F)}$	$Box_Err_Mul(x)$	†
$max_err_mul_{c(F)}$	$Err_Mul(x)$	†
$box_error_mode_div_{c(F)}$	$Box_Err_Div(x)$	†
$max_err_div_{c(F)}$	$Err_Div(x)$	†
$max_err_exp_{c(F)}$	$Err_Exp(x)$	†
$max_err_power_{c(F)}$	$Err_Power(x)$	†
$max_err_sin_{c(F)}$	$Err_Sin(x)$ or $Err_Sinh(x)$	†
$max_err_tan_{c(F)}$	$Err_Tan(x)$ or $Err_Tanh(x)$	†

where x is an expression of type $CFLT$. The parameter functions are constant for each type (and library); the argument is used only to differentiate among the floating point types for the overloading resolution.

The LIA-3 basic complex floating point operations are listed below, along with the syntax used to invoke them:

$eq_{i(F)}(x, y)$	$x = y$	★
$eq_{F,i(F)}(x, y)$	$x = y$	†
$eq_{i(F),F}(x, y)$	$x = y$	†
$eq_{F,c(F)}(x, y)$	$x = y$	†
$eq_{c(F),F}(x, y)$	$x = y$	†
$eq_{i(F),c(F)}(x, y)$	$x = y$	†
$eq_{c(F),i(F)}(x, y)$	$x = y$	†
$eq_{c(F)}(x, y)$	$x = y$	★
$neq_{i(F)}(x, y)$	$x \neq y$	★
$neq_{F,i(F)}(x, y)$	$x \neq y$	†
$neq_{i(F),F}(x, y)$	$x \neq y$	†
$neq_{F,c(F)}(x, y)$	$x \neq y$	†
$neq_{c(F),F}(x, y)$	$x \neq y$	†
$neq_{i(F),c(F)}(x, y)$	$x \neq y$	†
$neq_{c(F),i(F)}(x, y)$	$x \neq y$	†
$neq_{c(F)}(x, y)$	$x \neq y$	★
$lss_{i(F)}(x, y)$	$x < y$	★
$leq_{i(F)}(x, y)$	$x \leq y$	★
$gtr_{i(F)}(x, y)$	$x > y$	★
$geq_{i(F)}(x, y)$	$x \geq y$	★

$itimes_{F \rightarrow i(F)}(x)$	$i * x$ or $j * x$	*
$itimes_{i(F) \rightarrow F}(x)$	$i * x$ or $j * x$	*
$itimes_{c(F)}(x)$	$i * x$ or $j * x$	*
$re_F(x)$	$\text{Re}(x)$	†
$re_{i(F)}(x)$	$\text{Re}(x)$	†
$re_{c(F)}(x)$	$\text{Re}(x)$	*
$im_F(x)$	$\text{Im}(x)$	†
$im_{i(F)}(x)$	$\text{Im}(x)$	*
$im_{c(F)}(x)$	$\text{Im}(x)$	*
$plusitimes_F(x, y)$	$\text{Compose_From_Cartesian}(x, y)$ or $x + i * y$ or $x + j * y$	*
$neg_{i(F)}(x)$	$-x$	*
$neg_{c(F)}(x)$	$-x$	*
$conj_F(x)$	$\text{Conjugate}(x)$	†
$conj_{i(F)}(x)$	$\text{Conjugate}(x)$	*
$conj_{c(F)}(x)$	$\text{Conjugate}(x)$	*
$add_{i(F)}(x, y)$	$x + y$	*
$add_{F, i(F)}(x, y)$	$x + y$	*
$add_{i(F), F}(x, y)$	$x + y$	*
$add_{F, c(F)}(x, y)$	$x + y$	*
$add_{c(F), F}(x, y)$	$x + y$	*
$add_{i(F), c(F)}(x, y)$	$x + y$	*
$add_{c(F), i(F)}(x, y)$	$x + y$	*
$add_{c(F)}(x, y)$	$x + y$	*
$sub_{i(F)}(x, y)$	$x - y$	*
$sub_{F, i(F)}(x, y)$	$x - y$	*
$sub_{i(F), F}(x, y)$	$x - y$	*
$sub_{F, c(F)}(x, y)$	$x - y$	*
$sub_{c(F), F}(x, y)$	$x - y$	*
$sub_{i(F), c(F)}(x, y)$	$x - y$	*
$sub_{c(F), i(F)}(x, y)$	$x - y$	*
$sub_{c(F)}(x, y)$	$x - y$	*
$mul_{i(F)}(x, y)$	$x * y$	*
$mul_{F, i(F)}(x, y)$	$x * y$	*
$mul_{i(F), F}(x, y)$	$x * y$	*
$mul_{F, c(F)}(x, y)$	$x * y$	*
$mul_{c(F), F}(x, y)$	$x * y$	*
$mul_{i(F), c(F)}(x, y)$	$x * y$	*
$mul_{c(F), i(F)}(x, y)$	$x * y$	*
$mul_{c(F)}(x, y)$	$x * y$	*
$div_{i(F)}(x, y)$	x / y	*

$div_{F,i(F)}(x, y)$	x / y	★
$div_{i(F),F}(x, y)$	x / y	★
$div_{F,c(F)}(x, y)$	x / y	★
$div_{c(F),F}(x, y)$	x / y	★
$div_{i(F),c(F)}(x, y)$	x / y	★
$div_{c(F),i(F)}(x, y)$	x / y	★
$div_{c(F)}(x, y)$	x / y	★
$abs_{i(F)}(x)$	$abs\ x$	★
$abs_{c(F)}(x)$	$abs\ x\ \text{ or } \text{Modulus}(x)$	★
$phase_F(x)$	$\text{Argument}(x)$	†
$phase_{i(F)}(x)$	$\text{Argument}(x)$	†
$phase_{c(F)}(x)$	$\text{Argument}(x)$	★
$phaseu_F(u, x)$	$\text{Argument}(x, u)$	†
$phaseu_{i(F)}(u, x)$	$\text{Argument}(x, u)$	†
$phaseu_{c(F)}(u, x)$	$\text{Argument}(x, u)$	★
$signum_F(x)$	$\text{Signum}(x)$	†
$signum_{i(F)}(x)$	$\text{Signum}(x)$	†
$signum_{c(F)}(x)$	$\text{Signum}(x)$	†
$floor_{i(F)}(x)$	$IFLT'\text{Floor}(x)$	†
$floor_{c(F)}(x)$	$CFLT'\text{Floor}(x)$	†
$rounding_{i(F)}(x)$	$IFLT'\text{UnbiasedRounding}(x)$	†
$rounding_{c(F)}(x)$	$CFLT'\text{UnbiasedRounding}(x)$	†
$ceiling_{i(F)}(x)$	$IFLT'\text{Ceiling}(x)$	†
$ceiling_{c(F)}(x)$	$CFLT'\text{Ceiling}(x)$	†
$max_{i(F)}(x, y)$	$IFLT'\text{Max}(x, y)$	†
$mmax_{i(F)}(x, y)$	$Mmax(x, y)$	†
$min_{i(F)}(x, y)$	$IFLT'\text{Min}(x, y)$	†
$mmin_{i(F)}(x, y)$	$Mmin(x, y)$	†
$max_seq_{i(F)}(xs)$	$\text{Max}(xs)$	†
$mmax_seq_{i(F)}(xs)$	$Mmax(xs)$	†
$min_seq_{i(F)}(xs)$	$\text{Min}(xs)$	†
$mmin_seq_{i(F)}(xs)$	$Mmin(xs)$	†
$polar_F(x, y)$	$\text{Compose_From_Polar}(x, y)$	★
$polaru_F(u, x, y)$	$\text{Compose_From_Polar}(x, y, u)$	★

where x and y are expressions of type FLT , $IFLT$, or $CFLT$ as appropriate, and u is an expression of type FLT .

The LIA-3 elementary floating point operations are listed below, along with the syntax used to invoke them:

$power_{i(F),I}(x, y)$	$x ** y$	★
$power_{c(F),I}(x, y)$	$x ** y$	★ Not LIA-3

$exp_{i(F) \rightarrow c(F)}(x)$	<code>Exp(x)</code>	★
$exp_{c(F)}(x)$	<code>Exp(x)</code>	★
$power_{F \rightarrow c(F)}(x, y)$	<code>Compose_From_Cartesian(x) ** y</code>	★
$power_{c(F)}(x, y)$	<code>x ** y</code>	★ Not LIA-3
$sqr_{F \rightarrow c(F)}(x)$	<code>Sqrtc(x)</code>	†
$sqr_{i(F) \rightarrow c(F)}(x)$	<code>Sqrt(x)</code>	†
$sqr_{c(F)}(x)$	<code>Sqrt(x)</code>	★
$ln_{F \rightarrow c(F)}(x)$	<code>Logc(x)</code> , or <code>Log(abs x) + i*Arctan(Im(x), x)</code>	† ★
$ln_{i(F) \rightarrow c(F)}(x)$	<code>Log(x)</code>	†
$ln_{i(F) \rightarrow c(F)}(x)$	<code>Log(abs x) + i*Arctan(Im(x), Re(x))</code>	★
$ln_{c(F)}(x)$	<code>Log(x)</code>	★
$rad_{i(F)}(x)$	<code>Rad(x)</code>	†
$rad_{c(F)}(x)$	<code>Rad(x)</code>	†
$sin_{i(F)}(x)$	<code>Sin(x)</code> or <code>i * Sinh(Im(x))</code>	† ★
$sin_{c(F)}(x)$	<code>Sin(x)</code>	★
$cos_{i(F)}(x)$	<code>Cos(x)</code> or <code>Cosh(Im(x))</code>	† ★
$cos_{c(F)}(x)$	<code>Cos(x)</code>	★
$tan_{i(F)}(x)$	<code>Tan(x)</code> or <code>i * Tanh(Im(x))</code>	† ★
$tan_{c(F)}(x)$	<code>Tan(x)</code>	★
$cot_{i(F)}(x)$	<code>Cot(x)</code> or <code>-i * Coth(Im(x))</code>	† ★
$cot_{c(F)}(x)$	<code>Cot(x)</code>	★
$sec_{i(F)}(x)$	<code>Sec(x)</code>	†
$sec_{c(F)}(x)$	<code>Sec(x)</code>	†
$csc_{i(F)}(x)$	<code>Csc(x)</code>	†
$csc_{c(F)}(x)$	<code>Csc(x)</code>	†
$arcsin_{F \rightarrow c(F)}(x)$	<code>Arcsinc(x)</code>	†
$arcsin_{i(F)}(x)$	<code>Arcsin(x)</code>	†
$arcsin_{c(F)}(x)$	<code>Arcsin(x)</code>	★
$arccos_{F \rightarrow c(F)}(x)$	<code>Arccosc(x)</code>	†
$arccos_{i(F) \rightarrow c(F)}(x)$	<code>Arccos(x)</code>	†
$arccos_{c(F)}(x)$	<code>Arccos(x)</code>	★
$arctan_{i(F)}(x)$	<code>Arctan(x)</code>	†
$arctan_{i(F) \rightarrow c(F)}(x)$	<code>Arctanc(x)</code>	†
$arctan_{c(F)}(x)$	<code>Arctan(x)</code>	★
$arccot_{i(F)}(x)$	<code>Arccot(x)</code>	†

$arccot_{i(F) \rightarrow c(F)}(x)$	$Arccotc(x)$	†
$arccot_{c(F)}(x)$	$Arccot(x)$	★
$arcsec_{F \rightarrow c(F)}(x)$	$Arcsecc(x)$	†
$arcsec_{i(F) \rightarrow c(F)}(x)$	$Arcsecc(x)$	†
$arcsec_{c(F)}(x)$	$Arcsec(x)$	†
$arccsc_{F \rightarrow c(F)}(x)$	$Arccsc(x)$	†
$arccsc_{i(F)}(x)$	$Arccsc(x)$	†
$arccsc_{c(F)}(x)$	$Arccsc(x)$	†
$radh_F(x)$	$Radh(x)$	†
$radh_{i(F)}(x)$	$Radh(x)$	†
$radh_{c(F)}(x)$	$Radh(x)$	†
$sinh_{i(F)}(x)$	$Sinh(x)$	†
$sinh_{c(F)}(x)$	$Sinh(x)$	★
$cosh_{i(F)}(x)$	$Cosh(x)$	†
$cosh_{c(F)}(x)$	$Cosh(x)$	★
$tanh_{i(F)}(x)$	$Tanh(x)$	†
$tanh_{c(F)}(x)$	$Tanh(x)$	★
$coth_{i(F)}(x)$	$Coth(x)$	†
$coth_{c(F)}(x)$	$Coth(x)$	★
$sech_{i(F)}(x)$	$Sech(x)$	†
$sech_{c(F)}(x)$	$Sech(x)$	†
$csch_{i(F)}(x)$	$Csch(x)$	†
$csch_{c(F)}(x)$	$Csch(x)$	†
$arcsinh_{i(F)}(x)$	$Arcsinh(x)$	†
$arcsinh_{i(F) \rightarrow c(F)}(x)$	$Arcsinhc(x)$	†
$arcsinh_{c(F)}(x)$	$ArcSinh(x)$	★
$arccosh_{F \rightarrow c(F)}(x)$	$Arccoshc(x)$	†
$arccosh_{i(F) \rightarrow c(F)}(x)$	$Arccoshc(x)$	†
$arccosh_{c(F)}(x)$	$Arccosh(x)$	★
$arctanh_{F \rightarrow c(F)}(x)$	$Arctanhc(x)$	†
$arctanh_{i(F)}(x)$	$Arctanh(x)$	†
$arctanh_{c(F)}(x)$	$Arctanh(x)$	★
$arcoth_{F \rightarrow c(F)}(x)$	$Arcothc(x)$	†
$arcoth_{i(F)}(x)$	$Arcoth(x)$	†
$arcoth_{c(F)}(x)$	$Arcoth(x)$	★
$arcsech_{F \rightarrow c(F)}(x)$	$Arcsechc(x)$	†
$arcsech_{i(F) \rightarrow c(F)}(x)$	$Arcsechc(x)$	†
$arcsech_{c(F)}(x)$	$Arcsech(x)$	†
$arccsch_{i(F)}(x)$	$Arccsch(x)$	†
$arccsch_{i(F) \rightarrow c(F)}(x)$	$Arccschc(x)$	†
$arccsch_{c(F)}(x)$	$Arccsch(x)$	†

where x and y are expressions of type *FLT*, *IFLT*, or *CFLT* as appropriate.

Arithmetic value conversions in Ada are always explicit.

$convert_{I \rightarrow c(I)}(x)$	<code>Compose_From_Cartesian(x)</code>	†
$convert_{i(I) \rightarrow c(I)}(x)$	<code>Compose_From_Cartesian(x)</code>	†
$convert_{i(I) \rightarrow i(I')}(x)$	<code>ii * INT2(Im(x))</code>	†
$convert_{c(I) \rightarrow c(I')}(x)$	<code>Compose_From_Cartesian(INT2(Re(x)), INT2(Im(x)))</code>	†
$convert_{F \rightarrow c(F)}(x)$	<code>Compose_From_Cartesian(x)</code>	*
$convert_{F \rightarrow c(F)}(x)$	<code>x + i * Im(x) or x + j * Im(x)</code>	†
$convert_{i(F) \rightarrow c(F)}(x)$	<code>Compose_From_Cartesian(x)</code>	*
$convert_{i(F) \rightarrow c(F)}(x)$	<code>Re(x) + x</code>	†
$convert_{i(F) \rightarrow i(F')}(x)$	<code>i * FLT2(Im(x))</code>	*
$convert_{c(F) \rightarrow c(F')}(x)$	<code>Compose_From_Cartesian(FLT2(Re(x)), FLT2(Im(x)))</code>	*
$convert_{c(F) \rightarrow c(F')}(x)$	<code>FLT2(Re(x)) + i * FLT2(Im(x))</code>	*

where x is of type *IFLT* or *CFLT* as appropriate.

Ada has standard support for input of complex numerals and output of complex values, though not for imaginary values, both to/from strings as well as to/from files, both in “narrow” characters (Latin-1) and “wide” characters (UCS-2/UTF-16):

$convert_{c(F'') \rightarrow c(F)}(s)$	<code>Get(s, g, w?);</code>	*
$convert_{c(F'') \rightarrow c(F)}(f)$	<code>Get(f?, g, w?);</code>	*
$convert_{c(F) \rightarrow c(F'')}(y)$	<code>Put(s, y, Aft=>a?, Exp=>e?);</code>	*
$convert_{c(F) \rightarrow c(F'')}(y)$	<code>Put(h?, y, Fore=>i?, Aft=>a?, Exp=>e?);</code>	*
$convert_{c(D') \rightarrow c(F)}(s)$	<code>Get(s, g, w?);</code>	*
$convert_{c(D') \rightarrow c(F)}(f)$	<code>Get(f?, g, w?);</code>	*
$convert_{c(F) \rightarrow c(D')}(y)$	<code>Put(s, y, Aft=>a?, Exp=>0);</code>	*
$convert_{c(F) \rightarrow c(D')}(y)$	<code>Put(h?, y, Fore=>i?, Aft=>a?, Exp=>0);</code>	*

where y is an expression of type *CFLT*, $base$, w , i , a , and e are expressions for non-negative integers. e is greater than 0. $base$ is greater than 1. g is a variable of type *CFLT*. A ? above indicates that the parameter is optional. f is an opened input file (default is the default input file). h is an opened output file (default is the default output file). s is of type **String** or **Wide_String**. For **Get** of a floating point or fixed point numeral, the base is indicated in the numeral (default 10). For **Put** of a floating point or fixed point numeral, only base 10 is required to be supported. For details on **Get** and **Put**, see clause G.1.3 Complex Input-Output, and G.1.4, of ISO/IEC 8652:1995. The Ada standard is not explicit about the use of UCS-2/UTF-16 for input/output of “wide” characters. However, there is an implicit assumption of no data loss for data output and later input again. Though there are several standard external encodings that satisfy that (e.g. UTF-8), one may also assume that there is no encoding change during IO. Hence, big-endian UCS-2/UTF-16 is the only reasonable option (otherwise explicit mention in the Ada standard would have been necessary).

Ada does not have any special programming language syntax (as opposed to I/O syntax) for complex, or imaginary, numerals in general. Instead imaginary and complex numerals are expressed as (compile-time evaluable) expressions involving a notation for the imaginary unit.

$imaginary_unit_{i(I)}$	<code>ii</code>	†
$imaginary_unit_{c(I)}$	<code>0+ii</code>	†
$imaginary_unit_{i(F)}$	<code>i or j</code>	*
$imaginary_unit_{c(F)}$	<code>0+i or 0+j</code>	*

C.2 C

The programming language C is defined by ISO/IEC 9899:1999, *Information technology – Programming languages – C* [4].

An implementation should follow all the requirements of LIA-3 unless otherwise specified by this language binding.

The operations or parameters marked “†” are not part of the language and should be provided by an implementation that wishes to conform to the LIA-3 for that operation. For each of the marked items a suggested identifier is provided.

The LIA datatype **Boolean** is implemented by the C datatype **int** (1 = **true** and 0 = **false**), or the new **_Bool** datatype (usually referred to via the macro **bool**).

C defines numerous integer datatypes. They may be aliases of each other in an implementation defined way. The description here is not complete. See the C standard. Some of the integer datatypes have a predetermined bit width, and the signed ones use 2's complement for representation of negative values: **int_n.t** and **uint_n.t**, where *n* is the bit width expressed as a decimal numeral. Some bit widths are required. There are also minimum width, fastest minimum width, and special purpose integer datatypes (like **size_t**). Also provided are the more well-known integer datatypes **char**, **short int**, **int**, **long int**, **long long int**, **unsigned char**, **unsigned short int**, **unsigned int**, **unsigned long int**, and **unsigned long long int**. Finally there are the integer datatypes **intmax_t** and **uintmax_t** that are the largest provided signed and unsigned integer datatypes. **intmax_t** and **uintmax_t** may even be unbounded with a negative integer infinity for the signed variety as **INTMAX_MIN** and a positive integer infinity as **INTMAX_MAX** and **UINTMAX_MAX**. **INT** is used below to designate one of the integer datatypes.

C has no standard imaginary integer or complex integer datatypes. However, in the suggested binding below, **IINT** and **CINT** are used to designate such potential datatypes.

C names three floating point datatypes: **float**, **double**, and **long double**. **FLT** is used below to designate one of the floating point datatypes.

C has three complex floating point datatypes: **float _Complex**, **double _Complex**, and **long double _Complex**, and *may* have three imaginary floating point datatypes: **float _Imaginary**, **double _Imaginary**, and **long double _Imaginary**. **_Complex** and **_Imaginary** are usually written **complex** and **imaginary** respectively. **CFLT** and **IFLT** are used below to designate one of the complex or imaginary (respectively) floating point datatypes.

For some of the operations below, the C standard defines ‘type generic macros’. These are fixed by the C standard, and new ones cannot be defined in program libraries.

The operations marked “(★)” are available only when the implementation provides imaginary floating point datatypes.

The LIA-3 integer operations are listed below, along with the syntax used to invoke them (the binding for **II** is given towards the end of this binding example):

$eq_{i(I)}(x, y)$	$x == y$	†
$eq_{I,i(I)}(x, y)$	$x == y$	†
$eq_{i(I),I}(x, y)$	$x == y$	†
$eq_{I,c(I)}(x, y)$	$x == y$	†
$eq_{c(I),I}(x, y)$	$x == y$	†
$eq_{i(I),c(I)}(x, y)$	$x == y$	†
$eq_{c(I),i(I)}(x, y)$	$x == y$	†

$eq_{c(I)}(x, y)$	$x == y$	†
$neq_{i(I)}(x, y)$	$x != y$	†
$neq_{I,i(I)}(x, y)$	$x != y$	†
$neq_{i(I),I}(x, y)$	$x != y$	†
$neq_{I,c(I)}(x, y)$	$x != y$	†
$neq_{c(I),I}(x, y)$	$x != y$	†
$neq_{i(I),c(I)}(x, y)$	$x != y$	†
$neq_{c(I),i(I)}(x, y)$	$x != y$	†
$neq_{c(I)}(x, y)$	$x != y$	†
$lss_{i(I)}(x, y)$	$x < y$	†
$leq_{i(I)}(x, y)$	$x <= y$	†
$gtr_{i(I)}(x, y)$	$x > y$	†
$geq_{i(I)}(x, y)$	$x >= y$	†
$itimes_{I \rightarrow i(I)}(x)$	$II * x$	†
$itimes_{i(I) \rightarrow I}(x)$	$II * x$	†
$itimes_{c(I)}(x)$	$II * x$	†
$re_I(x)$	$crealit(x)$	†
$re_{i(I)}(x)$	$crealit(x)$	†
$re_{c(I)}(x)$	$crealit(x)$	†
$im_I(x)$	$cimagit(x)$	†
$im_{i(I)}(x)$	$cimagit(x)$	†
$im_{c(I)}(x)$	$cimagit(x)$	†
$plusitimes_I(x, y)$	$x + II * y$	†
$neg_{i(I)}(x)$	$-x$	†
$neg_{c(I)}(x)$	$-x$	†
$conj_I(x)$	$conjit(x)$	†
$conj_{i(I)}(x)$	$conjit(x)$	†
$conj_{c(I)}(x)$	$conjit(x)$	†
$add_{i(I)}(x, y)$	$x + y$	†
$add_{I,i(I)}(x, y)$	$x + y$	†
$add_{i(I),I}(x, y)$	$x + y$	†
$add_{I,c(I)}(x, y)$	$x + y$	†
$add_{c(I),I}(x, y)$	$x + y$	†
$add_{i(I),c(I)}(x, y)$	$x + y$	†
$add_{c(I),i(I)}(x, y)$	$x + y$	†
$add_{c(I)}(x, y)$	$x + y$	†
$sub_{i(I)}(x, y)$	$x - y$	†
$sub_{I,i(I)}(x, y)$	$x - y$	†
$sub_{i(I),I}(x, y)$	$x - y$	†
$sub_{I,c(I)}(x, y)$	$x - y$	†
$sub_{c(I),I}(x, y)$	$x - y$	†

$sub_{i(I),c(I)}(x, y)$	$x - y$	†
$sub_{c(I),i(I)}(x, y)$	$x - y$	†
$sub_{c(I)}(x, y)$	$x - y$	†
$mul_{i(I)}(x, y)$	$x * y$	†
$mul_{I,i(I)}(x, y)$	$x * y$	†
$mul_{i(I),I}(x, y)$	$x * y$	†
$mul_{I,c(I)}(x, y)$	$x * y$	†
$mul_{c(I),I}(x, y)$	$x * y$	†
$mul_{i(I),c(I)}(x, y)$	$x * y$	†
$mul_{c(I),i(I)}(x, y)$	$x * y$	†
$mul_{c(I)}(x, y)$	$x * y$	†
$abs_{i(I)}(x)$	$iabst(x)$	†
$signum_I(x)$	$signumt(x)$	†
$signum_{i(I)}(x)$	$signumt(x)$	†
$divides_{i(I)}(x, y)$	$dividest(x, y)$	†
$divides_{I,i(I)}(x, y)$	$dividest(x, y)$	†
$divides_{i(I),I}(x, y)$	$dividest(x, y)$	†
$divides_{I,c(I)}(x, y)$	$dividest(x, y)$	†
$divides_{c(I),I}(x, y)$	$dividest(x, y)$	†
$divides_{i(I),c(I)}(x, y)$	$dividest(x, y)$	†
$divides_{c(I),i(I)}(x, y)$	$dividest(x, y)$	†
$divides_{c(I)}(x, y)$	$dividest(x, y)$	†
$quot_{i(I)}(x, y)$	$quott(x, y)$	†
$quot_{I,i(I)}(x, y)$	$quott(x, y)$	†
$quot_{i(I),I}(x, y)$	$quott(x, y)$	†
$quot_{I,c(I)}(x, y)$	$quott(x, y)$	†
$quot_{c(I),I}(x, y)$	$quott(x, y)$	†
$quot_{i(I),c(I)}(x, y)$	$quott(x, y)$	†
$quot_{c(I),i(I)}(x, y)$	$quott(x, y)$	†
$quot_{c(I)}(x, y)$	$quott(x, y)$	†
$mod_{i(I)}(x, y)$	$modt(x, y)$	†
$mod_{I,i(I)}(x, y)$	$modt(x, y)$	†
$mod_{i(I),I}(x, y)$	$modt(x, y)$	†
$mod_{I,c(I)}(x, y)$	$modt(x, y)$	†
$mod_{c(I),I}(x, y)$	$modt(x, y)$	†
$mod_{i(I),c(I)}(x, y)$	$modt(x, y)$	†
$mod_{c(I),i(I)}(x, y)$	$modt(x, y)$	†
$mod_{c(I)}(x, y)$	$modt(x, y)$	†
$ratio_{i(I)}(x, y)$	$ratio_t(x, y)$	†
$ratio_{I,i(I)}(x, y)$	$ratio_t(x, y)$	†
$ratio_{i(I),I}(x, y)$	$ratio_t(x, y)$	†

$ratio_{I,c(I)}(x, y)$	$ratio_t(x, y)$	†
$ratio_{c(I),I}(x, y)$	$ratio_t(x, y)$	†
$ratio_{i(I),c(I)}(x, y)$	$ratio_t(x, y)$	†
$ratio_{c(I),i(I)}(x, y)$	$ratio_t(x, y)$	†
$ratio_{c(I)}(x, y)$	$ratio_t(x, y)$	†
$residue_{i(I)}(x, y)$	$iremaindert(x, y)$	†
$residue_{I,i(I)}(x, y)$	$iremaindert(x, y)$	†
$residue_{i(I),I}(x, y)$	$iremaindert(x, y)$	†
$residue_{I,c(I)}(x, y)$	$iremaindert(x, y)$	†
$residue_{c(I),I}(x, y)$	$iremaindert(x, y)$	†
$residue_{i(I),c(I)}(x, y)$	$iremaindert(x, y)$	†
$residue_{c(I),i(I)}(x, y)$	$iremaindert(x, y)$	†
$residue_{c(I)}(x, y)$	$iremaindert(x, y)$	†
$group_{i(I)}(x, y)$	$group_t(x, y)$	†
$group_{I,i(I)}(x, y)$	$group_t(x, y)$	†
$group_{i(I),I}(x, y)$	$group_t(x, y)$	†
$group_{I,c(I)}(x, y)$	$group_t(x, y)$	†
$group_{c(I),I}(x, y)$	$group_t(x, y)$	†
$group_{i(I),c(I)}(x, y)$	$group_t(x, y)$	†
$group_{c(I),i(I)}(x, y)$	$group_t(x, y)$	†
$group_{c(I)}(x, y)$	$group_t(x, y)$	†
$pad_{i(I)}(x, y)$	$pad_t(x, y)$	†
$pad_{I,i(I)}(x, y)$	$pad_t(x, y)$	†
$pad_{i(I),I}(x, y)$	$pad_t(x, y)$	†
$pad_{I,c(I)}(x, y)$	$pad_t(x, y)$	†
$pad_{c(I),I}(x, y)$	$pad_t(x, y)$	†
$pad_{i(I),c(I)}(x, y)$	$pad_t(x, y)$	†
$pad_{c(I),i(I)}(x, y)$	$pad_t(x, y)$	†
$pad_{c(I)}(x, y)$	$pad_t(x, y)$	†
$max_{i(I)}(x, y)$	$imax_t(x, y)$	†
$min_{i(I)}(x, y)$	$imint(x, y)$	†
$max_seq_{i(I)}(xs)$	$imax_arrt(xs, nr_of_items)$	†
$min_seq_{i(I)}(xs)$	$imax_arrt(xs, nr_of_items)$	†

where x and y are expressions of type INT , $IINT$, or $CINT$ as appropriate, and xs is an expression of type array of $IINT$, and t is a string (part of the operation name in C), that is the empty string for `int`, is `l` for `long int`, is `u` for `unsigned int`, and is `ul` for `unsigned long int`.

The parameters for LIA-3 operations approximating complex real valued transcendental functions can be accessed by the following syntax:

$box_error_mode_mul_{c(F)}$	box_err_cmult	†
$max_err_mul_{c(F)}$	err_cmult	†
$box_error_mode_div_{c(F)}$	box_err_cdivt	†
$max_err_div_{c(F)}$	err_cdivt	†

$max_err_exp_{c(F)}$	<code>err_cexpt</code>	†
$max_err_power_{c(F)}$	<code>err_cpowers</code>	†
$max_err_sin_{c(F)}$	<code>err_csint</code>	†
$max_err_tan_{c(F)}$	<code>err_ctant</code>	†

where t is a string (part of the parameter name in C), that is the empty string for `int`, is `l` for `long int`, is `u` for `unsigned int`, and is `ul` for `unsigned long int`.

The LIA-3 non-transcendental complex floating point operations are listed below, along with the syntax used to invoke them (the binding for `I` is given towards the end of this binding example):

$eq_{i(F)}(x, y)$	<code>x == y</code>	(*)
$eq_{F,i(F)}(x, y)$	<code>x == y</code>	(*)
$eq_{i(F),F}(x, y)$	<code>x == y</code>	(*)
$eq_{F,c(F)}(x, y)$	<code>x == y</code>	*
$eq_{c(F),F}(x, y)$	<code>x == y</code>	*
$eq_{i(F),c(F)}(x, y)$	<code>x == y</code>	(*)
$eq_{c(F),i(F)}(x, y)$	<code>x == y</code>	(*)
$eq_{c(F)}(x, y)$	<code>x == y</code>	*
$neq_{i(F)}(x, y)$	<code>x != y</code>	(*)
$neq_{F,i(F)}(x, y)$	<code>x != y</code>	(*)
$neq_{i(F),F}(x, y)$	<code>x != y</code>	(*)
$neq_{F,c(F)}(x, y)$	<code>x != y</code>	*
$neq_{c(F),F}(x, y)$	<code>x != y</code>	*
$neq_{i(F),c(F)}(x, y)$	<code>x != y</code>	(*)
$neq_{c(F),i(F)}(x, y)$	<code>x != y</code>	(*)
$neq_{c(F)}(x, y)$	<code>x != y</code>	*
$lss_{i(F)}(x, y)$	<code>x < y</code>	†
$leq_{i(F)}(x, y)$	<code>x <= y</code>	†
$gtr_{i(F)}(x, y)$	<code>x > y</code>	†
$geq_{i(F)}(x, y)$	<code>x >= y</code>	†
$itimes_{F \rightarrow i(F)}(x)$	<code>I * x</code>	(*)
$itimes_{i(F) \rightarrow F}(x)$	<code>I * x</code>	(*)
$itimes_{c(F)}(x)$	<code>I * x</code>	*
$re_F(x)$	<code>crealt(x)</code>	†
$re_{i(F)}(x)$	<code>crealt(x)</code>	†
$re_{c(F)}(x)$	<code>crealt(x)</code> or <code>creal(x)</code>	*
$im_F(x)$	<code>cimagt(x)</code>	†
$im_{i(F)}(x)$	<code>cimagt(x)</code>	†
$im_{c(F)}(x)$	<code>cimagt(x)</code> or <code>cimag(x)</code>	*
$plusitimes_F(x, y)$	<code>x + I * y</code>	*
$neg_{i(F)}(x)$	<code>- x</code>	(*)
$neg_{c(F)}(x)$	<code>- x</code>	*
$conj_F(x)$	<code>conjt(x)</code> or <code>conj(x)</code>	†

$conj_{i(F)}(x)$	$conj_t(x)$	or	$conj(x)$	(★)
$conj_{c(F)}(x)$	$conj_t(x)$	or	$conj(x)$	★
$add_{i(F)}(x, y)$	$x + y$			(★)
$add_{F,i(F)}(x, y)$	$x + y$			(★)
$add_{i(F),F}(x, y)$	$x + y$			(★)
$add_{F,c(F)}(x, y)$	$x + y$			★
$add_{c(F),F}(x, y)$	$x + y$			★
$add_{i(F),c(F)}(x, y)$	$x + y$			(★)
$add_{c(F),i(F)}(x, y)$	$x + y$			(★)
$add_{c(F)}(x, y)$	$x + y$			★
$sub_{i(F)}(x, y)$	$x - y$			(★)
$sub_{F,i(F)}(x, y)$	$x - y$			(★)
$sub_{i(F),F}(x, y)$	$x - y$			(★)
$sub_{F,c(F)}(x, y)$	$x - y$			★
$sub_{c(F),F}(x, y)$	$x - y$			★
$sub_{i(F),c(F)}(x, y)$	$x - y$			(★)
$sub_{c(F),i(F)}(x, y)$	$x - y$			(★)
$sub_{c(F)}(x, y)$	$x - y$			★
$mul_{i(F)}(x, y)$	$x * y$			(★)
$mul_{F,i(F)}(x, y)$	$x * y$			(★)
$mul_{i(F),F}(x, y)$	$x * y$			(★)
$mul_{F,c(F)}(x, y)$	$x * y$			★
$mul_{c(F),F}(x, y)$	$x * y$			★
$mul_{i(F),c(F)}(x, y)$	$x * y$			(★)
$mul_{c(F),i(F)}(x, y)$	$x * y$			(★)
$mul_{c(F)}(x, y)$	$x * y$			★
$div_{i(F)}(x, y)$	x / y			(★)
$div_{F,i(F)}(x, y)$	x / y			(★)
$div_{i(F),F}(x, y)$	x / y			(★)
$div_{F,c(F)}(x, y)$	x / y			★
$div_{c(F),F}(x, y)$	x / y			★
$div_{i(F),c(F)}(x, y)$	x / y			(★)
$div_{c(F),i(F)}(x, y)$	x / y			(★)
$div_{c(F)}(x, y)$	x / y			★
$abs_{i(F)}(x)$	$fabs(x)$			(★)
$abs_{c(F)}(x)$	$cabst(x)$	or	$fabs(x)$	★
$phase_F(x)$	$cargt(x)$	or	$carg(x)$	†
$phase_{i(F)}(x)$	$carg(x)$			(★)
$phase_{c(F)}(x)$	$cargt(x)$	or	$carg(x)$	★
$phaseu_F(u, x)$	$cargut(x, u)$			†
$phaseu_{i(F)}(u, x)$	$cargut(x, u)$			†

$phaseu_{c(F)}(u, x)$	$cargut(x, u)$	†
$signum_F(x)$	$signbit(x)$	★
$signum_{i(F)}(x)$	$I*signbit(cimag(x))$	★
$signum_{c(F)}(x)$	$csignt(x)$	†
$floor_{i(F)}(x)$	$floort(x)$	†
$floor_{c(F)}(x)$	$floort(x)$	†
$rounding_{i(F)}(x)$	$nearbyintt(x)$	†
$rounding_{c(F)}(x)$	$nearbyintt(x)$	†
$ceiling_{i(F)}(x)$	$ceilingt(x)$	†
$ceiling_{c(F)}(x)$	$ceilingt(x)$	†
$max_{i(F)}(x, y)$	$nmaxt(x, y)$	†
$mma_{i(F)}(x, y)$	$fmaxt(x, y)$	†
$min_{i(F)}(x, y)$	$nmint(x, y)$	†
$mmin_{i(F)}(x, y)$	$fmint(x, y)$	†
$max_seq_{i(F)}(xs)$	$nmax_arrt(xs, nr_of_items)$	†
$mma_seq_{i(F)}(xs)$	$fmax_arrt(xs, nr_of_items)$	†
$min_seq_{i(F)}(xs)$	$nmin_arrt(xs, nr_of_items)$	†
$mmin_seq_{i(F)}(xs)$	$fmin_arrt(xs, nr_of_items)$	†
$polar_F(x, y)$	$polar_t(x, y)$	†
$polaru_F(u, x, y)$	$polarut(x, y, u)$	†
$proj_F(x)$	$projt(x)$ or $cproj(x)$	† Not LIA-3
$proj_{i(F)}(x)$	$cproj(x)$	† Not LIA-3
$proj_{c(F)}(x)$	$cprojt(x)$ or $cproj(x)$	★ Not LIA-3

where x and y are expressions of type *FLT*, *IFLT*, or *CFLT* as appropriate, u is an expression of type *FLT*, and t is a string (part of the operation name in C), that is the empty string for *double*, is “l” for *long double*, and is “f” for *float*.

The LIA-3 elementary complex floating point operations are listed below, along with the syntax used to invoke them:

$power_{i(F),I}(b, i)$	$cpowit(b, i)$	†
$exp_{i(F)}(x)$	$exp(x)$	(★)
$exp_{c(F)}(x)$	$cexpt(x)$ or $exp(x)$	★
$power_{F \rightarrow c(F)}(x, y)$	$powc(x, y)$	†
$power_{c(F)}(b, y)$	$cpowt(x, y)$ or $pow(x, y)$	★ Not LIA-3
$sqrt_{F \rightarrow c(F)}(x)$	$sqrctc(x)$	†
$sqrt_{i(F) \rightarrow c(F)}(x)$	$sqrt(x)$	(★)
$sqrt_{c(F)}(x)$	$csqrctt(x)$ or $sqrt(x)$	★
$ln_{F \rightarrow c(F)}(x)$	$logct(x)$	†

$\ln_{F \rightarrow c(F)}(x)$	$\log(\text{abs}(x)) + I * \text{atan2}(\text{cimag}(x), x)$	*
$\ln_{i(F) \rightarrow c(F)}(x)$	$\log(x)$	(*)
$\ln_{i(F) \rightarrow c(F)}(x)$	$\log(\text{abs}(x)) + I * \text{atan2}(\text{cimag}(x), \text{creal}(x))$	(*)
$\ln_{c(F)}(x)$	$\text{clogt}(x)$ or $\log(x)$	*
$\text{rad}_{i(F)}(x)$	$\text{radianht}(x)$	†
$\text{rad}_{c(F)}(x)$	$\text{radianht}(x)$	†
$\sin_{i(F)}(x)$	$\sin(x)$	(*)
$\sin_{c(F)}(x)$	$\text{csint}(x)$ or $\sin(x)$	*
$\cos_{i(F)}(x)$	$\cos(x)$	(*)
$\cos_{c(F)}(x)$	$\text{ccost}(x)$ or $\cos(x)$	*
$\tan_{i(F)}(x)$	$\tan(x)$	(*)
$\tan_{c(F)}(x)$	$\text{ctant}(x)$ or $\tan(x)$	*
$\cot_{i(F)}(x)$	$\cot(x)$	†
$\cot_{c(F)}(x)$	$\text{ccott}(x)$ or $\cot(x)$	†
$\sec_{i(F)}(x)$	$\sec(x)$	†
$\sec_{c(F)}(x)$	$\text{csect}(x)$ or $\sec(x)$	†
$\csc_{i(F)}(x)$	$\csc(x)$	†
$\csc_{c(F)}(x)$	$\text{ccsct}(x)$ or $\csc(x)$	†
$\arcsin_{F \rightarrow c(F)}(x)$	$\text{casinht}(x)$	†
$\arcsin_{i(F)}(x)$	$\text{asin}(x)$	(*)
$\arcsin_{c(F)}(x)$	$\text{casint}(x)$ or $\text{asin}(x)$	*
$\arccos_{F \rightarrow c(F)}(x)$	$\text{acosct}(x)$	†
$\arccos_{i(F) \rightarrow c(F)}(x)$	$\text{acos}(x)$	(*)
$\arccos_{c(F)}(x)$	$\text{cacost}(x)$ or $\text{acos}(x)$	*
$\arctan_{i(F)}(x)$	$\text{atan}(x)$	(*)
$\arctan_{i(F) \rightarrow c(F)}(x)$	$\text{atanht}(x)$	†
$\arctan_{c(F)}(x)$	$\text{catant}(x)$ or $\text{atan}(x)$	*
$\text{arccot}_{i(F)}(x)$	$\text{acot}(x)$	†
$\text{arccot}_{i(F) \rightarrow c(F)}(x)$	$\text{acot}(x)$	†
$\text{arccot}_{c(F)}(x)$	$\text{cacott}(x)$ or $\text{acot}(x)$	†
$\text{arcsec}_{F \rightarrow c(F)}(x)$	$\text{asecct}(x)$	†
$\text{arcsec}_{i(F) \rightarrow c(F)}(x)$	$\text{asecct}(x)$	†
$\text{arcsec}_{c(F)}(x)$	$\text{csect}(x)$ or $\text{asec}(x)$	†
$\text{arccsc}_{F \rightarrow c(F)}(x)$	$\text{acscct}(x)$	†
$\text{arccsc}_{i(F)}(x)$	$\text{acsc}(x)$	†
$\text{arccsc}_{c(F)}(x)$	$\text{cacsct}(x)$ or $\text{acsc}(x)$	†
$\text{radh}_F(x)$	$\text{radianht}(x)$	†
$\text{radh}_{i(F)}(x)$	$\text{radianht}(x)$	†
$\text{radh}_{c(F)}(x)$	$\text{radianht}(x)$	†
$\sinh_{i(F)}(x)$	$\sinh(x)$	(*)
$\sinh_{c(F)}(x)$	$\text{csinht}(x)$ or $\sinh(x)$	*
$\cosh_{i(F)}(x)$	$\cosh(x)$	(*)

$\cosh_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\cosh(x)$	★
$\tanh_{i(F)}(x)$	$\tanh(x)$	(★)
$\tanh_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\tanh(x)$	★
$\coth_{i(F)}(x)$	$\coth(x)$	†
$\coth_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\coth(x)$	†
$\operatorname{sech}_{i(F)}(x)$	$\operatorname{sech}(x)$	†
$\operatorname{sech}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{sech}(x)$	†
$\operatorname{csch}_{i(F)}(x)$	$\operatorname{csch}(x)$	†
$\operatorname{csch}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{csch}(x)$	†
$\operatorname{arcsinh}_{i(F)}(x)$	$\operatorname{asinh}(x)$	(★)
$\operatorname{arcsinh}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{asinh}_{ct}(x)$	†
$\operatorname{arcsinh}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{asinh}(x)$	★
$\operatorname{arccosh}_{F \rightarrow c(F)}(x)$	$\operatorname{acosh}_{ct}(x)$	†
$\operatorname{arccosh}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{acosh}(x)$	(★)
$\operatorname{arccosh}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{acosh}(x)$	★
$\operatorname{arctanh}_{F \rightarrow c(F)}(x)$	$\operatorname{atanh}_{ct}(x)$	†
$\operatorname{arctanh}_{i(F)}(x)$	$\operatorname{atanh}(x)$	(★)
$\operatorname{arctanh}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{atanh}(x)$	★
$\operatorname{arcoth}_{F \rightarrow c(F)}(x)$	$\operatorname{acoth}_{ct}(x)$	†
$\operatorname{arcoth}_{i(F)}(x)$	$\operatorname{acoth}(x)$	†
$\operatorname{arcoth}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{acoth}(x)$	†
$\operatorname{arcsech}_{F \rightarrow c(F)}(x)$	$\operatorname{asech}_{ct}(x)$	†
$\operatorname{arcsech}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{asech}(x)$	†
$\operatorname{arcsech}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{asech}(x)$	†
$\operatorname{arccsch}_{i(F)}(x)$	$\operatorname{acsch}(x)$	†
$\operatorname{arccsch}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{acsch}_{ct}(x)$	†
$\operatorname{arccsch}_{c(F)}(x)$	$\cosh_{ct}(x)$ or $\operatorname{acsch}(x)$	†

where x and y are expressions of type *FLT*, *IFLT*, or *CFLT* as appropriate. t is a string, part of the operation name, and is “f” for *float*, the empty string for *double*, and “l” for *long double*.

Arithmetic value conversions in C can be explicit or implicit. The explicit arithmetic value conversions are usually expressed as ‘casts’, except when converting to/from strings. The rules for when implicit conversions are applied is not repeated here, but work as if a cast had been applied.

$\operatorname{convert}_{I \rightarrow c(I)}(x)$	$(INT \text{ _Complex})x$	†
$\operatorname{convert}_{i(I) \rightarrow c(I)}(x)$	$(INT \text{ _Complex})x$	†
$\operatorname{convert}_{i(I) \rightarrow i(I')}(x)$	$(INT2 \text{ _Imaginary})x$	†
$\operatorname{convert}_{c(I) \rightarrow c(I')}(x)$	$(INT2 \text{ _Complex})x$	†
$\operatorname{convert}_{F \rightarrow c(F)}(y)$	$(FLT \text{ _Complex})x$	★
$\operatorname{convert}_{i(F) \rightarrow c(F)}(y)$	$(FLT \text{ _Complex})x$	(★)
$\operatorname{convert}_{i(F) \rightarrow i(F')}(y)$	$(FLT2 \text{ _Imaginary})x$	(★)
$\operatorname{convert}_{c(F) \rightarrow c(F')}(y)$	$(FLT2 \text{ _Complex})x$	★

where x is an expression of type *INT*, y is an expression of type *FLT*. *INT2* is the integer datatype that corresponds to I' . *FLT2* is the floating point datatype that corresponds to F' .

When converting to/from string formats, format strings are used. The format string is used as a pattern for the string format generated or parsed. Please see the C standard for a full

description. There are no special format indicators for imaginary or complex values. Instead, the real part and the imaginary part are formatted separately, and any syntax before/between/after those numerals must be specified in the format string.

C does not have any special syntax for complex, or imaginary, numerals in general. Instead imaginary and complex numerals are expressed as (compile-time evaluable) expressions involving a notation for the imaginary unit (of types `const int _Imaginary` and `const float _Imaginary` or `const int _Complex` and `const float _Complex`).

$imaginary_unit_{i(I)}$	II or <code>_Imaginary_II</code>	†
$imaginary_unit_{c(I)}$	II or <code>_Complex_II</code>	†
$imaginary_unit_{i(F)}$	I or <code>_Imaginary_I</code>	(★)
$imaginary_unit_{c(F)}$	I or <code>_Complex_I</code>	★

The I macro is defined as `_Imaginary_I` if imaginary floating point datatypes are provided by the implementations, but defined as `_Complex_I` if imaginary floating point datatypes are not provided. Similarly can be done for the II macro.

C.3 C++

The programming language C++ is defined by ISO/IEC 14882:2003, *Programming languages – C++* [5].

An implementation should follow all the requirements of LIA-3 unless otherwise specified by this language binding.

The operations or parameters marked “†” are not part of the language and should be provided by an implementation that wishes to conform to the LIA-3 for that operation. For each of the marked items a suggested identifier is provided.

This example binding recommends that all identifiers suggested here be defined in the namespace `std::math`.

The LIA datatype **Boolean** is implemented by the C++ datatype `bool`.

Every implementation of C++ has integral datatypes `int`, `long int`, `unsigned int`, and `unsigned long int`. *INT* is used below to designate one of the integer datatypes.

C++ has no standard imaginary integer or complex integer datatypes. However, in the suggested binding below, *IINT* and *CINT* are used to designate such potential datatypes.

C++ names two floating point datatypes: `float` and `double`. *FLT* is used below to designate one of the floating point datatypes.

C++ has three complex floating point datatypes: `complex<float>`, `complex<double>`, and `complex<long double>`. *CFLT* is used below to designate one of the complex floating point datatypes.

C++ does not have imaginary floating point, nor any imaginary or complex integer datatype.

The LIA-3 integer operations are listed below, along with the syntax used to invoke them:

$eq_{i(I)}(x, y)$	<code>x == y</code>	†
$eq_{I,i(I)}(x, y)$	<code>x == y</code>	†
$eq_{i(I),I}(x, y)$	<code>x == y</code>	†
$eq_{I,c(I)}(x, y)$	<code>x == y</code>	†
$eq_{c(I),I}(x, y)$	<code>x == y</code>	†

$eq_{i(I),c(I)}(x, y)$	$x == y$	†
$eq_{c(I),i(I)}(x, y)$	$x == y$	†
$eq_{c(I)}(x, y)$	$x == y$	†
$neq_{i(I)}(x, y)$	$x != y$	†
$neq_{I,i(I)}(x, y)$	$x != y$	†
$neq_{i(I),I}(x, y)$	$x != y$	†
$neq_{I,c(I)}(x, y)$	$x != y$	†
$neq_{c(I),I}(x, y)$	$x != y$	†
$neq_{i(I),c(I)}(x, y)$	$x != y$	†
$neq_{c(I),i(I)}(x, y)$	$x != y$	†
$neq_{c(I)}(x, y)$	$x != y$	†
$lss_{i(I)}(x, y)$	$x < y$	†
$leq_{i(I)}(x, y)$	$x <= y$	†
$gtr_{i(I)}(x, y)$	$x > y$	†
$geq_{i(I)}(x, y)$	$x >= y$	†
$itimes_{I \rightarrow i(I)}(x)$	$II * x$	†
$itimes_{i(I) \rightarrow I}(x)$	$II * x$	†
$itimes_{c(I)}(x)$	$II * x$	†
$re_I(x)$	$x.real()$ or $real(x)$	†
$re_{i(I)}(x)$	$x.real()$ or $real(x)$	†
$re_{c(I)}(x)$	$x.real()$ or $real(x)$	†
$im_I(x)$	$x.imag()$ or $imag(x)$	†
$im_{i(I)}(x)$	$x.imag()$ or $imag(x)$	†
$im_{c(I)}(x)$	$x.imag()$ or $imag(x)$	†
$plusitimes_I(x, y)$	$complex(x, y)$	†
$neg_{i(I)}(x)$	$-x$	†
$neg_{c(I)}(x)$	$-x$	†
$conj_I(x)$	$conj(x)$	†
$conj_{i(I)}(x)$	$conj(x)$	†
$conj_{c(I)}(x)$	$conj(x)$	†
$add_{i(I)}(x, y)$	$x + y$	†
$add_{I,i(I)}(x, y)$	$x + y$	†
$add_{i(I),I}(x, y)$	$x + y$	†
$add_{I,c(I)}(x, y)$	$x + y$	†
$add_{c(I),I}(x, y)$	$x + y$	†
$add_{i(I),c(I)}(x, y)$	$x + y$	†
$add_{c(I),i(I)}(x, y)$	$x + y$	†
$add_{c(I)}(x, y)$	$x + y$	†
$sub_{i(I)}(x, y)$	$x - y$	†
$sub_{I,i(I)}(x, y)$	$x - y$	†
$sub_{i(I),I}(x, y)$	$x - y$	†

$sub_{I,c(I)}(x, y)$	$x - y$	†
$sub_{c(I),I}(x, y)$	$x - y$	†
$sub_{i(I),c(I)}(x, y)$	$x - y$	†
$sub_{c(I),i(I)}(x, y)$	$x - y$	†
$sub_{c(I)}(x, y)$	$x - y$	†
$mul_{i(I)}(x, y)$	$x * y$	†
$mul_{I,i(I)}(x, y)$	$x * y$	†
$mul_{i(I),I}(x, y)$	$x * y$	†
$mul_{I,c(I)}(x, y)$	$x * y$	†
$mul_{c(I),I}(x, y)$	$x * y$	†
$mul_{i(I),c(I)}(x, y)$	$x * y$	†
$mul_{c(I),i(I)}(x, y)$	$x * y$	†
$mul_{c(I)}(x, y)$	$x * y$	†
$abs_{i(I)}(x)$	$abs(x)$	†
$signum_I(x)$	$signum(x)$	†
$signum_{i(I)}(x)$	$signum(x)$	†
$divides_{i(I)}(x, y)$	$divides(x, y)$	†
$divides_{I,i(I)}(x, y)$	$divides(x, y)$	†
$divides_{i(I),I}(x, y)$	$divides(x, y)$	†
$divides_{I,c(I)}(x, y)$	$divides(x, y)$	†
$divides_{c(I),I}(x, y)$	$divides(x, y)$	†
$divides_{i(I),c(I)}(x, y)$	$divides(x, y)$	†
$divides_{c(I),i(I)}(x, y)$	$divides(x, y)$	†
$divides_{c(I)}(x, y)$	$divides(x, y)$	†
$quot_{i(I)}(x, y)$	$quot(x, y)$	†
$quot_{I,i(I)}(x, y)$	$quot(x, y)$	†
$quot_{i(I),I}(x, y)$	$quot(x, y)$	†
$quot_{I,c(I)}(x, y)$	$quot(x, y)$	†
$quot_{c(I),I}(x, y)$	$quot(x, y)$	†
$quot_{i(I),c(I)}(x, y)$	$quot(x, y)$	†
$quot_{c(I),i(I)}(x, y)$	$quot(x, y)$	†
$quot_{c(I)}(x, y)$	$quot(x, y)$	†
$mod_{i(I)}(x, y)$	$mod(x, y)$	†
$mod_{I,i(I)}(x, y)$	$mod(x, y)$	†
$mod_{i(I),I}(x, y)$	$mod(x, y)$	†
$mod_{I,c(I)}(x, y)$	$mod(x, y)$	†
$mod_{c(I),I}(x, y)$	$mod(x, y)$	†
$mod_{i(I),c(I)}(x, y)$	$mod(x, y)$	†
$mod_{c(I),i(I)}(x, y)$	$mod(x, y)$	†
$mod_{c(I)}(x, y)$	$mod(x, y)$	†
$ratio_{i(I)}(x, y)$	$ratio(x, y)$	†

$ratio_{I,i(I)}(x, y)$	$ratio(x, y)$	†
$ratio_{i(I),I}(x, y)$	$ratio(x, y)$	†
$ratio_{I,c(I)}(x, y)$	$ratio(x, y)$	†
$ratio_{c(I),I}(x, y)$	$ratio(x, y)$	†
$ratio_{i(I),c(I)}(x, y)$	$ratio(x, y)$	†
$ratio_{c(I),i(I)}(x, y)$	$ratio(x, y)$	†
$ratio_{c(I)}(x, y)$	$ratio(x, y)$	†
$residue_{i(I)}(x, y)$	$iremainder(x, y)$	†
$residue_{I,i(I)}(x, y)$	$iremainder(x, y)$	†
$residue_{i(I),I}(x, y)$	$iremainder(x, y)$	†
$residue_{I,c(I)}(x, y)$	$iremainder(x, y)$	†
$residue_{c(I),I}(x, y)$	$iremainder(x, y)$	†
$residue_{i(I),c(I)}(x, y)$	$iremainder(x, y)$	†
$residue_{c(I),i(I)}(x, y)$	$iremainder(x, y)$	†
$residue_{c(I)}(x, y)$	$iremainder(x, y)$	†
$group_{i(I)}(x, y)$	$group(x, y)$	†
$group_{I,i(I)}(x, y)$	$group(x, y)$	†
$group_{i(I),I}(x, y)$	$group(x, y)$	†
$group_{I,c(I)}(x, y)$	$group(x, y)$	†
$group_{c(I),I}(x, y)$	$group(x, y)$	†
$group_{i(I),c(I)}(x, y)$	$group(x, y)$	†
$group_{c(I),i(I)}(x, y)$	$group(x, y)$	†
$group_{c(I)}(x, y)$	$group(x, y)$	†
$pad_{i(I)}(x, y)$	$pad(x, y)$	†
$pad_{I,i(I)}(x, y)$	$pad(x, y)$	†
$pad_{i(I),I}(x, y)$	$pad(x, y)$	†
$pad_{I,c(I)}(x, y)$	$pad(x, y)$	†
$pad_{c(I),I}(x, y)$	$pad(x, y)$	†
$pad_{i(I),c(I)}(x, y)$	$pad(x, y)$	†
$pad_{c(I),i(I)}(x, y)$	$pad(x, y)$	†
$pad_{c(I)}(x, y)$	$pad(x, y)$	†
$max_{i(I)}(x, y)$	$max(x, y)$	†
$min_{i(I)}(x, y)$	$min(x, y)$	†
$max_seq_{i(I)}(xs)$	$xs.max()$	†
$min_seq_{i(I)}(xs)$	$xs.min()$	†

where x and y are expressions of imaginary or complex integer type, or an integer datatype as appropriate. xs is an expression having the type of an array of an imaginary integer datatype.

The parameters for operations approximating complex valued transcendental functions can be accessed by the following syntax:

$box_error_mode_mul_{c(F)}$	$numeric_limits<CFLT>::box_err_cmul()$	†
$max_err_mul_{c(F)}$	$numeric_limits<CFLT>::err_cmul()$	†
$box_error_mode_div_{c(F)}$	$numeric_limits<CFLT>::box_err_cdiv()$	†

$max_err_div_{c(F)}$	<code>numeric_limits<CFLT>::err_cdiv()</code>	†
$max_err_exp_{c(F)}$	<code>numeric_limits<CFLT>::err_exp()</code>	†
$max_err_power_{c(F)}$	<code>numeric_limits<CFLT>::err_power()</code>	†
$max_err_sin_{c(F)}$	<code>numeric_limits<CFLT>::err_sin()</code>	†
$max_err_tan_{c(F)}$	<code>numeric_limits<CFLT>::err_tan()</code>	†

The LIA-3 non-transcendental complex floating point operations are listed below, along with the syntax used to invoke them:

$eq_{i(F)}(x, y)$	<code>x == y</code>	†
$eq_{F,i(F)}(x, y)$	<code>x == y</code>	†
$eq_{i(F),F}(x, y)$	<code>x == y</code>	†
$eq_{F,c(F)}(x, y)$	<code>x == y</code>	*
$eq_{c(F),F}(x, y)$	<code>x == y</code>	*
$eq_{i(F),c(F)}(x, y)$	<code>x == y</code>	†
$eq_{c(F),i(F)}(x, y)$	<code>x == y</code>	†
$eq_{c(F)}(x, y)$	<code>x == y</code>	*
$neq_{i(F)}(x, y)$	<code>x != y</code>	†
$neq_{F,i(F)}(x, y)$	<code>x != y</code>	†
$neq_{i(F),F}(x, y)$	<code>x != y</code>	†
$neq_{F,c(F)}(x, y)$	<code>x != y</code>	*
$neq_{c(F),F}(x, y)$	<code>x != y</code>	*
$neq_{i(F),c(F)}(x, y)$	<code>x != y</code>	†
$neq_{c(F),i(F)}(x, y)$	<code>x != y</code>	†
$neq_{c(F)}(x, y)$	<code>x != y</code>	*
$lss_{i(F)}(x, y)$	<code>x < y</code>	†
$leq_{i(F)}(x, y)$	<code>x <= y</code>	†
$gtr_{i(F)}(x, y)$	<code>x > y</code>	†
$geq_{i(F)}(x, y)$	<code>x >= y</code>	†
$itimes_{F \rightarrow i(F)}(x)$	<code>I * x</code>	†
$itimes_{i(F) \rightarrow F}(x)$	<code>I * x</code>	†
$itimes_{c(F)}(x)$	<code>I * x</code>	†
$re_F(x)$	<code>x.real()</code> or <code>real(x)</code>	†
$re_{i(F)}(x)$	<code>x.real()</code> or <code>real(x)</code>	†
$re_{c(F)}(x)$	<code>x.real()</code> or <code>real(x)</code>	*
$im_F(x)$	<code>x.imag()</code> or <code>imag(x)</code>	†
$im_{i(F)}(x)$	<code>x.imag()</code> or <code>imag(x)</code>	†
$im_{c(F)}(x)$	<code>x.imag()</code> or <code>imag(x)</code>	*
$plusitimes_F(x, y)$	<code>complex(x, y)</code>	*
$plusitimes_F(x, y)$	<code>x + I * y</code>	†
$neg_{i(F)}(x)$	<code>- x</code>	†
$neg_{c(F)}(x)$	<code>- x</code>	*
$conj_F(x)$	<code>conj(x)</code>	†

$conj_{i(F)}(x)$	$conj(x)$	†
$conj_{c(F)}(x)$	$conj(x)$	*
$add_{i(F)}(x, y)$	$x + y$	†
$add_{F,i(F)}(x, y)$	$x + y$	†
$add_{i(F),F}(x, y)$	$x + y$	†
$add_{F,c(F)}(x, y)$	$x + y$	*
$add_{c(F),F}(x, y)$	$x + y$	*
$add_{i(F),c(F)}(x, y)$	$x + y$	†
$add_{c(F),i(F)}(x, y)$	$x + y$	†
$add_{c(F)}(x, y)$	$x + y$	*
$sub_{i(F)}(x, y)$	$x - y$	†
$sub_{F,i(F)}(x, y)$	$x - y$	†
$sub_{i(F),F}(x, y)$	$x - y$	†
$sub_{F,c(F)}(x, y)$	$x - y$	*
$sub_{c(F),F}(x, y)$	$x - y$	*
$sub_{i(F),c(F)}(x, y)$	$x - y$	†
$sub_{c(F),i(F)}(x, y)$	$x - y$	†
$sub_{c(F)}(x, y)$	$x - y$	*
$mul_{i(F)}(x, y)$	$x * y$	†
$mul_{F,i(F)}(x, y)$	$x * y$	†
$mul_{i(F),F}(x, y)$	$x * y$	†
$mul_{F,c(F)}(x, y)$	$x * y$	*
$mul_{c(F),F}(x, y)$	$x * y$	*
$mul_{i(F),c(F)}(x, y)$	$x * y$	†
$mul_{c(F),i(F)}(x, y)$	$x * y$	†
$mul_{c(F)}(x, y)$	$x * y$	*
$div_{i(F)}(x, y)$	x / y	†
$div_{F,i(F)}(x, y)$	x / y	†
$div_{i(F),F}(x, y)$	x / y	†
$div_{F,c(F)}(x, y)$	x / y	*
$div_{c(F),F}(x, y)$	x / y	*
$div_{i(F),c(F)}(x, y)$	x / y	†
$div_{c(F),i(F)}(x, y)$	x / y	†
$div_{c(F)}(x, y)$	x / y	*
$abs_{i(F)}(x)$	$abs(x)$	*
$abs_{c(F)}(x)$	$abs(x)$	*
$sqr_abs_{c(F)}(x)$	$norm(x)$	* Not LIA-3
$phase_F(x)$	$arg(x)$	†
$phase_{i(F)}(x)$	$arg(x)$	†
$phase_{c(F)}(x)$	$arg(x)$	*
$phaseu_F(u, x)$	$argu(x, u)$	†

$phase_{i(F)}(u, x)$	<code>argu(x, u)</code>	†
$phase_{c(F)}(u, x)$	<code>argu(x, u)</code>	†
$signum_F(x)$	<code>sign(x)</code>	†
$signum_{i(F)}(x)$	<code>sign(x)</code>	†
$signum_{c(F)}(x)$	<code>sign(x)</code>	†
$floor_{i(F)}(x)$	<code>floor(x)</code>	†
$floor_{c(F)}(x)$	<code>floor(x)</code>	†
$rounding_{i(F)}(x)$	<code>nearbyint(x)</code>	†
$rounding_{c(F)}(x)$	<code>nearbyint(x)</code>	†
$ceiling_{i(F)}(x)$	<code>ceiling(x)</code>	†
$ceiling_{c(F)}(x)$	<code>ceiling(x)</code>	†
$max_{i(F)}(x, y)$	<code>nmax(x, y)</code>	†
$mmax_{i(F)}(x, y)$	<code>fmax(x, y)</code>	†
$min_{i(F)}(x, y)$	<code>nmin(x, y)</code>	†
$mmin_{i(F)}(x, y)$	<code>fmin(x, y)</code>	†
$max_seq_{i(F)}(xs)$	<code>nmax(xs, nr_of_items)</code>	†
$mmax_seq_{i(F)}(xs)$	<code>fmax(xs, nr_of_items)</code>	†
$min_seq_{i(F)}(xs)$	<code>nmin(xs, nr_of_items)</code>	†
$mmin_seq_{i(F)}(xs)$	<code>fmin(xs, nr_of_items)</code>	†
$polar_F(x, y)$	<code>polar(x, y)</code>	★
$polaru_F(u, x, y)$	<code>polaru(x, y, u)</code>	†

where x and y are expressions of an imaginary or complex floating point type, and u is of the corresponding real floating point type.

The LIA-3 elementary complex floating point operations are listed below, along with the syntax used to invoke them:

$power_{i(F),I}(x, i)$	<code>pow(x, i)</code>	†
$power_{c(F),I}(x, i)$	<code>pow(x, i)</code>	★ (for int) Not LIA-3
$exp_{i(F)}(x)$	<code>exp(x)</code>	†
$exp_{c(F)}(x)$	<code>exp(x)</code>	★
$power_{F \rightarrow c(F)}(x, y)$	<code>powc(x, y)</code>	†
$power_{c(F),F}(x, y)$	<code>pow(x, y)</code>	★ Not LIA-3
$power_{F,c(F)}(x, y)$	<code>pow(x, y)</code>	★ Not LIA-3
$power_{c(F)}(x, y)$	<code>pow(x, y)</code>	★ Not LIA-3
$sqrt_{F \rightarrow c(F)}(x)$	<code>sqrtc(x)</code>	†
$sqrt_{i(F) \rightarrow c(F)}(x)$	<code>sqrt(x)</code>	†
$sqrt_{c(F)}(x)$	<code>sqrt(x)</code>	★
$ln_{F \rightarrow c(F)}(x)$	<code>logc(x)</code>	†
$ln_{i(F) \rightarrow c(F)}(x)$	<code>complex(log(abs(x)), atan2(imag(x), real(x)))</code>	★

$ln_{i(F) \rightarrow c(F)}(x)$	$\log(x)$	†
$ln_{c(F)}(x)$	$\log(x)$	*
$rad_{i(F)}(x)$	$\text{radian}(x)$	†
$rad_{c(F)}(x)$	$\text{radian}(x)$	†
$sin_{i(F)}(x)$	$\sin(x)$	†
$sin_{c(F)}(x)$	$\sin(x)$	*
$cos_{i(F)}(x)$	$\cos(x)$	†
$cos_{c(F)}(x)$	$\cos(x)$	*
$tan_{i(F)}(x)$	$\tan(x)$	†
$tan_{c(F)}(x)$	$\tan(x)$	*
$cot_{i(F)}(x)$	$\cot(x)$	†
$cot_{c(F)}(x)$	$\cot(x)$	†
$sec_{i(F)}(x)$	$\sec(x)$	†
$sec_{c(F)}(x)$	$\sec(x)$	†
$csc_{i(F)}(x)$	$\csc(x)$	†
$csc_{c(F)}(x)$	$\csc(x)$	†
$arcsin_{F \rightarrow c(F)}(x)$	$\text{asinc}(x)$	†
$arcsin_{i(F)}(x)$	$\text{asin}(x)$	†
$arcsin_{c(F)}(x)$	$\text{asin}(x)$	*
$arccos_{F \rightarrow c(F)}(x)$	$\text{acosc}(x)$	†
$arccos_{c(F) \rightarrow c(F)}(x)$	$\text{acos}(x)$	†
$arccos_{c(F)}(x)$	$\text{acos}(x)$	*
$arctan_{i(F)}(x)$	$\text{atan}(x)$	†
$arctan_{i(F) \rightarrow c(F)}(x)$	$\text{atanc}(x)$	†
$arctan_{c(F)}(x)$	$\text{atan}(x)$	*
$arccot_{i(F)}(x)$	$\text{acot}(x)$	†
$arccot_{i(F) \rightarrow c(F)}(x)$	$\text{acotc}(x)$	†
$arccot_{c(F)}(x)$	$\text{acot}(x)$	†
$arcsec_{F \rightarrow c(F)}(x)$	$\text{asecc}(x)$	†
$arcsec_{i(F) \rightarrow c(F)}(x)$	$\text{asec}(x)$	†
$arcsec_{c(F)}(x)$	$\text{asec}(x)$	†
$arccsc_{F \rightarrow c(F)}(x)$	$\text{acsc}(x)$	†
$arccsc_{i(F)}(x)$	$\text{acsc}(x)$	†
$arccsc_{c(F)}(x)$	$\text{acsc}(x)$	†
$radh_F(x)$	$\text{radianh}(x)$	†
$radh_{i(F)}(x)$	$\text{radianh}(x)$	†
$radh_{c(F)}(x)$	$\text{radianh}(x)$	†
$sinh_{i(F)}(x)$	$\sinh(x)$	†
$sinh_{c(F)}(x)$	$\sinh(x)$	*
$cosh_{i(F)}(x)$	$\cosh(x)$	†
$cosh_{c(F)}(x)$	$\cosh(x)$	*
$tanh_{i(F)}(x)$	$\tanh(x)$	†

$\tanh_{c(F)}(x)$	$\tanh(x)$	★
$\coth_{i(F)}(x)$	$\coth(x)$	†
$\coth_{c(F)}(x)$	$\coth(x)$	†
$\operatorname{sech}_{i(F)}(x)$	$\operatorname{sech}(x)$	†
$\operatorname{sech}_{c(F)}(x)$	$\operatorname{sech}(x)$	†
$\operatorname{csch}_{i(F)}(x)$	$\operatorname{csch}(x)$	†
$\operatorname{csch}_{c(F)}(x)$	$\operatorname{csch}(x)$	†
$\operatorname{arcsinh}_{i(F)}(x)$	$\operatorname{asinh}(x)$	†
$\operatorname{arcsinh}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{asinhc}(x)$	†
$\operatorname{arcsinh}_{c(F)}(x)$	$\operatorname{asinh}(x)$	★
$\operatorname{arccosh}_{F \rightarrow c(F)}(x)$	$\operatorname{acoshc}(x)$	†
$\operatorname{arccosh}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{acoshc}(x)$	†
$\operatorname{arccosh}_{c(F)}(x)$	$\operatorname{acosh}(x)$	★
$\operatorname{artanh}_{F \rightarrow c(F)}(x)$	$\operatorname{atanhc}(x)$	†
$\operatorname{artanh}_{i(F)}(x)$	$\operatorname{atanh}(x)$	†
$\operatorname{artanh}_{c(F)}(x)$	$\operatorname{atanh}(x)$	★
$\operatorname{arcoth}_{F \rightarrow c(F)}(x)$	$\operatorname{acothc}(x)$	†
$\operatorname{arcoth}_{i(F)}(x)$	$\operatorname{acoth}(x)$	†
$\operatorname{arcoth}_{c(F)}(x)$	$\operatorname{acoth}(x)$	†
$\operatorname{arcsech}_{F \rightarrow c(F)}(x)$	$\operatorname{asechc}(x)$	†
$\operatorname{arcsech}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{asech}(x)$	†
$\operatorname{arcsech}_{c(F)}(x)$	$\operatorname{asech}(x)$	†
$\operatorname{arccsch}_{i(F)}(x)$	$\operatorname{acsch}(x)$	†
$\operatorname{arccsch}_{i(F) \rightarrow c(F)}(x)$	$\operatorname{acschc}(x)$	†
$\operatorname{arccsch}_{c(F)}(x)$	$\operatorname{acsch}(x)$	†

where x , and y are expressions of an imaginary, complex, or plain floating point type (as appropriate).

When converting to/from string formats when using the C printf/scanf family of routines, format strings are used. The format string is used as a pattern for the string format generated or parsed. Please see the C99 standard for a full description. There are no special format indicators for imaginary or complex values. Instead, the real part and the imaginary part are formatted separately, and any syntax before/between/after those numerals must be specified in the format string. In addition, C++ has a predefined format for complex (but not imaginary) values for I/O when using the stream formatting operators “<<” and “>>”. For details, see clause 26.2.5 of the C++ standard [5].

$\operatorname{convert}_{c(F) \rightarrow c(F')}(x)$	$is \gg y$	★
$\operatorname{convert}_{c(F') \rightarrow c(F)}(x)$	$os \ll x$	★

C++ does not have any special syntax for complex numerals in general. Instead the `complex` function is used on floating point values (not necessarily syntactically numerals).

C.4 Fortran

The programming language Fortran is defined by ISO/IEC 1539-1:1997, *Information technology – Programming languages – Fortran – Part 1: Base language* [6].

An implementation should follow all the requirements of LIA-3 unless otherwise specified by this language binding.

The operations or parameters marked “†” are not part of the language and should be provided by an implementation that wishes to conform to the LIA-3 for that operation. For each of the marked items a suggested identifier is provided.

The Fortran datatype **LOGICAL** corresponds to the LIA datatype **Boolean**.

Every implementation of Fortran has at least one integer datatype, denoted as **INTEGER**, and at least two floating point datatypes denoted as **REAL** (single precision, also denoted **REAL(KIND(0.0))**) and **DOUBLE PRECISION** (also denoted **REAL(KIND(0d0))**, and at least two complex floating point datatypes denoted as **COMPLEX** (also denoted **COMPLEX(KIND(0.0))**, and **COMPLEX(KIND(0d0))**. Standard Fortran does not have any imaginary datatypes, nor any complex integer datatypes.

An implementation is permitted to offer additional **INTEGER** types with a different range and additional **REAL** and **COMPLEX** types with different precision or range, parameterised with the **KIND** parameter.

The LIA-3 integer operations are listed below, along with the syntax used to invoke them:

$eq_{i(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{I,i(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{i(I),I}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{I,c(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{c(I),I}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{i(I),c(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{c(I),i(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{c(I)}(x, y)$	$x == y$ or $x .EQ. y$	†
$neq_{i(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{I,i(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{i(I),I}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{I,c(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{c(I),I}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{i(I),c(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{c(I),i(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{c(I)}(x, y)$	$x /= y$ or $x .NE. y$	†
$lss_{i(I)}(x, y)$	$x < y$ or $x .LT. y$	†
$leq_{i(I)}(x, y)$	$x <= y$ or $x .LE. y$	†
$gtr_{i(I)}(x, y)$	$x > y$ or $x .GT. y$	†
$geq_{i(I)}(x, y)$	$x >= y$ or $x .GE. y$	†
$itimes_{I \rightarrow i(I)}(x)$	IMUL (x)	†
$itimes_{i(I) \rightarrow I}(x)$	IMUL (x)	†
$itimes_{c(I)}(x)$	IMUL (x)	†
$re_I(x)$	IREAL (x)	†
$re_{i(I)}(x)$	IREAL (x)	†
$re_{c(I)}(x)$	IREAL (x)	†
$im_I(x)$	IMAG (x)	†
$im_{i(I)}(x)$	IMAG (x)	†

$im_{c(I)}(x)$	IMAG(x)	†
$plusitimes_I(x, y)$	CMPLX(x, y)	†
$neg_{i(I)}(x)$	$-x$	†
$neg_{c(I)}(x)$	$-x$	†
$conj_I(x)$	CONJG(x)	†
$conj_{i(I)}(x)$	CONJG(x)	†
$conj_{c(I)}(x)$	CONJG(x)	†
$add_{i(I)}(x, y)$	$x + y$	†
$add_{I,i(I)}(x, y)$	$x + y$	†
$add_{i(I),I}(x, y)$	$x + y$	†
$add_{I,c(I)}(x, y)$	$x + y$	†
$add_{c(I),I}(x, y)$	$x + y$	†
$add_{i(I),c(I)}(x, y)$	$x + y$	†
$add_{c(I),i(I)}(x, y)$	$x + y$	†
$add_{c(I)}(x, y)$	$x + y$	†
$sub_{i(I)}(x, y)$	$x - y$	†
$sub_{I,i(I)}(x, y)$	$x - y$	†
$sub_{i(I),I}(x, y)$	$x - y$	†
$sub_{I,c(I)}(x, y)$	$x - y$	†
$sub_{c(I),I}(x, y)$	$x - y$	†
$sub_{i(I),c(I)}(x, y)$	$x - y$	†
$sub_{c(I),i(I)}(x, y)$	$x - y$	†
$sub_{c(I)}(x, y)$	$x - y$	†
$mul_{i(I)}(x, y)$	$x * y$	†
$mul_{I,i(I)}(x, y)$	$x * y$	†
$mul_{i(I),I}(x, y)$	$x * y$	†
$mul_{I,c(I)}(x, y)$	$x * y$	†
$mul_{c(I),I}(x, y)$	$x * y$	†
$mul_{i(I),c(I)}(x, y)$	$x * y$	†
$mul_{c(I),i(I)}(x, y)$	$x * y$	†
$mul_{c(I)}(x, y)$	$x * y$	†
$abs_{i(I)}(x)$	ABS(x)	†
$signum_I(x)$	SIGN(1, x)	★
$signum_{i(I)}(x)$	SIGN(1, x)	†
$mul_I(abs_I(y), signum_I(x))$	SIGN(y, x)	★
	(SIGN(y, x) is the copysign operation of IEC 60559)	
$divides_{i(I)}(x, y)$	DIVIDES(x, y)	†
$divides_{I,i(I)}(x, y)$	DIVIDES(x, y)	†
$divides_{i(I),I}(x, y)$	DIVIDES(x, y)	†
$divides_{I,c(I)}(x, y)$	DIVIDES(x, y)	†
$divides_{c(I),I}(x, y)$	DIVIDES(x, y)	†
$divides_{i(I),c(I)}(x, y)$	DIVIDES(x, y)	†

$divides_{c(I),i(I)}(x, y)$	DIVIDES(x, y)	†
$divides_{c(I)}(x, y)$	DIVIDES(x, y)	†
$quot_{i(I)}(x, y)$	QUOTIENT(x, y)	†
$quot_{I,i(I)}(x, y)$	QUOTIENT(x, y)	†
$quot_{i(I),I}(x, y)$	QUOTIENT(x, y)	†
$quot_{I,c(I)}(x, y)$	QUOTIENT(x, y)	†
$quot_{c(I),I}(x, y)$	QUOTIENT(x, y)	†
$quot_{i(I),c(I)}(x, y)$	QUOTIENT(x, y)	†
$quot_{c(I),i(I)}(x, y)$	QUOTIENT(x, y)	†
$quot_{c(I)}(x, y)$	QUOTIENT(x, y)	†
$mod_{i(I)}(x, y)$	MODULO(x, y)	†
$mod_{I,i(I)}(x, y)$	MODULO(x, y)	†
$mod_{i(I),I}(x, y)$	MODULO(x, y)	†
$mod_{I,c(I)}(x, y)$	MODULO(x, y)	†
$mod_{c(I),I}(x, y)$	MODULO(x, y)	†
$mod_{i(I),c(I)}(x, y)$	MODULO(x, y)	†
$mod_{c(I),i(I)}(x, y)$	MODULO(x, y)	†
$mod_{c(I)}(x, y)$	MODULO(x, y)	†
$ratio_{i(I)}(x, y)$	RATIO(x, y)	†
$ratio_{I,i(I)}(x, y)$	RATIO(x, y)	†
$ratio_{i(I),I}(x, y)$	RATIO(x, y)	†
$ratio_{I,c(I)}(x, y)$	RATIO(x, y)	†
$ratio_{c(I),I}(x, y)$	RATIO(x, y)	†
$ratio_{i(I),c(I)}(x, y)$	RATIO(x, y)	†
$ratio_{c(I),i(I)}(x, y)$	RATIO(x, y)	†
$ratio_{c(I)}(x, y)$	RATIO(x, y)	†
$residue_{i(I)}(x, y)$	RESIDUE(x, y)	†
$residue_{I,i(I)}(x, y)$	RESIDUE(x, y)	†
$residue_{i(I),I}(x, y)$	RESIDUE(x, y)	†
$residue_{I,c(I)}(x, y)$	RESIDUE(x, y)	†
$residue_{c(I),I}(x, y)$	RESIDUE(x, y)	†
$residue_{i(I),c(I)}(x, y)$	RESIDUE(x, y)	†
$residue_{c(I),i(I)}(x, y)$	RESIDUE(x, y)	†
$residue_{c(I)}(x, y)$	RESIDUE(x, y)	†
$group_{i(I)}(x, y)$	GROUP(x, y)	†
$group_{I,i(I)}(x, y)$	GROUP(x, y)	†
$group_{i(I),I}(x, y)$	GROUP(x, y)	†
$group_{I,c(I)}(x, y)$	GROUP(x, y)	†
$group_{c(I),I}(x, y)$	GROUP(x, y)	†
$group_{i(I),c(I)}(x, y)$	GROUP(x, y)	†
$group_{c(I),i(I)}(x, y)$	GROUP(x, y)	†
$group_{c(I)}(x, y)$	GROUP(x, y)	†

$pad_{i(I)}(x, y)$	PAD(x, y)	†
$pad_{I,i(I)}(x, y)$	PAD(x, y)	†
$pad_{i(I),I}(x, y)$	PAD(x, y)	†
$pad_{I,c(I)}(x, y)$	PAD(x, y)	†
$pad_{c(I),I}(x, y)$	PAD(x, y)	†
$pad_{i(I),c(I)}(x, y)$	PAD(x, y)	†
$pad_{c(I),i(I)}(x, y)$	PAD(x, y)	†
$pad_{c(I)}(x, y)$	PAD(x, y)	†
$max_{i(I)}(x, y)$	MAX(x, y)	†
$min_{i(I)}(x, y)$	MIN(x, y)	†
$max_seq_{i(I)}([x_1, \dots, x_n])$	MAX($x_1, \dots, x_n$)	†
$min_seq_{i(I)}([x_1, \dots, x_n])$	MIN($x_1, \dots, x_n$)	†
$max_seq_{i(I)}(xs)$	MAXVAL(xs)	†
$min_seq_{i(I)}(xs)$	MINVAL(xs)	†

where x and y are expressions of complex integer, imaginary integer, or plain integer type (as appropriate) and where xs is an expression of type array of imaginary integer.

The LIA-3 parameters for operations approximating real valued transcendental functions can be accessed by the following syntax:

$box_error_mode_mul_{c(F)}$	BOX_ERR_MUL(x)	†
$max_err_mul_{c(F)}$	ERR_MUL(x)	†
$box_error_mode_div_{c(F)}$	BOX_ERR_DIV(x)	†
$max_err_div_{c(F)}$	ERR_DIV(x)	†
$max_err_exp_{c(F)}$	ERR_EXP(x)	†
$max_err_power_{c(F)}$	ERR_POWER(x)	†
$max_err_sin_{c(F)}$	ERR_SIN(x)	†
$max_err_tan_{c(F)}$	ERR_TAN(x)	†

where x is an expression of type COMPLEX. Several of the parameter functions are constant for each type (and library), the argument is then used only to differentiate among the floating point types.

The LIA-3 non-transcendental floating point operations are listed below, along with the syntax used to invoke them:

$eq_{i(F)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{F,i(F)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{i(F),F}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{F,c(F)}(x, y)$	$x == y$ or $x .EQ. y$	★
$eq_{c(F),F}(x, y)$	$x == y$ or $x .EQ. y$	★
$eq_{i(F),c(F)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{c(F),i(F)}(x, y)$	$x == y$ or $x .EQ. y$	†
$eq_{c(F)}(x, y)$	$x == y$ or $x .EQ. y$	★
$neq_{i(F)}(x, y)$	$x /= y$ or $x .NE. y$	†
$neq_{F,i(F)}(x, y)$	$x /= y$ or $x .NE. y$	†

$neq_{i(F),F}(x, y)$	$x \neq y$ or $x .NE. y$	†
$neq_{F,c(F)}(x, y)$	$x \neq y$ or $x .NE. y$	*
$neq_{c(F),F}(x, y)$	$x \neq y$ or $x .NE. y$	*
$neq_{i(F),c(F)}(x, y)$	$x \neq y$ or $x .NE. y$	†
$neq_{c(F),i(F)}(x, y)$	$x \neq y$ or $x .NE. y$	†
$neq_{c(F)}(x, y)$	$x \neq y$ or $x .NE. y$	*
$lss_{i(F)}(x, y)$	$x < y$ or $x .LT. y$	†
$leq_{i(F)}(x, y)$	$x \leq y$ or $x .LE. y$	†
$gtr_{i(F)}(x, y)$	$x > y$ or $x .GT. y$	†
$geq_{i(F)}(x, y)$	$x \geq y$ or $x .GE. y$	†
$itimes_{F \rightarrow i(F)}(x)$	IMUL(x)	†
$itimes_{i(F) \rightarrow F}(x)$	IMUL(x)	†
$itimes_{c(F)}(x)$	IMUL(x)	†
$re_F(x)$	REAL(x)	*
$re_{i(F)}(x)$	REAL(x)	†
$re_{c(F)}(x)$	REAL(x)	*
$im_F(x)$	AIMAG(x)	†
$im_{i(F)}(x)$	AIMAG(x)	*
$im_{c(F)}(x)$	AIMAG(x)	*
$plusitimes_F(x, y)$	CMPLX(x, y)	*
$plusitimes_F(x, 0)$	CMPLX(x)	*
$neg_{i(F)}(x)$	$-x$	†
$neg_{c(F)}(x)$	$-x$	*
$conj_F(x)$	CONJG(x)	†
$conj_{i(F)}(x)$	CONJG(x)	†
$conj_{c(F)}(x)$	CONJG(x)	*
$add_{i(F)}(x, y)$	$x + y$	†
$add_{F,i(F)}(x, y)$	$x + y$	†
$add_{i(F),F}(x, y)$	$x + y$	†
$add_{F,c(F)}(x, y)$	$x + y$	*
$add_{c(F),F}(x, y)$	$x + y$	*
$add_{i(F),c(F)}(x, y)$	$x + y$	†
$add_{c(F),i(F)}(x, y)$	$x + y$	†
$add_{c(F)}(x, y)$	$x + y$	*
$sub_{i(F)}(x, y)$	$x - y$	†
$sub_{F,i(F)}(x, y)$	$x - y$	†
$sub_{i(F),F}(x, y)$	$x - y$	†
$sub_{F,c(F)}(x, y)$	$x - y$	*
$sub_{c(F),F}(x, y)$	$x - y$	*
$sub_{i(F),c(F)}(x, y)$	$x - y$	†
$sub_{c(F),i(F)}(x, y)$	$x - y$	†
$sub_{c(F)}(x, y)$	$x - y$	*