



On the robustness of the 2Sum and Fast2Sum algorithms

Sylvie Boldo, Stef Graillat, Jean-Michel Muller

► To cite this version:

Sylvie Boldo, Stef Graillat, Jean-Michel Muller. On the robustness of the 2Sum and Fast2Sum algorithms. 2016. ensl-01310023v1

HAL Id: ensl-01310023

<https://ens-lyon.hal.science/ensl-01310023v1>

Preprint submitted on 1 May 2016 (v1), last revised 22 Feb 2017 (v2)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

On the robustness of the 2Sum and Fast2Sum algorithms

Sylvie Boldo

Inria, LRI, CNRS & Univ. Paris-Sud, Université Paris-Saclay, France

Stef Graillat

Sorbonne Universités, UPMC Univ Paris 06, CNRS, LIP6 UMR 7606, Paris

Jean-Michel Muller

CNRS, LIP, Université de Lyon, France

May 1, 2016

Abstract

The 2Sum and Fast2Sum algorithms are important building blocks in numerical computing. They are used (implicitly or explicitly) in many *compensated* algorithms (such as compensated summation or compensated polynomial evaluation). They are also used for manipulating floating-point *expansions*. We show that these algorithms are much more robust than it is usually believed: the returned result makes sense even when the rounding function is not round-to-nearest, and they are almost immune to overflow.

Keywords floating-point, error-free transformation, rounding errors, faithful rounding, 2Sum, Fast2Sum

1 Motivations

One easily shows that, provided that no overflow occurred, the error of a rounded-to-nearest floating-point addition or subtraction is exactly representable by a floating-point number. The 2Sum [9] and Fast2Sum [2] algorithms make it possible to compute that error, under some conditions that will be reminded below. That error can therefore be used later on in a calculation: this is the underlying idea behind *compensated* algorithms. This has allowed for the development of various techniques, such as very accurate (compensated) summation [8, 14, 18, 17, 3], accurate polynomial evaluation [5], efficient manipulation of floating-point expansions [15, 19, 7], etc. However, these techniques suffer from some limitations:

- as noticed, among others, by Boldo and Daumas [1], when the rounding function differs from round-to-nearest, the error of floating-point addition/subtraction may not be exactly representable. And yet, rounding functions such as round towards $\pm\infty$ and round towards zero are very useful. They allow one to get certain lower and/or upper bounds on the exact result of a computation, and to easily implement *interval arithmetic* or *stochastic arithmetic*. With these rounding functions, if we cannot always obtain the “exact” error of floating-point addition, it would still be useful to obtain a value close to that error. This problem was partly dealt with by Demmel and Nguyen [3], and later on by Graillat, Jézéquel, and Picot [4] for the Fast2Sum algorithm, and by Martin-Dorel et al. [11] in the case of “double roundings”. We aim at tackling this issue in a more general context, and we wish to study the behaviour of 2Sum and Fast2Sum just assuming “general” rounding functions (see definition 2.1 below).
- in the literature, these algorithms are usually considered as returning a correct result provided that no underflow or overflow occurs. The case of underflow can be handled fairly intuitively, using a property mentioned by Hauser [6]—see below. The case of overflow is more problematic: the central question is: *can we have a “spurious” overflow?*, i.e., can we have situations where the initial addition does not overflow, and yet one of the arithmetic operations performed in the 2Sum or Fast2Sum algorithm overflows? We will see in the following that such a situation almost never arises.

There exist some kind of error-free transformations for summation with faithful rounding (see Priest [16]). But these algorithms are costly and are not necessary to achieve a good accuracy in many compensated algorithms [4].

The rest of the paper is organized as follows. In Section 2, we introduce some notation, definitions and preliminary remarks used in the sequel. The accuracy of Fast2Sum with no overflow is analyzed in Section 3 while the accuracy of 2Sum is dealt with in Section 4. Section 5 is devoted to show that Fast2Sum is immune to overflow and Section 6 to show that 2Sum is almost immune to overflow.

2 Notation, definitions, preliminary remarks

Throughout this paper, we assume a radix-2, precision- p , floating-point (FP) arithmetic, of extremal exponents $e_{\min}$ and $e_{\max}$. We also assume that subnormal numbers are available. We denote by Ω the largest representable floating-point number:

$$\Omega = (2 - 2^{1-p}) \cdot 2^{e_{\max}}.$$

The floating-point predecessor of a FP number x will be noted $\text{pred}(x)$. Its successor will be noted $\text{succ}(x)$. When an arithmetic operation τ is performed, with input FP operands a and b , what is actually computed is $\circ(a\tau b)$, where $\circ$

is a *rounding function*. The IEEE-754 Standard for Floating-point arithmetic defines 5 rounding functions (round towards $+\infty$ —denoted RU below—, round towards $-\infty$ —denoted RD below—, round towards zero, round to nearest ties to even, and round to nearest ties to infinity). The two round-to-nearest functions will be denoted RN in the following (the choice of the tie-breaking rule is not important here). We say that the FP number $\hat{x}$ is a *faithful rounding* of the real number x if $\hat{x} \in \{\text{RD}(x), \text{RU}(x)\}$. If x is a real number, satisfying $2^k \leq |x| < 2^{k+1}$, where k is an integer, we define $\text{ulp}(x)$ as follows:

$$\text{ulp}(x) = 2^{\max(k, e_{\min}) - p + 1}.$$

The rounding functions considered in this paper satisfy the following definition (introduced by Kulisch [10] under the name of *optimal rounding*).

Definition 2.1 (Rounding function—“optimal rounding” in [10]). *Function $\circ$ from $\mathbb{R}$ to F_p is a rounding function if*

- $\forall x \in F_p, \circ(x) = x$;
- $\forall (x, y) \in \mathbb{R}^2, x \leq y \Rightarrow \circ(x) \leq \circ(y)$.

Remark 2.2. *If $\circ$ is a rounding function, then for any x , $\circ(x) \in \{\text{RD}(x), \text{RU}(x)\}$, where RD and RU are the rounds-towards $-\infty$ and round-towards $+\infty$ rounding functions.*

The Fast2Sum algorithm was first introduced by Dekker [2]. It allows one to compute the error of a (rounded to nearest) floating-point addition. That algorithm is

ALGORITHM 1: Conventional Fast2Sum Algorithm.

$s \leftarrow \text{RN}(a + b)$
 $z \leftarrow \text{RN}(s - a)$
 $t \leftarrow \text{RN}(b - z)$

The conventional 2Sum algorithm, due to Knuth [9] and Møller [12], is

ALGORITHM 2: Conventional 2Sum algorithm.

(1) $s \leftarrow \text{RN}(a + b)$
(2) $a' \leftarrow \text{RN}(s - b)$
(3) $b' \leftarrow \text{RN}(s - a')$
(4) $\delta_a \leftarrow \text{RN}(a - a')$
(5) $\delta_b \leftarrow \text{RN}(b - b')$
(6) $t \leftarrow \text{RN}(\delta_a + \delta_b)$

We know that, in the absence of overflow, if the radix β of the floating-point system being used is less than or equal to 3, and if the floating-point exponents e_a and e_b of a and b satisfy $e_a \geq e_b$, then the values s and t returned by Algorithm 1 satisfy $s + t = a + b$, i.e., t is the error of the floating-point addition $s \leftarrow \text{RN}(a + b)$. Depending on the environment, testing the exponents of a and b may prove difficult. However if $|a| \geq |b|$ then $e_a \geq e_b$. Algorithm 2 gives the

same results as Algorithm 1, but without any requirement on β or on the exponents of a and b : it works in all cases provided that no overflow occurs. Due to the large penalty of a wrong branch prediction on modern architectures, if we do not have preliminary information on the respective orders of magnitude of a and b , calling the 6-operation 2Sum algorithm (Algorithm 2) is, in general, more efficient than comparing $|a|$ and $|b|$, swapping them if needed, and calling the 3-operation algorithm Fast2Sum (Algorithm 1).

Algorithms 1 and 2 allow one to compute the error of a floating-point addition, provided that that addition was performed using a round-to-nearest rounding function. The computed error can be re-injected later on in a calculation to compensate for it. This makes these “error free transformations” very useful. However, when a rounding function different from round-to-nearest is used, the error of a floating-point addition is not always equal to a floating-point number. For instance [13], in a radix-2 and precision- p arithmetic, assuming rounding toward $-\infty$, if $a = 1$ and $b = -2^{-3p}$, then

$$\begin{aligned} s &= \text{RD}(a + b) = 0.\underbrace{11111 \dots 11}_p \\ &= 1 - 2^{-p}, \end{aligned}$$

and

$$a + b - s = \underbrace{1.111111111 \dots 11}_{2p} \times 2^{-p-1},$$

which cannot be exactly represented with precision p (it would require precision $2p$).

Therefore, with rounding functions different from RN, it is important to know what Algorithms 1 and 2 (or, rather, a modified version, with different rounding functions, of these algorithms) will return, to know if they are still of any use.

This issue was already dealt with by Martin-Dorel, Melquiond, and Muller [11] in the restricted case where the rounding function is round to nearest with a possible “double rounding”.¹ Demmel and Nguyen show that if $4\text{ulp}(a) \leq |b| \leq a$ then Algorithm 1 returns the error of the floating-point addition of a and b when directed rounding functions are used.

Graillat, Jézéquel, and Picot [4] give an error bound on the value returned by Algorithm 1 when directed rounding functions are used. We will improve on their bound, showing that the algorithm always returns the best possible result, namely a floating-point number t nearest the error of the floating-point addition of a and b . We will perform a similar analysis with the 2Sum algorithm.

There is another issue with these two algorithms. One can rather easily convince oneself that they are immune to underflow. The main reason for that is that, as shown by Hauser [6], if the sum $a + b$ of two floating-point numbers

¹This happens when the result of an operation is first rounded to a wider floating-point format, before being rounded to the destination format.

is below the underflow threshold, then that sum is a floating-point number, which implies that it is computed exactly, with any rounding function (it can be viewed as a consequence of Lemma 2.4 below). It is, however, much more difficult to know if these algorithms are, at least for some restricted input domain, immune to *overflow*. More precisely, if the first operation (namely the floating-point addition of a and b) does not overflow, can one of the following operations overflow ?

The goal of this paper is to deal with these two issues, and to show that Fast2Sum and 2Sum (Algorithms 1 and 2) are much more robust than it is in general believed: for any combination of rounding functions (the rounding functions can vary at each step) they are immune to overflow (except for a very limited number of “extreme” cases that are easy to detect), and they always produce a very accurate estimate of the error of the floating-point addition $a + b$. The algorithms that we will analyze are the following:

ALGORITHM 3: Fast2Sum with faithful roundings: $\circ_1, \circ_2, \circ_3$ are rounding functions (see Definition 2.1).

$s \leftarrow \circ_1(a + b)$
 $z \leftarrow \circ_2(s - a)$
 $t \leftarrow \circ_3(b - z)$

ALGORITHM 4: 2Sum with faithful roundings: $\circ_i$, for $i = 1, \dots, 6$, are rounding functions (see Definition 2.1).

(1) $s \leftarrow \circ_1(a + b)$
(2) $a' \leftarrow \circ_2(s - b)$
(3) $b' \leftarrow \circ_3(s - a')$
(4) $\delta_a \leftarrow \circ_4(a - a')$
(5) $\delta_b \leftarrow \circ_5(b - b')$
(6) $t \leftarrow \circ_6(\delta_a + \delta_b)$

We will make much use of the following result, due to Sterbenz [20] (see for instance [6] or [13] for a proof).

Lemma 2.3 (Sterbenz). *In a radix- β floating-point system with subnormal numbers available, if x and y are finite floating-point numbers such that*

$$\frac{y}{2} \leq x \leq 2y,$$

then $x - y$ is a floating-point number.

Lemma 2.4 below is common computer arithmetic folklore. We give a proof of it for the sake of completeness.

Lemma 2.4. *Let a and b be two binary FP numbers of respective exponents e_a and e_b . Let $s \in \{\text{RD}(a + b), \text{RU}(a + b)\}$. If the exponent e_s of s is less than or equal to $\min(e_a, e_b)$ then $s = a + b$ exactly.*

Proof. First a and b are multiple of $2^{e_a - p + 1}$ and $2^{e_b - p + 1}$, respectively. Since $e_s \leq \min(e_a, e_b)$, the number $a + b$ is an integer multiple of $2^{e_s - p + 1}$. Hence, by

rounding it (through any rounding function) to a multiple of 2^{e_s-p+1} we just get it. Hence $s = a + b$. $\square$

The following lemma allows one to understand the behavior of the first two lines of Fast2Sum.

Lemma 2.5. *Let a and b be two binary FP numbers, with $e_a \geq e_b$. Let $s \in \{\text{RD}(a + b), \text{RU}(a + b)\}$. The number $s - a$ is a floating-point number (which implies that it will be computed exactly, with any rounding function).*

Notice that Lemma 2.5 only holds in radix 2. With floating-point systems of higher radices, we can build counter-examples. For instance, in radix 3 with $p = 4$ and $\circ = \text{RU}$, if $a = 1002_3 = 29_{10}$ and $b = 2222_3 = 80_{10}$, then $s = \text{RU}(a + b) = 11010_3 = 111_{10}$, so that $s - a = 10001_3 = 82_{10}$ is not exactly representable with precision 4.

Proof. We have $a = M_a \cdot 2^{e_a-p+1}$ and $b = M_b \cdot 2^{e_b-p+1}$, with $|M_a|, |M_b| \leq 2^p - 1$. Without loss of generality, we assume $M_a \geq 0$. Let M_s and e_s be the integral significand and the exponent of s , respectively.

1. if $e_s = e_a + 1$, then

$$M_s \in \left\{ \left\lfloor \frac{M_a}{2} + \frac{M_b}{2^{1+(e_a-e_b)}} \right\rfloor; \left\lceil \frac{M_a}{2} + \frac{M_b}{2^{1+(e_a-e_b)}} \right\rceil \right\}. \quad (1)$$

Defining $\mu = 2M_s - M_a$, from (1), we obtain

$$\frac{M_b}{2^{e_a-e_b}} - 2 < \mu < \frac{M_b}{2^{e_a-e_b}} + 2,$$

which implies $|\mu| \leq M_b + 1 \leq 2^p$. An immediate consequence is that $s - a = \mu \cdot 2^{e_a-p+1}$ is a floating-point number.

2. if $e_s \leq e_a$ then first notice that if $e_s \leq e_b$ then $s = a + b$ exactly by Lemma 2.4 so that $s - a = b$ is a floating-point number. Therefore we need only to focus on the case $e_s > e_b$. In that case

$$s \in \left\{ \left\lfloor 2^{e_a-e_s} M_a + 2^{e_b-e_s} M_b \right\rfloor \cdot 2^{e_s-p+1}; \left\lceil 2^{e_a-e_s} M_a + 2^{e_b-e_s} M_b \right\rceil \cdot 2^{e_s-p+1} \right\};$$

so that

$$(2^{e_b-e_s} M_b - 1) \cdot 2^{e_s-p+1} < s - a < (2^{e_b-e_s} M_b + 1) \cdot 2^{e_s-p+1}.$$

Hence $|s - a|$ is of the form $K \cdot 2^{e_s-p+1}$, with

$$|K| \leq \frac{|M_b|}{2} + 1 \leq 2^p,$$

which implies that it is a floating-point number. $\square$

The following remark shows that even when the rounding function is not round-to-nearest, the error of a floating-point addition will very frequently be exactly representable by a floating-point number.

Remark 2.6. *Let a and b be binary, precision- p , floating-point numbers. Let $s \in \{\text{RD}(a + b), \text{RU}(a + b)\}$. If the difference $|e_a - e_b|$ of the exponents of a and b does not exceed $p - 1$, then $s - (a + b)$ is a binary, precision- p , floating-point number.*

Proof. Without loss of generality, we assume $|a| \geq |b|$, and $e_a - e_b \leq p - 1$. The numbers a and b are multiple of $2^{e_b - p + 1}$, therefore $a + b$ and s are multiple of $2^{e_b - p + 1}$ too. Therefore, there exists an integer Z such that

$$(a + b) - s = Z \cdot 2^{e_b - p + 1}. \quad (2)$$

Let e_s be the FP exponent of s . Since $|s - (a + b)| < \text{ulp}(s)$, we have $|(a + b) - s| < 2^{e_s - p + 1}$. Since $|b| \leq |a|$, we have $|s| \leq 2a$, which implies $e_s \leq e_a + 1$. Therefore

$$|(a + b) - s| < 2^{e_a - p + 2} \leq 2^{e_b + 1}. \quad (3)$$

By combining (2) and (3) we deduce that $|Z| \leq 2^p - 1$, therefore $(a + b) - s$ is a FP number. $\square$

3 Accuracy of Fast2Sum in the absence of overflow

Let us first deal with Algorithm Fast2Sum with arbitrary rounding functions (Algorithm 3).

Theorem 3.1. *If no overflow occurs, and $e_a \geq e_b$ then the values s and t returned by Algorithm 3 satisfy*

$$t = \circ_3((a + b) - s),$$

i.e., t is a faithful rounding of the error of the FP addition $s \leftarrow \circ_1(a + b)$.

Notice that if we combine this theorem with Remark 2.6, we deduce that if the difference of the exponents of a and b does not exceed $p - 1$ (which will occur in many practical cases), then t is exactly $(a + b) - s$.

Proof. Lemma 2.5 above implies that $s - a$ is a floating-point number. Hence, $z = s - a$, so that

$$t = \circ_3(b - z) = \circ_3((a + b) - s).$$

$\square$

4 Accuracy of TwoSum in the absence of overflow

We now consider Algorithm 2Sum with arbitrary rounding functions (Algorithm 4). Contrarily to what happened in the previous section with Algorithm 3, we will not always obtain a final value t equal to a faithful rounding of $(a + b) - s$. Consider the following example, in binary32/single-precision arithmetic ($p = 24$):

- $a = 3076485 \cdot 2^{-21}, b = -6130317 \cdot 2^{-49};$
- $\circ_1 = \circ_2 = \circ_5 = \text{RU}, \circ_3 = \circ_4 = \circ_6 = \text{RD}.$

We successively obtain:

$$\begin{aligned} s &= a = 3076485 \cdot 2^{-21}; \\ a' &= 12305941 \cdot 2^{-23}; \\ b' &= -2^{-23}; \\ \delta_a &= -2^{-23}; \\ \delta_b &= 15244637 \cdot 2^{-47}; \\ t &= -1532579 \cdot 2^{-47}. \end{aligned}$$

and since $(a+b) - s = b$ is a floating-point number, with any rounding function $\circ$, $\circ((a+b) - s) = b$ will be different from t . However, $(a+b-s) - t = -2^{-49}$, so that t remains a very good approximation to $(a+b) - s$. As we are going to see, this will always be true. More precisely, we will prove together the following two results. The first one (Theorem 4.1) is the main result of this section. The second one (Lemma 4.2) will be used in the proof of Theorem 6.2.

Theorem 4.1. *If $p \geq 4$ and no overflow occurs, then the values s and t returned by Algorithm 4 satisfy*

$$t = (a+b) - s + \alpha,$$

with $|\alpha| < 2^{-p+1} \cdot \text{ulp}(a+b) \leq 2^{-p+1} \cdot \text{ulp}(s)$. Furthermore, if the floating-point exponents e_s and e_b of s and b satisfy $e_s - e_b \leq p-1$ then t is a faithful rounding of $(a+b) - s$.

Lemma 4.2. *If $p \geq 4$ and no overflow occurs in lines (1) to (5) of Algorithm 4, then the variables δ_a and δ_b computed at lines (4) and (5) satisfy*

$$|\delta_a + \delta_b| \leq \text{ulp}(a+b).$$

Proof. We prove together Theorem 4.1 and Lemma 4.2. This means the case split and intermediate results are the same, but they do not rely one on another. Without loss of generality, we assume $a \geq 0$. The following figure illustrates the various cases that are discussed in the proof.

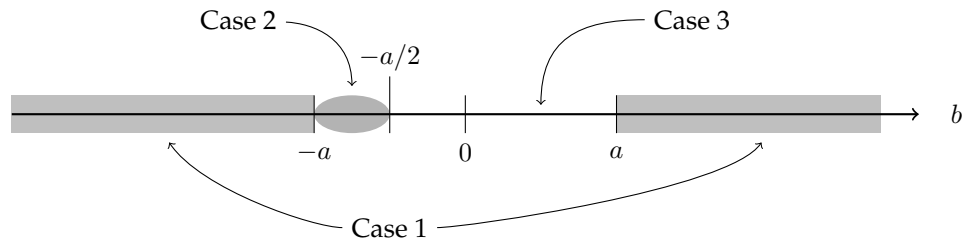


Figure 1: Various cases discussed in the proof.

1. If $|b| \geq a$ then lines (1), (2), and (4) of Algorithm 4 are equivalent to Fast2Sum(b, a). Therefore, from Theorem 3.1, we have $a' = s - b$ and $\delta_a = \circ_4(a + b - s)$, hence $|\delta_a| \leq \text{ulp}(a + b)$. An immediate consequence of $a' = s - b$ is $b' = b$ and $\delta_b = 0$. From this, we find

$$t = \circ_4(a + b - s)$$

and $|\delta_a + \delta_b| \leq \text{ulp}(a + b)$, so that the result of Lemma 4.2 holds.

2. If $|b| < a$ and $|s| \leq |b|$ (which is equivalent to saying that $-a < b \leq -a/2$) then by Sterbenz Lemma, $s = a + b$. An immediate consequence is $a' = a$, $b' = b$, $\delta_a = \delta_b = 0$ (so that, obviously, the result of Lemma 4.2 holds), $t = 0$. Hence $t = (a + b) - s$.
3. If $|b| < a$ and $|s| > |b|$ (which is equivalent to saying that $-a/2 < b < a$), notice that we have $s > 0$. Let $u = 2^{1-p}$ (i.e., u is the rounding unit for directed roundings). We have

$$\begin{aligned} s &= (a + b) \cdot (1 + \epsilon_1); \text{ with } |\epsilon_1| \leq u; \\ a' &= (s - b) \cdot (1 + \epsilon_2); \text{ with } |\epsilon_2| \leq u. \end{aligned}$$

Thus $a' = (a + a\epsilon_1 + b\epsilon_1) \cdot (1 + \epsilon_2)$. Since $|b| < a$, $a\epsilon_1 + b\epsilon_1$ can be written $2a\epsilon_3$, with $|\epsilon_3| \leq u$. Therefore

$$a' = a \cdot (1 + \eta),$$

with $|\eta| \leq 3u + 2u^2$. As soon as $p \geq 4$, we have $|\eta| < 1/2$, so that $a' \geq 0$ and $a/2 \leq a' \leq 2a$. Therefore, Sterbenz Lemma applies to line (4) of Algorithm 4, and

$$\delta_a = a - a'. \quad (4)$$

Also, since $s > |b|$, Lines (2), (3), and (5) of Algorithm 4 are equivalent to Fast2Sum($s, -b$), so that

$$b' = s - a', \quad (5)$$

and

$$\delta_b = \circ_5(a' - (s - b)). \quad (6)$$

Notice that, from Remark 2.6, as soon as the exponents e_s and e_b of s and b satisfy $e_s - e_b \leq p - 1$, (6) implies $\delta_b = a' - (s - b)$, from which we easily deduce $t = \circ_6(a + b - s)$. Also, in that case, $\delta_a + \delta_b = (a + b) - s$, so that Lemma 4.2 holds. Hence, *let us now assume that $e_s - e_b \geq p$* . Notice that this implies

$$|b| < 2^{e_b+1} \leq 2^{e_s-p+1} = \text{ulp}(s).$$

Hence,

$$a' \in \{\text{succ}(s), s, \text{pred}(s), \text{pred}(\text{pred}(s))\}.$$

Notice that the case $a' = \text{pred}(\text{pred}(s))$ can occur only when s is a power of 2.

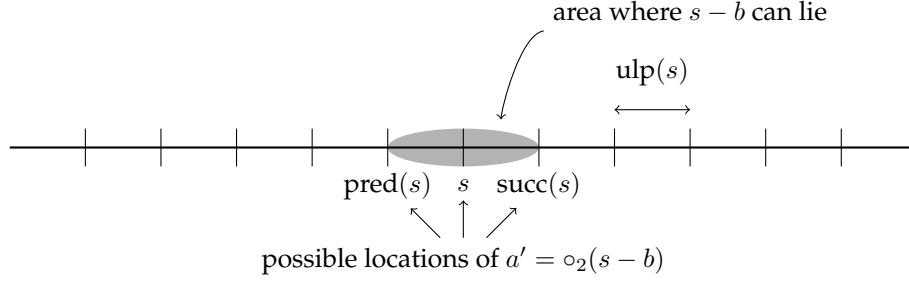


Figure 2: General case: s is not a power of 2

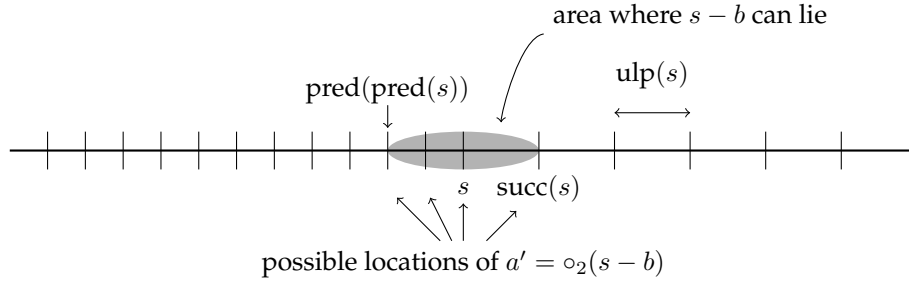


Figure 3: Special case: s is a power of 2

- (a) **If $a' = s$ then $b' = 0$.** It follows that $\delta_b = b$ and $\delta_a = a - s$, for which we deduce $t = \circ_6(\delta_a + \delta_b) = \circ_6(a + b - s)$, and $|\delta_a + \delta_b| = |a + b - s| < \text{ulp}(a + b)$, so that Lemma 4.2 holds.
- (b) **If $a' \neq s$ then**

$$a' = s - \sigma \cdot \text{ulp}(s), \text{ with } \sigma \in \{-1, 1/2, 1\},$$

(the case $\sigma = 1/2$ can occur only when s is a power of 2), and we have

$$\begin{aligned} b' &= \sigma \cdot \text{ulp}(s) \\ \delta_a &= a - s + \sigma \cdot \text{ulp}(s) \\ \delta_b &= \circ_5(b - \sigma \cdot \text{ulp}(s)). \end{aligned}$$

We know that $|b| < \text{ulp}(s)$. Furthermore, b has the same sign as σ . Therefore

- either $|b| \geq |\sigma|/2 \cdot \text{ulp}(s)$, in which case Sterbenz Lemma implies $\delta_b = b - \sigma \cdot \text{ulp}(s)$, so that $\delta_a + \delta_b = a + b - s$, so that $t = \circ_6(a + b - s)$;
- or $|b| < |\sigma|/2 \cdot \text{ulp}(s)$, in which case

$$|b - \sigma \cdot \text{ulp}(s)| < |\sigma| \cdot \text{ulp}(s)$$

(unless $b = 0$ but that case is straightforwardly handled), so that (since $|\sigma| \cdot \text{ulp}(s)$ is a power of 2)

$$|\delta_b - (b - \sigma \cdot \text{ulp}(s))| < \frac{|\sigma|}{2} \text{ulp}(\text{ulp}(s)) = |\sigma| \cdot 2^{-p} \text{ulp}(s)$$

(since $\text{ulp}(s)$ is a power of 2). An immediate consequence is

$$|(\delta_a + \delta_b) - (a + b - s)| < |\sigma| \cdot 2^{-p} \text{ulp}(s). \quad (7)$$

Since we already know that $|(a + b) - s| < \text{ulp}(a + b)$, we deduce

$$|\delta_a + \delta_b| < \text{ulp}(a + b) + |\sigma| \cdot 2^{-p} \text{ulp}(s). \quad (8)$$

Let us try to slightly improve on the bound (8). First, from $|b| \leq \text{ulp}(s)$ one easily deduces $a > s/2$ (otherwise, we would have $a + b \leq s/2 + \text{ulp}(s)$, which would imply $s = \circ_1(a + b) \leq \circ_1(s/2 + \text{ulp}(s)) = s/2 + \text{ulp}(s)$). Hence δ_a is a multiple of $\frac{1}{2} \text{ulp}(s)$. Also, $\text{ulp}(a + b)$ is equal to $\text{ulp}(s)$ or $\frac{1}{2} \text{ulp}(s)$. Finally, $|b - b'| > \frac{|\sigma|}{2} \text{ulp}(s)$, so that $|\delta_b| \geq \frac{|\sigma|}{2} \text{ulp}(s)$, which implies that δ_b is a multiple of $|\sigma| \cdot 2^{-p} \text{ulp}(s)$. All this implies that $\delta_a + \delta_b$ is a multiple of $|\sigma| \cdot 2^{-p} \text{ulp}(s)$. Hence, from (8), we deduce

$$|\delta_a + \delta_b| \leq \text{ulp}(a + b).$$

First, this proves Lemma 4.2. Furthermore, since $\text{ulp}(a + b)$ is a power of 2, we obtain

$$|\circ_6(\delta_a + \delta_b) - (\delta_a + \delta_b)| \leq \frac{1}{2} \text{ulp}(\text{ulp}(a + b)) = 2^{-p} \text{ulp}(a + b).$$

Combined with (7), this gives

$$|t - (a + b - s)| < 2^{-p} \cdot (\text{ulp}(a + b) + |\sigma| \cdot \text{ulp}(s)). \quad (9)$$

This already gives $|t - (a + b - s)| < 2^{-p+1} \cdot \text{ulp}(s)$. Let us now try to express a bound on $|t - (a + b - s)|$ as a function of $\text{ulp}(a + b)$ only. We have four cases to consider

- i. if s is not a power of 2, or if $a + b \geq s$, then $\text{ulp}(a + b) = \text{ulp}(s)$, which gives $|t - (a + b - s)| < 2^{-p+1} \cdot \text{ulp}(a + b)$;
- ii. if s is a power of 2 and $a + b < s$ and $\sigma = 1/2$, then $\text{ulp}(a + b) = 1/2 \cdot \text{ulp}(s)$, and (9) implies $|t - (a + b - s)| < 2^{-p+1} \cdot \text{ulp}(a + b)$;
- iii. the case when s is a power of 2, $a + b < s$, and $\sigma = 1$ is impossible: we assumed $|b| < |\sigma|/2 \cdot \text{ulp}(s) = 1/2 \cdot \text{ulp}(s)$, which implies $s - b \geq s - 1/2 \cdot \text{ulp}(s) = \text{pred}(s)$, which implies $a' = \circ_2(s - b) \geq \text{pred}(s)$, which is not compatible with the assumption $\sigma = 1$, since $a' = s - \sigma \text{ulp}(s)$;

- iv. if s is a power of 2, $a + b < s$, and $\sigma = -1$, we have the following relations (see Fig. 4): $a' = \text{succ}(s) = s + \text{ulp}(s)$, $b' = -\text{ulp}(s)$, and $-1/2 \cdot \text{ulp}(s) < b < 0$. We deduce that $a > \text{pred}(s) = s - \frac{1}{2} \text{ulp}(s)$ (otherwise we would have $a + b < \text{pred}(s)$, which would imply $s = \circ_1(a + b) \leq \text{pred}(s)$). Similarly, we have $a < \text{succ}(s) = s + \text{ulp}(s)$ (otherwise, we would have $a + b \geq \text{succ}(s) - \frac{1}{2} \text{ulp}(s) > s$). Therefore $a = s$, from which we immediately deduce $\delta_a = -\text{ulp}(s)$ and $\delta_b = \circ_5(b + \text{ulp}(s))$. Now, δ_a and δ_b have opposite signs, and

$$\frac{1}{2} \text{ulp}(s) < b + \text{ulp}(s) < \text{ulp}(s),$$

from which we deduce

$$\frac{|\delta_a|}{2} = \frac{1}{2} \text{ulp}(s) \leq \circ_5(b + \text{ulp}(s)) = \delta_b \leq |\delta_a| = \text{ulp}(s),$$

hence we can apply Sterbenz lemma to the addition of δ_a and δ_b , which gives

$$\begin{aligned} t = \circ_6(\delta_a + \delta_b) &= \delta_a + \delta_b \\ &= -\text{ulp}(s) + \circ_5(b + \text{ulp}(s)) \\ &= b + \eta, \end{aligned}$$

with $|\eta| < 2^{-p} \cdot \text{ulp}(s) = 2^{-p+1} \text{ulp}(a + b)$.

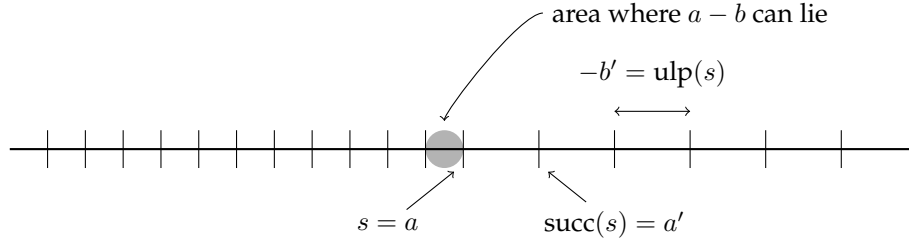


Figure 4: Last case: s is a power of 2, $a + b < s$, and $\sigma = -1$

□

5 Fast2Sum is immune to overflow

Let us now consider Algorithm 3, with $e_a \geq e_b$, where e_a and e_b are the floating-point exponents of a and b , and let us assume that no overflow occurred in the first line ($s \leftarrow \circ_1(a + b)$). Without loss of generality, we can assume $a > 0$. Let us first deal with the second line of the algorithm ($z \leftarrow \circ_2(s - a)$).

We have $s = a + b + \epsilon$, with $|\epsilon| < \text{ulp}(a + b) \leq 2\text{ulp}(a)$. Hence $s - a = b + \epsilon$. Therefore, if the computation of $z = \circ_2(s - a)$ overflows, it means that either $b < -\Omega + 2\text{ulp}(a)$ or $b > \Omega - 2\text{ulp}(a)$.

The second case is impossible: if $b > \Omega - 2\text{ulp}(a) \geq \Omega - 2\text{ulp}(\Omega)$, then (since $e_a \geq e_b$, which here implies $e_a = e_b = e_{\max}$) then $a + b \geq \Omega - 2\text{ulp}(\Omega) + 2^{e_{\max}} = 3 \cdot 2^{e_{\max}} - 3 \cdot 2^{e_{\max}-p+1}$, which implies that $a + b$ overflows. Let us consider the first case. In that case, we have $b < -\Omega + 2\text{ulp}(a) \leq -\Omega + 2\text{ulp}(\Omega)$ and (since $e_a \geq e_b$ which implies here $e_a = e_{\max}$), $\Omega/2 < 2^{e_{\max}} \leq a \leq \Omega$, in the first operation we are in the conditions of Sterbenz Lemma, so that $s = a + b$, which implies $z = b$: in that case the computation of z does not overflow.

Hence, in all cases, the second line of Algorithm 3 cannot overflow. Let us now deal with the last line ($t \leftarrow \circ_3(b - z)$). We know from Lemma 2.5 that $z = s - a$, so that $b - z = a + b - s$. The computation of t can overflow only if $|b - z| > \Omega$, but this is impossible since

$$|b - z| = |(a + b) - s| < \text{ulp}(s) < |s|.$$

We immediately deduce

Theorem 5.1. *Assume that we perform Algorithm 3 with input values a and b whose exponents satisfy $e_a \geq e_b$. If the computation of s (first line of the algorithm) does not overflow, then the other lines of the algorithm cannot overflow.*

6 2Sum is almost immune to overflow

The overflow analysis of Algorithm 4 will be significantly more difficult. Our main result is Theorem 6.2 below. To make its proof simpler, we first prove the following result.

Lemma 6.1. *If there are no overflows at lines (1) to (5) of Algorithm 4, there cannot be an overflow at line (6).*

Proof. From Lemma 4.2, we know that $|\delta_a + \delta_b| \leq \text{ulp}(a + b)$. Since no overflow occurs at line (1), $a + b$ is in the representable range, so that $\text{ulp}(a + b) \leq 2^{-p+1}|a + b|$ is much below the overflow threshold. Hence line (6) of Algorithm 4 (namely, $t \leftarrow \circ_6(\delta_a + \delta_b)$) cannot overflow. $\square$

Theorem 6.2. *If the first input value a of Algorithm 4 satisfy $|a| < \Omega$ and if there is no overflow at line (1) of the algorithm, then there will be no overflow at lines (2) to (6).*

Proof. Without loss of generality, we assume $a > 0$. Assume that no overflow occurred in the first line ($s \leftarrow \circ_1(a + b)$).

1. If $b \geq 0$

The monotonicity of the rounding functions implies: i) $s \geq \circ_1(b) = b$, so that $a' \geq \circ_2(0) = 0$; and ii) $a' \leq \circ_2(s) = s$. Therefore

$$0 \leq a' \leq s, \tag{10}$$

which implies that there is no overflow at line (2) of the algorithm. Now, (10) implies $0 \leq s - a' \leq s$, so that

$$0 \leq b' \leq s. \quad (11)$$

As a consequence, there is no overflow at line (3) of the algorithm.

Now, since $a \geq 0$ and $a' \geq 0$, we deduce $|a - a'| \leq \max\{a, a'\}$, hence line (4) cannot overflow.

Similarly, since $b \geq 0$ and $b' \geq 0$, we obtain $|b - b'| \leq \max\{b, b'\}$, hence line (5) cannot overflow.

Lemma 6.1 implies that line (6) cannot overflow.

2. If $b \leq 0$

Notice that there cannot be an overflow at line (1): $|a + b|$ (hence $|s|$) is less than or equal to $\max\{|a|, |b|\}$.

(a) if $-b < a$, then $a + b - \text{ulp}(a + b) < s < a + b + \text{ulp}(a + b)$, so that

$$a + b - \text{ulp}(a) < s < a + b + \text{ulp}(a) \quad (12)$$

(since $|a + b| < a$, which implies $\text{ulp}(a + b) \leq \text{ulp}(a)$). We therefore deduce

$$a - \text{ulp}(a) < s - b < a + \text{ulp}(a).$$

therefore, unless $a = \Omega$, there will be no overflow at line (2) of the algorithm, and $a' = \circ_2(s - b)$ will satisfy

$$a - \text{ulp}(a) \leq a' \leq a + \text{ulp}(a). \quad (13)$$

(this is deduced using the monotonicity of the rounding function $\circ_2$ and the fact that $a - \text{ulp}(a)$ and $a + \text{ulp}(a)$ are floating-point numbers). We now assume $a \neq \Omega$ (which, with our assumption $-b < a$, implies $-b < \text{pred}(\Omega)$), i.e., since b is a floating-point number,

$$|b| = -b \leq \text{pred}(\text{pred}(\Omega)). \quad (14)$$

From (12) and (13), we find

$$b - 2\text{ulp}(a) < s - a' < b + 2\text{ulp}(a),$$

This, along with (14) and $\text{ulp}(a) \leq \text{ulp}(\Omega)$ implies that line (3) of the algorithm cannot overflow. Furthermore, $-b < a$ implies $|b - 2\text{ulp}(a)| < 2a$ so that $\text{ulp}(b \pm 2\text{ulp}(a)) \leq 2\text{ulp}(a)$. This gives

$$b - 4\text{ulp}(a) \leq b' \leq b + 4\text{ulp}(a). \quad (15)$$

Now, from (13) and (15), we deduce $|a - a'| \leq \text{ulp}(a)$ and $|b - b'| \leq 4\text{ulp}(a)$, so that lines (4) and (5) of the algorithm cannot overflow. Lemma 6.1 implies that line (6) cannot overflow.

3. **if $-b \geq a$.** First, notice that the case $a \geq -b/2$ is easily handled, since Sterbenz Lemma applied to line (1) of Algorithm 4 implies $s = a + b$, so that $a = a'$, $b = b'$, and $\delta_a = \delta_b = t = 0$. Hence we only need to focus on the case $a < -b/2$.

From $0 \leq a < -b/2$ we deduce $b \leq a + b < b/2$, which implies²

$$b \leq s \leq b/2. \quad (16)$$

The consequence of (16) is twofold. First, we immediately deduce $0 \leq s - b \leq -b/2$, so that Line (2) of Algorithm 4 cannot overflow, and second, Sterbenz Lemma can be applied to line (2) of Algorithm 4, so that $a' = s - b$. It follows that $b' = b$ and Line (3) cannot overflow. Therefore $a - a' = a + b - s$, so that $|a - a'| < \text{ulp}(a + b)$, hence Line (4) cannot overflow. We finally have $\delta_b = b - b' = 0$ and $t = \delta_a$: lines (5) and (6) cannot overflow.

□

Notice that condition $|a| < \Omega$ is necessary. Assume all rounding functions are RN (with ties-to-even tie-breaking rule). The choice $a = \Omega$ and $b = -(3/2) \cdot \text{ulp}(\Omega)$ will give no overflow at line (1), and an overflow at line (2).

Conclusion

We have shown that, in binary floating-point arithmetic, the 2Sum and Fast2Sum algorithms are more “robust” than it is usually believed: even when the error of the initial floating-point addition is not exactly representable, they return a very good approximation to that error. Also, they are almost totally immune to overflow: the only case where a “spurious” overflow may occur is with 2Sum, when the absolute value of operand a is equal to the largest floating-point number.

Acknowledgement

This work was supported by the FastRelax (ANR-14-CE25-0018-01) project of the French National Agency for Research (ANR).

References

- [1] S. Boldo and M. Daumas. Representable correcting terms for possibly underflowing floating point operations. In J.-C. Bajard and M. Schulte, editors, *Proceedings of the 16th Symposium on Computer Arithmetic*, pages 79–86. IEEE Computer Society Press, Los Alamitos, CA, 2003.

²Unless $b/2$ is not a floating-point number: this can happen only if b is subnormal, and in that case, with $0 \leq a < -b/2$, overflow is of course impossible.

- [2] T. J. Dekker. A floating-point technique for extending the available precision. *Numerische Mathematik*, 18(3):224–242, 1971.
- [3] J. Demmel and H. D. Nguyen. Fast reproducible floating-point summation. In *21st IEEE Symposium on Computer Arithmetic, Austin, TX, USA, April 7-10*, pages 163–172, 2013.
- [4] S. Graillat, F. Jézéquel, and R. Picot. Numerical validation of compensated summation algorithms with stochastic arithmetic. *Electronic Notes in Theoretical Computer Science*, 317:55 – 69, 2015. The Seventh and Eighth International Workshops on Numerical Software Verification (NSV).
- [5] S. Graillat, P. Langlois, and N. Louvet. Algorithms for accurate, validated and fast computations with polynomials. *Japan Journal of Industrial and Applied Mathematics*, 26(2):215–231, 2009.
- [6] J. R. Hauser. Handling floating-point exceptions in numeric programs. *ACM Trans. Program. Lang. Syst.*, 18(2):139–174, 1996.
- [7] Y. Hida, X. S. Li, and D. H. Bailey. Algorithms for quad-double precision floating-point arithmetic. In N. Burgess and L. Ciminiera, editors, *Proceedings of the 15th IEEE Symposium on Computer Arithmetic (ARITH-16)*, pages 155–162, Vail, CO, June 2001.
- [8] W. Kahan. Pracniques: further remarks on reducing truncation errors. *Commun. ACM*, 8(1):40, 1965.
- [9] D. Knuth. *The Art of Computer Programming*, volume 2. Addison-Wesley, Reading, MA, 3rd edition, 1998.
- [10] U. W. Kulisch. An axiomatic approach to rounded computations. *Numerische Mathematik*, 19:1–17, 1971.
- [11] É. Martin-Dorel, G. Melquiond, and J.-M.s Muller. Some issues related to double rounding. *BIT Numerical Mathematics*, 53(4):897–924, 2013.
- [12] O. Møller. Quasi double-precision in floating-point addition. *BIT*, 5:37–50, 1965.
- [13] J.-M. Muller, N. Brisebarre, F. de Dinechin, C.-P. Jeannerod, V. Lefèvre, G. Melquiond, N. Revol, D. Stehlé, and S. Torres. *Handbook of Floating-Point Arithmetic*. Birkhäuser Boston, 2010.
- [14] A. Neumaier. Rundungsfehleranalyse einiger Verfahren zur Summation endlicher Summen. *ZAMM*, 54:39–51, 1974. In German.
- [15] D. M. Priest. Algorithms for arbitrary precision floating point arithmetic. In P. Kornerup and D. W. Matula, editors, *Proceedings of the 10th IEEE Symposium on Computer Arithmetic (Arith-10)*, pages 132–144. IEEE Computer Society Press, Los Alamitos, CA, June 1991.

- [16] D. M. Priest. *On Properties of Floating-Point Arithmetics: Numerical Stability and the Cost of Accurate Computations*. PhD thesis, University of California at Berkeley, 1992.
- [17] S. M. Rump, T. Ogita, and S. Oishi. Accurate floating-point summation part II: Sign, K-fold faithful and rounding to nearest. *SIAM Journal on Scientific Computing*, 2005–2008. Submitted for publication.
- [18] S. M. Rump, T. Ogita, and S. Oishi. Accurate floating-point summation part I: Faithful rounding. *SIAM Journal on Scientific Computing*, 31(1):189–224, 2008.
- [19] J. R. Shewchuk. Adaptive precision floating-point arithmetic and fast robust geometric predicates. *Discrete Computational Geometry*, 18:305–363, 1997.
- [20] P. H. Sterbenz. *Floating-Point Computation*. Prentice-Hall, Englewood Cliffs, NJ, 1974.