

How to Send a Real Number Using a Single Bit (and Some Shared Randomness)

Ran Ben Basat

University College London

Michael Mitzenmacher

Harvard University

Shay Vargaftik

VMware Research

Abstract

We consider the fundamental problem of communicating an estimate of a real number $x \in [0, 1]$ using a single bit. A sender that knows x chooses a value $X \in \{0, 1\}$ to transmit. In turn, a receiver estimates x based on the value of X . The goal is to minimize the cost, defined as the worst-case (over the choice of x) expected squared error.

We first overview common biased and unbiased estimation approaches and prove their optimality when no shared randomness is allowed. We then show how a small amount of shared randomness, which can be as low as a single bit, reduces the cost in both cases. Specifically, we derive lower bounds on the cost attainable by any algorithm with unrestricted use of shared randomness and propose optimal and near-optimal solutions that use a small number of shared random bits. Finally, we discuss open problems and future directions.

2012 ACM Subject Classification Theory of computation → Rounding techniques

Keywords and phrases Randomized Algorithms, Approximation Algorithms, Shared Randomness

1 Introduction

We consider the fundamental problem of communicating an estimate of a real number $x \in [0, 1]$ using a single bit. A sender, that we call *Buffy*, knows x , and chooses a value $X \in \{0, 1\}$ to transmit. In turn, a receiver, that we call *Angel*, estimates x based on the value of X .

This problem naturally appears in distributed computations where multiple machines perform parallel tasks and transmit their results to an aggregator. If communication to the aggregator is limited, the machines must send estimates of the results. In particular, we are motivated by recent work addressing the communication bottleneck in distributed and federated machine learning [9]. There, *clients* compute a local gradient and send it to a *parameter server* that computes the global gradient and updates the model [11]. For the typical large-scale federated learning problems over edge devices (e.g., mobile phones), the devices may only be able to communicate a small number of bits per gradient coordinate. In fact, solutions such as 1-Bit SGD [15] and signSGD [3], have recently been studied as appealing low-communication solutions that use a single bit per coordinate. Another common communication-efficient solution is TernGrad [17] that quantizes each coordinate to $\{-1, 0, 1\}$ instead of $\{-1, 1\}$ as commonly done by the 1-bit algorithms.

It is often desirable that each estimate be an *unbiased* random variable with a mean equal to corresponding x . For example, this provides that the estimates' average is an unbiased estimate of the average value. Alternatively, there are cases where it is beneficial to allow *biased* estimates if it reduces the error for that setting (e.g., see EF-signSGD [10]).

In this work, we consider several variations of the above problem. For algorithms that provide unbiased estimates for every value x , we use the *worst-case* (over all values of x) variance as the cost function to be minimized. For biased estimates, we consider the worst-case

expected squared error as the cost, as it coincides with variance for unbiased algorithms. That is, the worst-case is over the value of x , and the expectation of the cost is over the random choices used by the algorithm. Note that any lower bound for biased algorithms with these costs also applies to unbiased algorithms, and an upper bound for unbiased algorithms also applies to biased algorithms. We are interested in both lower and upper bounds in our work.

Beyond unbiased and biased variations, we also consider settings where Buffy and Angel have access to *shared randomness*. Shared randomness (often also referred to as public or common randomness) has been intensively studied in the field of communication complexity (e.g., see [12]). In our context, such shared randomness can arise naturally by having Buffy and Angel share a common seed for a pseudo-random number generator, for example. Here, we model the shared randomness as “perfectly random,” leaving issues related to pseudo-randomness aside. Nevertheless, we consider solutions using limited amounts of shared randomness, including the case of just one bit of shared randomness. Such solutions may be easier and cheaper to implement, including with pseudo-random generators.

We remark that there are known approaches to this problem. These include (deterministic) rounding, randomized rounding (also called stochastic quantization), and subtractive dithering [13]. For a detailed survey of such techniques, we refer the reader to [5]. We discuss these methods and compare our results with them in context throughout the paper.

Our contribution: In this paper, we study how to minimize the cost (i.e., the worst-case variance or worst-case expected squared error) for various settings. First, we consider the setting where there is no shared randomness. In this setting, we show that randomized rounding is the optimal unbiased algorithm and that deterministic rounding is optimal when biased estimations are allowed. While these algorithms are widely used in practice, the optimality proofs under these cost models have not appeared elsewhere to the best of our knowledge.

Next, we explore how to reduce the cost if Buffy and Angel have access to shared randomness. We prove upper and lower bounds on the attainable variance for unbiased algorithms and expected squared error for biased ones. For our upper bounds, we assume that Buffy and Angel have access to ℓ shared random bits, for some $\ell \in \mathbb{N}$. We also consider the *limiting algorithms* where ℓ is not restricted. Our work addresses several extensions for cases where unbounded private randomness is allowed and when it is not. Finally, we consider the special case where x is known to be in $\{0, 1/2, 1\}$, a setting that is of high interest, for example, for the sign-based federated learning algorithms (e.g., [3, 10]) and particularly for TernGrad [17] that uses 3-level quantization. We provide an improved algorithm and a matching lower bound for this setting, thus proving its optimality. A summary of our results appears in Table 1.

2 Preliminaries

We start with some notation. We use $[n]$ to denote $\{0, 1, \dots, n-1\}$, and $\Delta(S)$ to denote all possible probability distributions over the set S . (An element of $\Delta(S)$ will be expressed as a density function when S is uncountable, e.g., if $S = [0, 1]$.) We also use, for a binary predicate B , $\mathbb{1}_B$ as an indicator such that $\mathbb{1}_B = 1$ if B is true and 0 otherwise. Lastly, $\phi = (1 + \sqrt{5})/2$ denotes the Golden Ratio, which naturally comes up in many of our results.

Problem statement: Given a real number $x \in [0, 1]$, Buffy compresses it to a single bit value $X \in \{0, 1\}$ that is sent to Angel, who derives an estimate $\hat{x}$. We also consider the special case where x is known to be in $\{0, 1/2, 1\}$. Our objective is to minimize the *cost* that is defined as the *worst-case expected squared error*, i.e., $\max_{x \in [0, 1]} \mathbb{E}[(\hat{x} - x)^2]$. Note that the worst-case is taken over the value of x and the expectation is over the randomness of the algorithm. In the

Scenario	Unbiased (Variance)	Biased (Exp. Squared Error)
No shared randomness	$1/4 = 0.25$ (randomized rounding) Optimal (Section 3.1)	$1/16 = 0.0625$ (deterministic rounding) Optimal (Section 3.2)
ℓ -bit shared randomness Unbounded private randomness	$\ell = 1 : \frac{1}{8} = 0.125$ $\ell = 8 : \frac{1}{12} + \frac{1}{393216} \approx 0.08334$ In general: $1/6 \cdot (1/2 + 4^{-\ell})$ (Section 5)	Open
ℓ -bit shared randomness No private randomness	Impossible (Section 6)	$\ell = 1 : \frac{1}{20} = 0.05$ $\ell = 8 : \approx 0.04599$ (Section 6.2.4)
Lower Bounds for $x \in [0, 1]$ Unbounded shared randomness	$1/16 = 0.0625$ (Section 4.2)	≈ 0.0459 (Section 4.1.2)
$x \in \{0, 1/2, 1\}$ 1-bit shared randomness No private randomness	$1/16 = 0.0625$ (Section 5)	$3/4 - 1/\sqrt{2} \approx 0.04289$ (Section 6.1)
Lower Bounds for $x \in \{0, 1/2, 1\}$ Unbounded shared randomness	$1/16 = 0.0625$ (Section 4.2)	$3/4 - 1/\sqrt{2} \approx 0.04289$ (Section 4)

■ **Table 1** A summary of our results.

unbiased setting, we additionally require $\mathbb{E}[\hat{x}] = x$, in which case the cost becomes $\text{Var}[\hat{x}]$, i.e., the estimation variance. In some cases, we allow the parties to use ℓ bits of *shared randomness*. That is, we assume that they have access to a random value $h \in [2^\ell]$, known to both Buffy and Angel. When applicable, we use $r \in [0, 1]$ to denote the private randomness of Buffy.

3 Algorithms without Shared Randomness

We recap the performance of two standard algorithms – randomized and deterministic rounding. Interestingly, we show that when no shared randomness is allowed, randomized rounding is an optimal unbiased algorithm, and deterministic rounding is an optimal biased algorithm.

3.1 Randomized Rounding

In randomized rounding, Buffy uses private randomness to generate $X \sim \text{Bernoulli}(x)$ which is sent using a single bit. In turn, Angel estimates $\hat{x} = X$. Clearly, we have that $\mathbb{E}[\hat{x}] = \mathbb{E}[X] = x$, and thus the algorithm is unbiased. The variance of the algorithm is $\text{Var}[\hat{x}] = \text{Var}[X] = x(1-x)$, and thus the worst-case is reached at $x = 1/2$, which gives a cost of $1/4$. The following theorem, whose proof is deferred to Appendix A, shows that randomized rounding is optimal, in the sense that no unbiased algorithm without shared randomness can have a worst-case variance lower than $1/4$. Intuitively, requiring the estimate to be unbiased forces the algorithm to send 1 with a probability that is linear in x , maximizing its cost for $x = 1/2$. The proof also establishes the intuitive idea that it is not possible to benefit from randomness used solely by Angel.

► **Theorem 1.** *Any unbiased algorithm without shared randomness must have a worst-case variance of at least $1/4$.*

3.2 Deterministic Rounding

With deterministic rounding, Buffy sends $X = 1$ when $x \geq 1/2$. Angel then estimates $\hat{x} = X/2 + 1/4$. Deterministic rounding has an (absolute) error of at most $1/4$, which is achieved for $x \in \{0, 1/2, 1\}$. Therefore, its cost is $1/16$. The next theorem, whose proof appears in Appendix B, shows that deterministic rounding is optimal, as no algorithm that does not use shared randomness can have a lower cost (even with unrestricted private randomness). We show that any such algorithm must have an expected squared error of at least $1/16$ on at least one of $\{0, 1/2, 1\}$.

► **Theorem 2.** *Any algorithm without shared randomness must have a worst-case expected squared error of at least $1/16$.*

4 Lower Bounds

We next explore lower bounds for algorithms with shared randomness. We use Yao's minimax principle [18] to prove a lower bound on the cost of any biased shared randomness protocol. Then, we show a stronger lower bound for unbiased algorithms using a different approach.

4.1 Lower Bound for Biased Algorithms

We place Yao's general formulation in the context of our specific problem.

► **Theorem 3.** ([18]) *Consider our estimation problem over the inputs $x \in [0, 1]$, and let $\mathcal{A}$ be the set of all possible deterministic algorithms. For a (deterministic) algorithm $a \in \mathcal{A}$ and input $x \in [0, 1]$, let the function $c(a, x) = (a(x) - x)^2$ be its squared error. Then for any randomized algorithm A and input distribution $q \in \Delta([0, 1])$ such that $X \sim q$:*

$$\max_{x \in [0, 1]} \mathbb{E}[c(A, x)] \geq \min_{a \in \mathcal{A}} \mathbb{E}[c(a, X)] .$$

That is, the expected squared error (over the choice of x from distribution q) of the best deterministic algorithm (for q) lower bounds the expected squared error of any randomized (potentially biased) algorithm A for the worst-case x (i.e., its cost). Further, the inequality holds as an equality for the optimal distribution q and algorithm A , i.e.,

$$\min_A \max_{x \in [0, 1]} \mathbb{E}[c(A, x)] = \max_q \min_{a \in \mathcal{A}} \mathbb{E}[c(a, X)] .$$

We proceed by selecting distributions q to lower bound $\min_{a \in \mathcal{A}} \mathbb{E}[c(a, X)]$. Notice that a deterministic algorithm is defined using two values $v_0, v_1 \in [0, 1]$, such that if $|x - v_0| \leq |x - v_1|$ then Buffy sends 0 and Angel estimates x as v_0 . Similarly, if $|x - v_0| > |x - v_1|$ then Buffy sends 1 and the Angel estimates x as v_1 . In general, the above framework asserts that the expected cost, for the worst-case input, of any randomized algorithm is

$$\max_{q \in \Delta([0, 1])} \min_{v_0, v_1 \in [0, 1]} \int_0^1 \min \{(x - v_0)^2, (x - v_1)^2\} q(x) dx . \quad (1)$$

Our framework lower bounds the cost for any (biased or unbiased) algorithm that may use any amount of (shared or private) randomness. We now consider distributions q to lower bound the cost and later discuss the limitations of this approach.

4.1.1 The $\{0, 1/2, 1\}$ case

First, consider a discrete probability distribution q over $\{0, 1/2, 1\}$. Any deterministic algorithm cannot estimate all values exactly, and it must map at least two of the points to a single value, thus allowing us to lower bound its cost. In Appendix C, we prove the following.

► **Lemma 4.** *Any deterministic algorithm must incur a cost of at least $\frac{q(0) \cdot q(1/2)}{4(q(0) + q(1/2))}$.*

For $q(0) = q(1) = (2 - \sqrt{2})/2$ and $q(1/2) = \sqrt{2} - 1$, this lemma yields a lower bound of $3/4 - 1/\sqrt{2} \approx 0.04289$. In Section 6.1, we show that this is *an optimal* lower bound when x is known to be in $\{0, 1/2, 1\}$, by giving an algorithm with a matching cost.

4.1.2 The $[0, 1]$ Case

In the general case, where x can take on any value in $[0, 1]$, we can get a tighter bound by looking at mixed distributions. Specifically, for parameters $a, w \in [0, 1/2]$, we consider the distribution where:

$$q(x) = \begin{cases} 0 & \text{with probability } w \\ 1 & \text{with probability } w \\ \text{uniform on } [a, 1-a] & \text{otherwise} \end{cases}.$$

Directly analyzing the optimal deterministic algorithm for this distribution proves complex. Instead, we *hypothesize* that there exists an optimal algorithm for which either (1) $v_1 = 1 - v_0$ or (2) $v_1 = 1$. We first analyze what values of a, w maximize the cost of the best deterministic algorithm with the above form. Then, we verify that a lower bound for the resulting distribution (with the specific a, w values).

For case (1), we can express the cost as $2 \cdot \left(wv_0^2 + \frac{1/2-w}{1/2-a} \int_a^{1/2} (x - v_0)^2 dx \right)$. Similarly, for case (2), we get a cost of $wv_0^2 + \frac{1-2w}{1-2a} \cdot \left(\int_a^{\min\{1-a, (v_0+1)/2\}} (x - v_0)^2 dx + \int_{(v_0+1)/2}^{1-a} (x - 1)^2 dx \right)$. Therefore, the cost of the optimal algorithm from the above family is given as:

$$\min \left\{ 2 \cdot \left(wv_0^2 + \frac{1/2-w}{1/2-a} \int_a^{1/2} (x - v_0)^2 dx \right), \right. \\ \left. wv_0^2 + \frac{1-2w}{1-2a} \cdot \left(\int_a^{\min\{1-a, (v_0+1)/2\}} (x - v_0)^2 dx + \int_{(v_0+1)/2}^{1-a} (x - 1)^2 dx \right) \right\}.$$

This cost is maximized for $a = \frac{-2w^2 + w - 2\sqrt{w(1-w)} + 1}{4w^2 - 6w + 2}$, where the value of w satisfies

$$32w^3 - 56w^2 + \sqrt{w(1-w)} \cdot (8w^4 - 24w^3 + 38w^2 - 8w - 7) + 24w = 0.$$

The resulting bound is slightly larger than 0.0459. Next, we verify that for these a, w values, no deterministic algorithm can achieve a lower cost. Specifically, instead of using $q(x)$ as described above, we generate a finite discrete distribution. For a parameter $n \in \mathbb{N}$, we define:

$$q_n(x) = \begin{cases} \frac{\lfloor n \cdot w \rfloor}{n} & \text{if } x \in \{0, 1\} \\ \frac{1}{n} & \text{if } x \in \left\{ a + \frac{1-2a}{2(n-2\lfloor n \cdot w \rfloor)} + i \cdot \frac{1-2a}{n-2\lfloor n \cdot w \rfloor} \mid i \in [n - 2\lfloor n \cdot w \rfloor] \right\} \\ 0 & \text{otherwise} \end{cases}.$$

Note that $\lim_{n \rightarrow \infty} q_n(x) = q(x)$. Then, we find the optimal deterministic solution for this distribution by computing a deterministic k -means clustering for $k = 2$, e.g., using [6]. The

code that we used to obtain this result is available at [2]. The optimal deterministic algorithm for $q_n(x)$ tends to have either $v_1 = 1 - v_0$ or $v_1 = 1$ as hypothesized. Finally, for $n = 10^6$ we get a cost higher than 0.0459 which we use as a lower bound.

We do not believe that this bound is tight. Nonetheless, as we show in Section 6.2.4, our bound is within 0.2% of the optimum.

4.2 Lower Bound for Unbiased Algorithms - Beyond MiniMax

We now consider lower bounds for unbiased algorithms. Utilizing Yao's lemma does not appear to provide means to obtain sharper bounds when requiring the algorithm to be unbiased. We consider the case where $x \in \{0, 1/2, 1\}$ and directly prove that any unbiased algorithm must have a worst-case variance of at least $1/16$. This lower bound then also holds for $x \in [0, 1]$, although an improved bound of $\pi^2/64 - 1/12 \approx 0.07$ for this case, based on an average-case analysis, is presented in [7]. Assume that we have $h \in [0, 1]$. Buffy sends $X(x, h)$ to Angel, which determines an estimate $\hat{x}(X(x, h), h)$. For $x', x'' \in \{0, 1/2, 1\}$, let $p_{x', x''} = \Pr[X(x', h) = X(x'', h)]$ denote the probability (with respect to h) that the same bit is sent for x', x'' . Since we send a single bit, we have that $p_{0, 1/2} + p_{1/2, 1} + p_{0, 1} \geq 1$. For all $x', x'' \in \{0, 1/2, 1\}$, we define $H_{x', x''} = \{h \in [0, 1] : X(x', h) = X(x'', h)\}$ to be the set of shared-randomness values that would lead Buffy to send the same bit for both x and x' . Next, denote by $G_{x', x''} = \mathbb{E}[\hat{x}|h \in H_{x', x''}, x \in \{x', x''\}]$ the expected estimate value, conditioned on the shared randomness being in $H_{x', x''}$.

We have that:

$$\begin{aligned} \text{Var}[\hat{x}|x=0] &\geq p_{0, 1/2} \cdot (G_{0, 1/2} - 0)^2 + p_{0, 1} \cdot (G_{0, 1} - 0)^2 + p_{1/2, 1} \cdot (\mathbb{E}[\hat{x}|h \in H_{1/2, 1}, x=0] - 0)^2 \\ \text{Var}[\hat{x}|x=1/2] &\geq p_{0, 1/2} \cdot (G_{0, 1/2} - 1/2)^2 + p_{0, 1} \cdot (\mathbb{E}[\hat{x}|h \in H_{0, 1}, x=1/2] - 1/2)^2 \\ &\quad + p_{1/2, 1} \cdot (G_{1/2, 1} - 1/2)^2 \\ \text{Var}[\hat{x}|x=1] &\geq p_{0, 1/2} \cdot (\mathbb{E}[\hat{x}|h \in H_{0, 1/2}, x=1] - 1)^2 + p_{0, 1} \cdot (G_{0, 1} - 1)^2 + p_{1/2, 1} \cdot (G_{1/2, 1} - 1)^2. \end{aligned}$$

To proceed, we require the algorithm to be unbiased:

$$\begin{aligned} G_{0, 1/2}p_{0, 1/2} + G_{0, 1}p_{0, 1} + \mathbb{E}[\hat{x}|h \in H_{1/2, 1}, x=0]p_{1/2, 1} &= 0 \\ G_{0, 1/2}p_{0, 1/2} + \mathbb{E}[\hat{x}|h \in H_{0, 1}, x=1/2]p_{0, 1} + G_{1/2, 1}p_{1/2, 1} &= 1/2 \\ \mathbb{E}[\hat{x}|h \in H_{0, 1/2}, x=1]p_{0, 1/2} + G_{0, 1}p_{0, 1} + G_{1/2, 1}p_{1/2, 1} &= 1. \end{aligned}$$

This allows us to express the expectations $\{E[\hat{x}|h \in H_{x', x''}, x = x'''] | x', x'', x''' \in \{0, 1/2, 1\}\}$ using $p_{x', x''}, G_{x', x''}$ and obtain a set of three inequalities with six variables. Our full analysis, given in Appendix D, proceeds with a case analysis based on the value of $p_{0, 1}$, the probability that the sender would send the same bit for 0, 1. We show that there exists an optimal algorithm in which $p_{0, 1/2} + p_{1/2, 1} + p_{0, 1} = 1$, $p_{0, 1/2} = p_{1/2, 1}$, and $G_{1/2, 1} = 1 - G_{0, 1/2}$. This reduces the number of variables to three, allowing us to optimize the expression and show a lower bound of $1/16$ on any unbiased algorithm.

5 Algorithms with Unbounded Private Randomness

Here, we consider the case where the shared randomness is limited to ℓ bits, i.e., $h \in [2^\ell]$, but Buffy may use unbounded private randomness $r \sim U[0, 1]$ (that is independent of h).

We present the following algorithm: Buffy sends X to Angel, where

$$X \triangleq \begin{cases} 1 & \text{if } x \geq (r + h)2^{-\ell} \\ 0 & \text{otherwise} \end{cases}.$$

Angel then estimates

$$\hat{x} = X + (h - 0.5(2^\ell - 1)) \cdot 2^{-\ell}.$$

We first show that our protocol is unbiased. It holds that $\mathbb{E}[h] = 0.5(2^\ell - 1)$ and $(r + h) \sim U[0, 2^\ell]$ (i.e., $(r + h)2^{-\ell} \sim U[0, 1]$), and thus $\mathbb{E}[\hat{x}] = \mathbb{E}[X] = x$.

We now state theorem, whose proof appears in Appendix E, that bounds the variance:

► **Theorem 5.** $\text{Var}[\hat{x}] \leq 1/12 \cdot (1 - 4^{-\ell}) + 1/4 \cdot 4^{-\ell} = 1/6 \cdot (1/2 + 4^{-\ell})$.

In Appendix F, we describe a simple generalization of this algorithm, together with a lower bound, for sending $k > 1$ bits. We now explain the connection to subtractive dithering and explore the applicability of the algorithm for the $x \in \{0, 1/2, 1\}$ special case.

Connection to subtractive dithering: First invented for improving the visibility of quantized pictures [13], subtractive dithering aims to alleviate potential distortions that originate from quantization. Subtractive dithering was later extended for other domains such as speech [4], distributed deep learning [1], and federated learning [16].

In our setting, subtractive dithering corresponds to using shared randomness to add *noise* ς to x before applying a deterministic quantization and subtracting ς from the estimation. Specifically, let $\mathcal{Q} : [0, 1] \rightarrow \{0, 1\}$ be a two-level deterministic quantizer such that $\mathcal{Q}(g) = 1$ if $g \geq 1/2$ and 0 otherwise. Then, in subtractive dithering Buffy sends $X = \mathcal{Q}(x + \varsigma)$ and Angel estimates $\hat{x} = X - \varsigma$.

There are several noise classes that ς can be drawn from, as classified in [14], that yield $\hat{x} \sim U[x - 1/2, x + 1/2]$. For example, ς can be distributed uniformly on $[-1/2, 1/2]$.

Consider our algorithm of this section without restricting the number of random bits (i.e., $\ell \rightarrow \infty$, and rescale so $h \in U[0, 1]$). This would yield the following algorithm:

$$X \triangleq \begin{cases} 1 & \text{if } x \geq h \\ 0 & \text{otherwise} \end{cases}$$

and $\hat{x} = X + h - 0.5$. Similarly to subtractive dithering, we get that $\hat{x} \sim U[x - 1/2, x + 1/2]$, as we prove in Appendix G for completeness. To see that the two algorithms are equivalent (for $\varsigma \sim U[-1/2, 1/2]$), denote $h' = 1/2 - h$ (i.e., $h' \sim U[-1/2, 1/2]$). Then $X = 1$ if $x + h' \geq 1/2$ and $\hat{x} = X - h'$.

Therefore, we conclude that our algorithm provides a spectrum between randomized rounding ($\ell = 0$) and a form of subtractive dithering ($\ell \rightarrow \infty$). In practice, this means that a small number of shared random bits yields a variance that is close to that of subtractive dithering ($\text{Var}[\hat{x}] = 1/12$). For example, with a single shared random byte (i.e., $\ell = 8$), our algorithm has a worst-case variance that is within 0.02% of $1/12$.

The $x \in \{0, 1/2, 1\}$ case: Notice that if x is known to be in $\{0, 1/2, 1\}$, then our ($\ell = 1$) algorithm gives $\text{Var}[\hat{x}] = 1/16$, as evident from (25). Further, in this case, we do not require the private randomness as we can rewrite Buffy's algorithm as:

$$X \triangleq \begin{cases} 0 & \text{if } x = 0 \\ 1 - h & \text{if } x = 1/2 \\ 1 & \text{if } x = 1 \end{cases},$$

while Angel estimates $\hat{x} = X + (h - 0.5)/2$. This algorithm considerably improves over randomized rounding (which is optimal when no shared randomness is allowed, as shown in Appendix A), that has a variance of $1/4$ for $x = 1/2$; i.e., a single shared random bit

reduces the worst-case variance by a factor of 4. Further, it also improves over subtractive dithering, reducing the variance by a $4/3$ factor. Further, this result is optimal according to the Section 4.2 lower bound, even if unbounded shared randomness is allowed.

6 Algorithms without Private Randomness

In some cases, generating random bits may be expensive, e.g., when running on power-constrained devices. This is particularly acute when the device operates in an energy harvesting mode [19]. Past works have even considered how to “recycle” random bits (e.g., [8]). Therefore, it is important to study how to design algorithms that use just a few random bits. To address this need, we consider scenarios where Buffy and Angel have access to a shared ℓ -bit random value h , but no private randomness.

One thing to notice is that Angel can produce at most $2^{\ell+1}$ different values since Angel is deterministic after obtaining the $\ell + 1$ bits of h and X . In particular, this means there is no unbiased protocol for general $x \in [0, 1]$. Therefore, we focus on biased algorithms and study how shared randomness allows improving over deterministic rounding (which is optimal without shared randomness, as we show in Section 3.2).

We start by proposing an optimal algorithm for the case where x is known to be in $\{0, 1/2, 1\}$. Then, we present adaptations of the subtractive dithering estimation method for the biased $x \in [0, 1]$ setting. These improve over both (unbiased) subtractive dithering and deterministic rounding. To the best of our knowledge, these adaptations are novel. Next, we show how Buffy can further reduce the cost while, among other changes, using a small number of shared random bits. We conclude by giving design principles for numerically approximating the optimal algorithm and give realizations for small number of shared random bits.

6.1 The $x \in \{0, 1/2, 1\}$ case

We now consider the scenario where x is guaranteed to be in $\{0, 1/2, 1\}$ using a single shared randomness bit $h \in \{0, 1\}$. For some $\alpha \in [0, 1]$, Buffy sends

$$X = \begin{cases} 1 & \text{if } x = 1 \vee (x = 1/2 \wedge h = 0) \\ 0 & \text{otherwise} \end{cases}$$

while Angel estimates $\hat{x} = \alpha \cdot h + (1 - \alpha) \cdot X$.

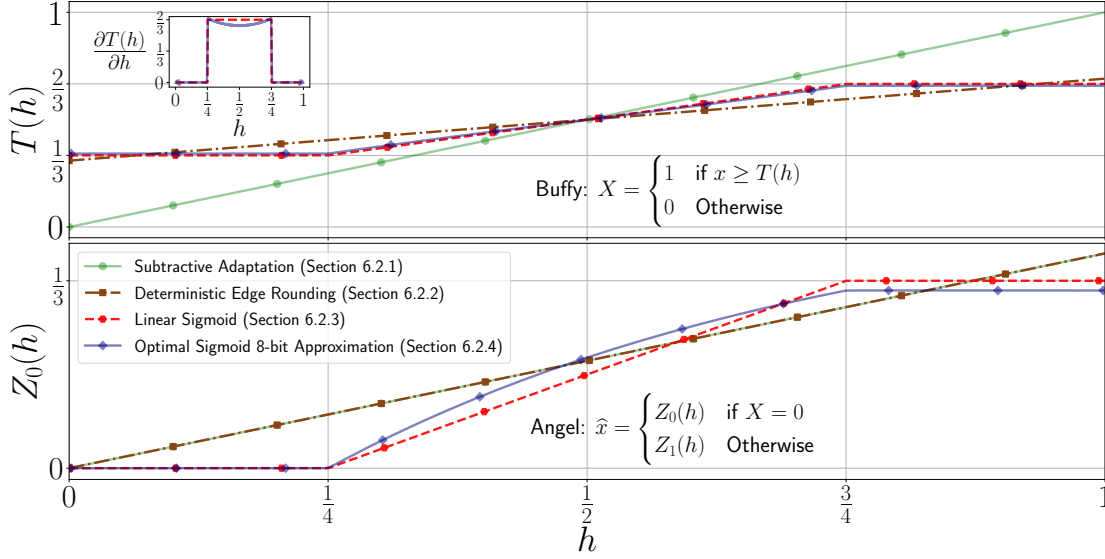
For example, this means that if $x = 0$, the squared error is 0 if $h = 0$ and α^2 otherwise. That is, the expected squared error is $\alpha^2/2$. We optimize over the α value to minimize the cost

$$\min_{\alpha \in [0, 1]} \max \left\{ \alpha^2/2, (1 - (1 - \alpha))^2/2, \mathbb{E} \left[(1/2 - (\alpha \cdot h + (1 - \alpha) \cdot (1 - h)))^2 \right] \right\}.$$

This is optimized for $\alpha = 1 - 1/\sqrt{2}$, yielding a cost of $3/4 - 1/\sqrt{2} \approx 0.04289$, which is *optimal* according to our Section 4 lower bound, even if unbounded shared randomness is allowed.

6.2 The $x \in [0, 1]$ case

An important observation regarding optimal biased algorithms is that they, without loss of generality, can be expressed as a pair of monotone increasing functions $T, Z_0 : [0, 1] \rightarrow [0, 1]$ as follows. Here T is a threshold function that determines whether 0 or 1 is sent, Z_0 is the estimator when 0 is received, and $Z_1 : [0, 1] \rightarrow [0, 1]$, given by $Z_1(h) = 1 - Z_0(1 - h)$, is the



■ **Figure 1** Illustration of the different biased algorithms.

estimator when 1 is received. That is, Buffy sends

$$X = \begin{cases} 1 & \text{if } x \geq T(h) \\ 0 & \text{otherwise} \end{cases}.$$

In turn, Angel estimates $\hat{x} = Z_X(h)$. We further explain this representation in Appendix H. Based on this observation, we next lay out a sequence of algorithmic improvements over deterministic rounding that leverage the shared randomness to reduce the cost. We visualize the algorithms resulting from each improvement in Figure 1.

6.2.1 Subtractive dithering adaptations

As subtractive dithering provides the lowest cost (albeit using unbounded shared randomness) of the previously mentioned unbiased algorithms, one may wonder if it is possible to adapt it to the biased scenario. Accordingly, we first briefly overview two natural adjustments that use unbounded shared randomness and improve over the $1/16$ cost of deterministic rounding. We then propose improved protocols that reduce the cost further despite using only a small number (e.g., $\ell = 3$) of random bits.

Intuitively, subtractive dithering may produce estimates that are outside the $[0, 1]$ range. Therefore, by *truncating* the estimates to $[0, 1]$ one may only reduce the expected squared error for any $x \neq 1/2$. However, it does not reduce the expected squared error for $x = 1/2$, and thus the cost would remain $1/12$.

To reduce the cost, one may further truncate the estimates to $[z, 1-z]$ for some $z \in [0, 1/2]$. Indeed, we show in Appendix I that this truncation reduces the cost to $\approx 0.0602 \approx 1/16.6$, for z satisfying $1/24 + z^2/2 + (2z^3)/3 = 0$ ($z \approx 0.17349$).

A better adaptation strategy is obtained by changing the estimation to a linear combination of X and h . Specifically, consider the protocol where Buffy sends (for a shared $h \sim U[0, 1]$)

$$X = \begin{cases} 1 & \text{if } x \geq h \\ 0 & \text{otherwise} \end{cases}$$

and Angel estimates, for some $\alpha \in [0, 1]$,

$$\hat{x} = \alpha \cdot h + (1 - \alpha) \cdot X.$$

Optimizing the parameters, we show in Appendix J that this algorithm achieves a cost of $5/3 - \phi \approx 0.04863$, which is obtained for $\alpha = 2 - \phi \approx 0.382$. Interestingly, this cost is achieved for *all* $x \in [0, 1]$.

In Figure 1, we illustrate this algorithm. As shown, the subtractive dithering adaption has $T(h) = h$ and $Z_0(h) = \alpha \cdot h$. This means that Buffy sends $X = 1$ if $x \geq h$ and Angel estimates $\hat{x} = \alpha \cdot h$ if $X = 0$ and $\hat{x} = (1 - \alpha) \cdot (1 - h) = (1 - \alpha) + \alpha \cdot h$ otherwise.

6.2.2 Deterministically rounding extreme values

We now show how to leverage a finite number of shared random bits ℓ to design improved algorithms. As we show, it is possible to benefit from *deterministically* rounding values that are “close” to 0 or 1 and use the shared randomness otherwise.

Similarly to the subtractive dithering adaptation above, Angel estimates x using a linear combination of h (with weight α) and X (with a weight of $1 - \alpha$), where $\alpha \in [0, 1]$ is chosen later. For all $i \in [2^\ell - 1]$, define the interval

$$\mathcal{I}_i = \left[(1 - \alpha)/2 + i \cdot \frac{\alpha}{2^\ell - 1}, (1 - \alpha)/2 + (i + 1) \cdot \frac{\alpha}{2^\ell - 1} \right). \quad (2)$$

In our algorithm, Buffy sends

$$X = \begin{cases} 0 & \text{if } x < (1 - \alpha)/2 \\ \mathbb{1}_{h \leq i} & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ 1 & \text{if } x \geq (1 + \alpha)/2 \end{cases},$$

and Angel estimates: $\hat{x} = \alpha \cdot h / (2^\ell - 1) + (1 - \alpha) \cdot X$.

Note that we deterministically partition the range $[(1 - \alpha)/2, (1 + \alpha)/2]$ into $2^\ell - 1$ equally spaced intervals. Intuitively, these intervals are chosen in a way that makes the expected squared error a continuous function of x , as our analysis, given in Appendix K, indicates.

As we show, minimizing $\text{cost} = \min_\alpha \max_x \mathbb{E}[(\hat{x} - x)^2]$ yields cumbersome expressions. For example, we get that with one shared random bit ($\ell = 1$), our algorithm has a cost of $1/18 \approx 0.05556$ (obtained for¹ $\alpha = 1/3$), lower than that of deterministic rounding (i.e., $1/16$). For $\ell = 2$, we obtain a cost of $\frac{259 - 140\sqrt{3}}{338} \approx 0.04885$ (reached for $\alpha = \frac{15 - 6\sqrt{3}}{13}$), and $\ell = 3$ bits further reduces the cost to $35/722 \approx 0.04848$ (when $\alpha = 7/19$). Additionally, with $\ell = 3$ bits, this improves over the subtractive dithering adaptations (that use unbounded shared randomness) for all $x \in [0, 1]$. Notice that these costs are $\approx 23.3\%$, $\approx 8.4\%$, and $\approx 7.6\%$ from the $\frac{5\phi - 8}{2}$ lower bound (see Section 4), and thus from the optimal algorithm. For completeness, we give the limiting algorithm (as $\ell \rightarrow \infty$) in Appendix L. For intuition, we illustrate the limiting algorithm ($h \in [0, 1]$) in Figure 1. As shown, we have $T(h) = \frac{1 - \alpha}{2} + \alpha \cdot h$ (where $\alpha = 2 - \phi \approx 0.38$) and $Z_0(h) = \alpha \cdot h$. Observe that Angel uses the same estimation function as in Section 6.2.1, but Buffy’s threshold function is different. Intuitively, the new threshold function ensures that each x is mapped to the closest estimate value. For example, if $x = 0.1$ and $h = 0$, the subtractive adaptation would have $X = 1$ and thus $\hat{x} = 1 - \alpha \approx 0.62$ while here we get $X = 0$ and $\hat{x} = 0$.

¹ Notice that it is different than the value used for the $x \in \{0, 1/2, 1\}$ case.

Interestingly, the cost slightly and monotonically *increases* when increasing the number of bits ℓ beyond 3. This phenomenon suggests that we need more complex algorithms to leverage additional available random bits. We explore several approaches; in Appendix M, we show that by probabilistically selecting between the above algorithm (for $\ell \rightarrow \infty$) and the $\{0, 1/2, 1\}$ algorithm from Section 6.1, we can reduce the error to $5/3 - \phi \approx 0.04863$. Intuitively, Buffy and Angel can implicitly agree on the chosen algorithm using the shared randomness. Here, we proceed by analyzing the potential benefits of non-uniform partitioning of the h values, which reduces the error further.

6.2.3 Non-uniform partitioning

Intuitively, the above algorithms have a threshold function that is linear in h ; i.e., it takes the form $T(h) = a \cdot h + b$. We now show that this can be improved by looking at *sigmoid*-like functions. For ease of exposition, in this section, we consider $h \in [0, 1]$ to represent unbounded shared randomness, although the algorithm can be discretized given sufficient random bits. Recall from Section 6.2 that an algorithm can be expressed as a pair of functions $T, Z_0 : [0, 1] \rightarrow [0, 1]$ such that Buffy sends 1 if $x \geq T(h)$ while Angel estimates $Z_0(h)$ when receiving $X = 0$ and $Z_1(h) = 1 - Z_0(1 - h)$ otherwise. Here, we consider a *linear sigmoid* function (also illustrated in Figure 1, which, for some $h_0 \in [0, 1/2]$, is defined as

$$T(h) = \begin{cases} \alpha & \text{if } h < h_0 \\ \alpha + (1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} & \text{if } h \in [h_0, 1 - h_0] \\ 1 - \alpha & \text{otherwise} \end{cases}$$

$$Z_0(h) = \begin{cases} 0 & \text{if } h < h_0 \\ (1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} & \text{if } h \in [h_0, 1 - h_0] \\ 1 - 2\alpha & \text{otherwise} \end{cases}$$

Notice that in this algorithm we have $Z_0(h) = T(h) - \alpha$.

Our analysis, given in Appendix N, shows that the cost is minimized for $h_0 = 1/4, \alpha = 1/3$, where the error is:

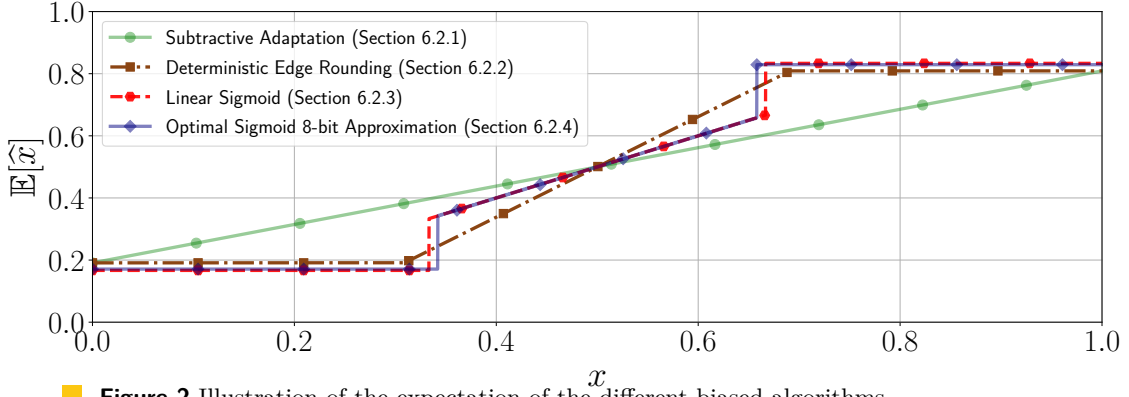
$$\mathbb{E}[(\hat{x} - x)^2] = \mathbb{E}[(\hat{x})^2] - 2x\mathbb{E}[\hat{x}] + x^2 = \begin{cases} 5/108 - x/3 + x^2 & \text{if } x < 1/3 \\ 5/108 & \text{if } x \in [1/3, 2/3] \\ 77/108 - 5x/3 + x^2 & \text{otherwise} \end{cases}$$

Therefore, the cost is $5/108 \approx 0.0463$, which is less than 0.9% higher than the 0.0459 lower bound (Section 4.1.2). The algorithm has two interesting properties. First, its expected squared error is constant for all $x \in \{0, 1\} \cup [1/3, 2/3]$ and, second, its expectation is not continuous as a function of x , as shown in Figure 2.

6.2.4 Towards the Optimal Algorithm

We now consider more general algorithms that have arbitrary estimate function Z_0 . To that end, we use a numerical solver that approximates the optimal solution. Clearly, to define the input problem, we need to limit the number of variables and constraints. We achieve this using several observations:

- We consider bounded shared randomness $h \in [2^\ell]$ for $\ell \in \mathbb{N}$ bits. In fact, bounded shared randomness is precisely what allows us to develop this numerical approach.



■ **Figure 2** Illustration of the expectation of the different biased algorithms.

- We use the observation that an optimal algorithm's T and Z_0 functions are not independent and satisfy $\forall h \in [0, 1] : T(h) = \frac{Z_0(h) + Z_1(h)}{2} = \frac{Z_0(h) + (1 - Z_0(1-h))}{2}$; this is because, that way, every $x \in [0, 1]$ is estimated using the value closer to it between $Z_0(h)$ and $Z_1(h) = 1 - Z_0(1 - h)$. In fact, the algorithms in sections 6.2.2-6.2.3 follow this rule, while the subtractive adaptation (Section 6.2.1) does not. As a result, we can define the variables $\{z_h | h \in [2^\ell]\}$ and derive the thresholds from the solver's output by $Z_0(h) = z_h$.
- For computing the maximal error for any $x \in [0, 1]$, it is enough to look at a discrete set of points. This is because the number of possible estimates is $2^{\ell+1}$. Therefore, given two estimates $z_h, 1 - z_{2^\ell-1-h}$ that correspond to the values Angel uses given h and $X = 0$ or $X = 1$, the worst expected squared error (for this h) is obtained for $y_h \triangleq \frac{z_h + z_{2^\ell-1-h}}{2}$. Therefore, by checking all $x \in \{y_h | h \in [2^\ell]\}$, we can compute the cost.

Using these observations, we formulate the input as:

$$\begin{aligned}
 & \underset{\{z_h | h \in [2^\ell]\}}{\text{minimize}} && C \\
 & \text{subject to} && C \geq \sum_{j=0}^h (y_h - (1 - z_{2^\ell-1-h}))^2 + \sum_{j=h+1}^{2^\ell-1} (y_h - z_h)^2, \quad h = 0, \dots, 2^\ell - 1 \\
 & && y_h = \frac{z_h + z_{2^\ell-1-h}}{2}, \quad z_h \in [0, 1] \quad h = 0, \dots, 2^\ell - 1
 \end{aligned}$$

In the above, we express the expected squared error at y_h by considering the h values for which $x \geq T(h)$ ($j \in [h]$) and those that $x < T(h)$. The output for the above problem does not seem to follow a compact representation. However, it is still possible to implement using a simple lookup table. For example, if $\ell = 8$, we can store all z_h when implementing Buffy and Angel. This algorithm's cost is lower than that of the linear sigmoid (that uses unbounded randomness) when using $\ell \geq 4$ bits. Specifically, using 4 shared random bits, the cost is ≈ 0.04611 , while using 8 bits, it further reduces to ≈ 0.04599 . Notice that these are less than 0.5% and 0.2% higher than the lower bound of Section 4.1.2. We note that this approach yields improvement even for a small number of shared random bits; for example, using $\ell = 1$ bit ($h \in \{0, 1\}$), we get a cost of $1/20$ for $z_0 = 0.1, z_1 = 0.3$ which is equivalent to the following algorithm:

$$X = \begin{cases} 1 & \text{if } x \geq 0.4 + 0.2h \\ 0 & \text{otherwise} \end{cases}, \quad \hat{x} = 0.1 + 0.2h + 0.6X.$$

We visualize the resulting algorithm, for $\ell = 8$, in figures 1 and 2. Notice that while the algorithm looks almost similar to our linear sigmoid, looking that the derivative $\frac{\partial T(h)}{\partial h}$ (Figure 1) shows that this optimal solution is not piece-wise linear.

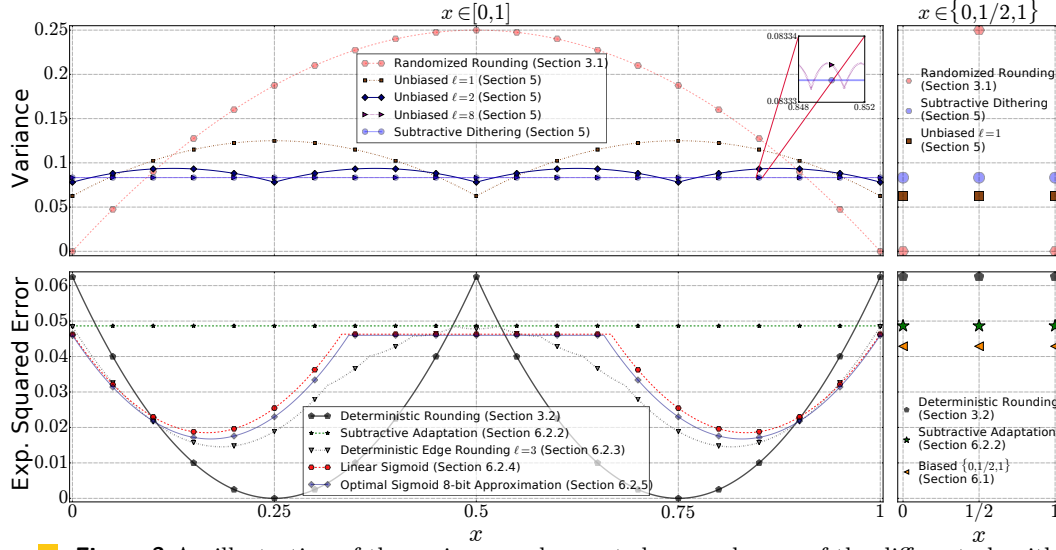


Figure 3 An illustration of the variance and expected squared error of the different algorithms.

7 Visual Comparison of the Algorithm Costs

We illustrate the various algorithms in Figure 3. In the unbiased case, notice how a single ($\ell = 1$) shared random bit significantly improves over randomized rounding (which is optimal when Buffy and Angel are restricted to private randomness). This further improves for larger ℓ values, where for $\ell = 8$ we have a cost that is only 0.02% higher than that of subtractive dithering, which uses unbounded shared randomness (the difference shown in zoom). When x is known to be in $\{0, 1/2, 1\}$ (right-hand side of the figure), it is evident how our unbiased $\ell = 1$ algorithm improves over both randomized rounding and subtractive dithering.

In the biased case, our adaptation to the subtractive dithering estimation (termed Subtractive Adaptation) improves over the cost of deterministic rounding. This is further improved by the algorithm of Section 6.2.2, termed Deterministic Edge Rounding, which is depicted using $\ell = 3$ bits as it minimizes its cost. Next, the Linear sigmoid (Section 6.2.3) shows how to lower the cost (using unbounded shared randomness) by non-uniform partitioning of the h values. Additionally, we show the optimal 8-bit algorithm (Section 6.2.4) that gets within 0.2% from the lower bound while using a single shared random byte. Finally, if x is known to be in $\{0, 1/2, 1\}$, our (optimal) biased $\{0, 1/2, 1\}$ algorithm improves over all other solutions while using only a single shared random bit.

8 Discussion

In this paper, we studied upper and lower bounds for the problem of sending a real number using a single bit. The goal is to minimize the cost, which is the worst-case variance for unbiased algorithms, or the worst-case expected squared error for biased ones. For all cases, we demonstrated how shared randomness helps to reduce the cost. Motivated by real-world applications, we derived algorithms with a bounded number of random bits that can be as low as a single shared bit. For example, in the unbiased case, using just one shared random bit reduces the variance two-fold compared to randomized rounding (which is optimal when no shared randomness is available). Further, using a single byte of shared randomness, our algorithm's variance is within 0.02% from the state of the art, which uses unbounded shared randomness. Our results are also near-optimal in the biased case, with a gap lower than 0.2% between the

upper and lower bounds with a single shared random byte. Our upper bound is presented as a sequence of algorithms, each generalizing the previous while reducing the cost further.

For the special case where x is known to be in $\{0, 1/2, 1\}$, we give optimal unbiased and biased algorithms, together with matching lower bounds. Our algorithms use a single shared random bit, and the lower bounds show that the cost cannot be improved even when unbounded shared randomness is allowed.

We conclude by identifying directions for future research, beyond settling the correct bounds. First, our lower bounds apply for algorithms that use unbounded shared randomness, and new techniques for developing sharper bounds for other cases are of interest. Another direction is looking into optimizing the cost when sending k bits, for some $k > 1$. We make a first small step in Appendix F, where we provide simple generalizations of our unbiased algorithm and lower bound to sending k bits. Finally, we are unclear on whether private randomness can help improve biased algorithms (see Table 1).

References

- 1 Afshin Abdi and Faramarz Fekri. Indirect stochastic gradient quantization and its application in distributed deep learning. In *AAAI*, 2020.
- 2 Ran Ben Basat, Gil Einziger, and Michael Mitzenmacher. Biased $[0,1]$ proof. URL: <https://gist.github.com/ranbenbasat/9959d1c70471fe870424eefbecd3e13c>.
- 3 Jeremy Bernstein, Yu-Xiang Wang, Kamyar Azizzadenesheli, and Animashree Anandkumar. signSGD: Compressed Optimisation for Non-Convex Problems. In *International Conference on Machine Learning*, pages 560–569, 2018.
- 4 MR Garey, David Johnson, and Hans Witsenhausen. The complexity of the generalized lloyd-max problem (corresp.). *IEEE Transactions on Information Theory*, 28(2):255–256, 1982.
- 5 Robert M. Gray and David L. Neuhoff. Quantization. *IEEE transactions on information theory*, 44(6):2325–2383, 1998.
- 6 Allan Grønlund, Kasper Green Larsen, Alexander Mathiasen, Jesper Sindahl Nielsen, Stefan Schneider, and Mingzhou Song. Fast exact k-means, k-medians and bregman divergence clustering in 1d. *arXiv preprint arXiv:1701.07204*, 2017.
- 7 Zarathustra Elessar Brady (<https://cstheory.stackexchange.com/users/50608/zeb>). Is subtractive dithering the optimal algorithm for sending a real number using one bit? URL: <https://cstheory.stackexchange.com/questions/48281>.
- 8 Russell Impagliazzo and David Zuckerman. How to recycle random bits. In *FOCS*, volume 30, pages 248–253, 1989.
- 9 Peter Kairouz, H. Brendan McMahan, Brendan Avent, Aurélien Bellet, Mehdi Bennis, Arjun Nitin Bhagoji, Keith Bonawitz, Zachary Charles, Graham Cormode, Rachel Cummings, Rafael G. L. D’Oliveira, Salim El Rouayheb, David Evans, Josh Gardner, Zachary Garrett, Adrià Gascón, Badi Ghazi, Phillip B. Gibbons, Marco Gruteser, Zaid Harchaoui, Chaoyang He, Lie He, Zhouyuan Huo, Ben Hutchinson, Justin Hsu, Martin Jaggi, Tara Javidi, Gauri Joshi, Mikhail Khodak, Jakub Konečný, Aleksandra Korolova, Farinaz Koushanfar, Sanmi Koyejo, Tancrede Lepoint, Yang Liu, Prateek Mittal, Mehryar Mohri, Richard Nock, Ayfer Özgür, Rasmus Pagh, Mariana Raykova, Hang Qi, Daniel Ramage, Ramesh Raskar, Dawn Song, Weikang Song, Sebastian U. Stich, Ziteng Sun, Ananda Theertha Suresh, Florian Tramèr, Praneeth Vepakomma, Jianyu Wang, Li Xiong, Zheng Xu, Qiang Yang, Felix X. Yu, Han Yu, and Sen Zhao. Advances and Open Problems in Federated Learning, 2019. [arXiv:1912.04977](https://arxiv.org/abs/1912.04977).
- 10 Sai Praneeth Karimireddy, Quentin Rebjock, Sebastian Stich, and Martin Jaggi. Error Feedback Fixes SignSGD and other Gradient Compression Schemes. In *International Conference on Machine Learning*, pages 3252–3261, 2019.
- 11 Jakub Konečný, Brendan McMahan, and Daniel Ramage. Federated optimization: Distributed optimization beyond the datacenter. *arXiv preprint arXiv:1511.03575*, 2015.
- 12 Ilan Newman. Private vs. common random bits in communication complexity. *Information processing letters*, 39(2):67–71, 1991.
- 13 Lawrence Roberts. Picture Coding Using Pseudo-Random Noise. *IRE Transactions on Information Theory*, 8(2):145–154, 1962.
- 14 Leonard Schuchman. Dither signals and their effect on quantization noise. *IEEE Transactions on Communication Technology*, 12(4):162–165, 1964.
- 15 Frank Seide, Hao Fu, Jasha Droppo, Gang Li, and Dong Yu. 1-Bit Stochastic Gradient Descent and its Application to Data-Parallel Distributed Training of Speech DNNs. In *Fifteenth Annual Conference of the International Speech Communication Association*, 2014.
- 16 Nir Shlezinger, Mingzhe Chen, Yonina C Eldar, H Vincent Poor, and Shuguang Cui. Uveqfed: Universal vector quantization for federated learning. *arXiv preprint arXiv:2006.03262*, 2020.
- 17 Wei Wen, Cong Xu, Feng Yan, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. Terngrad: Ternary Gradients to Reduce Communication in Distributed Deep Learning. In *Advances in neural information processing systems*, pages 1509–1519, 2017.

- 18 Andrew Chi-Chin Yao. Probabilistic Computations: Toward a Unified Measure of Complexity. In *IEEE FOCS*, 1977.
- 19 Pengyu Zhang and Deepak Ganesan. Enabling Bit-by-Bit Backscatter Communication in Severe Energy Harvesting Environments. In *11th {USENIX} Symposium on Networked Systems Design and Implementation ({NSDI} 14)*, pages 345–357, 2014.

A Optimality of randomized rounding

We show that without shared randomness, randomized rounding is optimal in the sense that it minimizes the worst-case estimation variance.

Consider an arbitrary protocol. We model it as follows: we have two (deterministic) parameters: $Y : [0, 1] \rightarrow [0, 1]$ and $\Gamma : \{0, 1\} \rightarrow \Delta([0, 1])$.

Buffy computes $p = Y(x)$ and sends $X \sim \text{Bernoulli}(p)$. In turn, Angel receives X , and estimates x by drawing from the distribution $\Gamma(X)$. We also denote by $Z_0 \sim \Gamma(0)$ and $Z_1 \sim \Gamma(1)$ random variables such that the final estimate is

$$\hat{x} = \mathbb{1}_{X=0} \cdot Z_0 + \mathbb{1}_{X=1} \cdot Z_1.$$

Notice that this formulation captures any protocol. For example, randomized rounding is defined as $Y(x) = x$ and

$$\Gamma(X)(y) = \begin{cases} 1 & \text{if } y = X \\ 0 & \text{otherwise} \end{cases}.$$

(this is a slight abuse of notation as the above definition assumes that $\Gamma(X)$ is a density function).

We demand that the protocol will produce unbiased estimates for any x . That is, it must satisfy:

$$\mathbb{E}[\hat{x}] = Y(x) \cdot \mathbb{E}[Z_1] + (1 - Y(x)) \cdot \mathbb{E}[Z_0] = x. \quad (3)$$

In particular, for $x = 0$, we have:

$$Y(0) \cdot \mathbb{E}[Z_1] + (1 - Y(0)) \cdot \mathbb{E}[Z_0] = 0.$$

and equivalently:

$$\mathbb{E}[Z_0] = -\frac{Y(0)}{1 - Y(0)} \cdot \mathbb{E}[Z_1]. \quad (4)$$

Similarly, plugging $x = 1$ into (3) gives:

$$Y(1) \cdot \mathbb{E}[Z_1] + (1 - Y(1)) \cdot \mathbb{E}[Z_0] = 1.$$

Using (4), we proceed with several simplifications:

$$\begin{aligned} Y(1) \cdot \mathbb{E}[Z_1] + (1 - Y(1)) \cdot -\frac{Y(0)}{1 - Y(0)} \cdot \mathbb{E}[Z_1] &= 1. \\ \mathbb{E}[Z_1] \cdot \left(Y(1) - (1 - Y(1)) \frac{Y(0)}{1 - Y(0)} \right) &= 1. \\ \mathbb{E}[Z_1] \cdot \left(\frac{Y(1) \cdot (1 - Y(0)) - (1 - Y(1))Y(0)}{1 - Y(0)} \right) &= 1. \\ \mathbb{E}[Z_1] \cdot \left(\frac{Y(1) - Y(0)}{1 - Y(0)} \right) &= 1. \end{aligned}$$

$$\mathbb{E}[Z_1] = \frac{1 - Y(0)}{Y(1) - Y(0)}. \quad (5)$$

Plugging (5) into (4), we also get:

$$\mathbb{E}[Z_0] = -\frac{Y(0)}{1 - Y(0)} \cdot \mathbb{E}[Z_1] = -\frac{Y(0)}{Y(1) - Y(0)}. \quad (6)$$

Substituting (5) and (6) in (3), we simplify the expression further:

$$\begin{aligned} Y(x) \cdot \mathbb{E}[Z_1] + (1 - Y(x)) \cdot \mathbb{E}[Z_0] &= x. \\ Y(x) \cdot \frac{1 - Y(0)}{Y(1) - Y(0)} + (1 - Y(x)) \cdot -\frac{Y(0)}{Y(1) - Y(0)} &= x. \\ \frac{Y(x) \cdot (1 - Y(0)) - (1 - Y(x)) \cdot Y(0)}{Y(1) - Y(0)} &= x. \\ \frac{Y(x) - Y(0)}{Y(1) - Y(0)} &= x. \\ Y(x) &= x \cdot (Y(1) - Y(0)) + Y(0). \\ Y(x) &= x \cdot Y(1) + (1 - x) \cdot Y(0). \end{aligned} \quad (7)$$

That is, we have that the probability to send $X = 1$ must be linear in x . We analyze the variance that results for $x = 0.5$.

$$\text{Var}[\hat{x}|x = 0.5] = \mathbb{E}[(\hat{x} - 0.5)^2 | x = 0.5] = \mathbb{E}[(\hat{x})^2 | x = 0.5] - \mathbb{E}[\hat{x} | x = 0.5] + 0.25.$$

Since $\hat{x}$ is unbiased, $\mathbb{E}[\hat{x} | x = 0.5] = 0.5$ and we get

$$\text{Var}[\hat{x}|x = 0.5] = \mathbb{E}[(\hat{x})^2 | x = 0.5] - 0.25. \quad (8)$$

Next, we analyze $\mathbb{E}[(\hat{x})^2 | x = 0.5]$:

$$\begin{aligned} \mathbb{E}[(\hat{x})^2 | x = 0.5] &= \mathbb{E}[(\mathbb{1}_{X=0} \cdot Z_0 + \mathbb{1}_{X=1} \cdot Z_1)^2 | x = 0.5] \\ &= \mathbb{E}[(Z_0)^2 \cdot \mathbb{1}_{X=0} | x = 0.5] + \mathbb{E}[(Z_1)^2 \cdot \mathbb{1}_{X=1} | x = 0.5]. \end{aligned}$$

We have that Z_0, Z_1 are independent of $\mathbb{1}_{X=0}, \mathbb{1}_{X=1}$ and of x , and thus

$$\begin{aligned} \mathbb{E}[(\hat{x})^2 | x = 0.5] &= \mathbb{E}[(Z_0)^2 | x = 0.5] \cdot (1 - Y(0.5)) + \mathbb{E}[(Z_1)^2 | x = 0.5] \cdot Y(0.5) \\ &= \mathbb{E}[(Z_0)^2] \cdot (1 - Y(0.5)) + \mathbb{E}[(Z_1)^2] \cdot Y(0.5) \\ &\geq (\mathbb{E}[Z_0])^2 \cdot (1 - Y(0.5)) + (\mathbb{E}[Z_1])^2 \cdot Y(0.5). \end{aligned} \quad (9)$$

Using (5) and (6), we have:

$$\begin{aligned} \mathbb{E}[(\hat{x})^2 | x = 0.5] &\geq \left(\frac{Y(0)}{Y(1) - Y(0)} \right)^2 \cdot (1 - Y(0.5)) + \left(\frac{1 - Y(0)}{Y(1) - Y(0)} \right)^2 \cdot Y(0.5) \\ &= \frac{(Y(0))^2 \cdot (1 - Y(0.5)) + (1 - Y(0))^2 \cdot Y(0.5)}{(Y(1) - Y(0))^2} \\ &= \frac{(Y(0))^2 + Y(0.5) - 2Y(0)Y(0.5)}{(Y(1) - Y(0))^2}. \end{aligned}$$

We now use (7) for $x = 0.5$ and get $Y(0.5) = 0.5 \cdot (Y(0) + Y(1))$, which means:

$$\mathbb{E}[(\hat{x})^2 | x = 0.5] \geq \frac{(Y(0))^2 + Y(0.5) - 2Y(0)Y(0.5)}{(Y(1) - Y(0))^2}$$

$$\begin{aligned}
&= \frac{(Y(0))^2 + 0.5 \cdot (Y(0) + Y(1)) - 2Y(0) \cdot 0.5 \cdot (Y(0) + Y(1))}{(Y(1) - Y(0))^2} \\
&= \frac{0.5 \cdot (Y(0) + Y(1)) - Y(0) \cdot Y(1)}{(Y(1) - Y(0))^2}.
\end{aligned}$$

Combined with (8), this gives:

$$\text{Var}[\hat{x}|x = 0.5] = \mathbb{E}[(\hat{x})^2|x = 0.5] - 0.25 \quad (10)$$

$$\begin{aligned}
&\geq \frac{0.5 \cdot (Y(0) + Y(1)) - Y(0) \cdot Y(1)}{(Y(1) - Y(0))^2} - 0.25 \\
&= \frac{0.5 \cdot (Y(0) + Y(1)) - Y(0) \cdot Y(1) - 0.25 (Y(1) - Y(0))^2}{(Y(1) - Y(0))^2} \\
&= \frac{0.5 \cdot (Y(0) + Y(1)) - Y(0) \cdot Y(1) - 0.25 (Y(1))^2 + 0.5Y(0)Y(1) - 0.25 (Y(0))^2}{(Y(1) - Y(0))^2} \\
&= \frac{0.5 \cdot (Y(0) + Y(1)) - 0.5Y(0) \cdot Y(1) - 0.25 (Y(1))^2 - 0.25 (Y(0))^2}{(Y(1) - Y(0))^2} \\
&= \frac{0.5 \cdot (Y(0) + Y(1)) - (0.5 \cdot (Y(0) + Y(1)))^2}{(Y(1) - Y(0))^2}. \quad (11)
\end{aligned}$$

Over the domain $Y(0), Y(1) \in [0, 1]$, (11) has two minima: $Y(0) = 0, Y(1) = 1$ and $Y(0) = 1, Y(1) = 0$. Indeed, the first corresponds to randomized rounding, while the second is using a simple transform that negates the randomized rounding's bit.

To conclude, we established that randomized rounding has a minimal worst-case variance. As a side note, by deterministically estimating $\hat{x} = X$, Inequality (9) holds as an equality and the variance is exactly 0.25.

B Optimality of deterministic rounding

We show that without shared randomness, deterministic rounding is an optimal biased solution. Notice that, in such a case, any protocol is defined by the probability of sending 1, denoted $Y(x)$, and the reconstruction distributions $V_0, V_1 \in \Delta([0, 1])$.

Let us examine $\mathbb{E}[V_0]$ and $\mathbb{E}[V_1]$. We assume, without loss of generality, that $\mathbb{E}[V_0] \leq \mathbb{E}[V_1]$.

We have that:

$$\mathbb{E}[\hat{x}] = Y(x)\mathbb{E}[V_1] + (1 - Y(x))\mathbb{E}[V_0].$$

That is, we have that *for any* $x \in [0, 1]$: $\mathbb{E}[V_0] \leq \mathbb{E}[\hat{x}] \leq \mathbb{E}[V_1]$. Next, we have that our cost $\mathbb{E}[(\hat{x} - x)^2]$ is bounded as

$$\mathbb{E}[(\hat{x} - x)^2] \geq (\mathbb{E}[(\hat{x} - x)])^2.$$

In particular, for $x = 0$, we get that

$$\mathbb{E}[(\hat{x} - x)^2|x = 0] \geq (\mathbb{E}[\hat{x}|x = 0])^2 \geq (\mathbb{E}[V_0])^2.$$

Similarly, for $x = 1$, we have

$$\mathbb{E}[(\hat{x} - x)^2|x = 1] \geq (\mathbb{E}[(\hat{x})|x = 1] - 1)^2 \geq (1 - \mathbb{E}[V_1])^2.$$

Notice that if $\mathbb{E}[V_0] \geq 0.25$ then $\mathbb{E}[(\hat{x} - x)^2 | x = 0] \geq 1/16$, and similarly, if $\mathbb{E}[V_1] \leq 0.75$ then $\mathbb{E}[(\hat{x} - x)^2 | x = 1] \geq 1/16$. Assume to the contrary that there exists an algorithm with a worst-case expected squared error lower than $1/16$, then we have $\mathbb{E}[V_0] \leq 0.25$ and $\mathbb{E}[V_1] \geq 0.75$. However, we have that $x = 0.5$ gives:

$$\begin{aligned} \mathbb{E}[(\hat{x} - x)^2 | x = 0.5] &= \mathbb{E}[\hat{x}^2 | x = 0.5] - 2x\mathbb{E}[\hat{x} | x = 0.5] + 0.25 \\ &= Y(0.5)\mathbb{E}[V_1^2] + (1 - Y(0.5))\mathbb{E}[V_0^2] - (Y(0.5)\mathbb{E}[V_1] + (1 - Y(0.5))\mathbb{E}[V_0]) + 0.25 \\ &\geq Y(0.5)(\mathbb{E}[V_1])^2 + (1 - Y(0.5))(\mathbb{E}[V_0])^2 - (Y(0.5)\mathbb{E}[V_1] + (1 - Y(0.5))\mathbb{E}[V_0]) + 0.25 \\ &= Y(0.5) \cdot \mathbb{E}[V_1] \cdot (\mathbb{E}[V_1] - 1) + (1 - Y(0.5)) \cdot \mathbb{E}[V_0] \cdot (\mathbb{E}[V_0] - 1) + 0.25 \\ &\geq Y(0.5) \cdot 0.75 \cdot (-0.25) + (1 - Y(0.5)) \cdot 0.25 \cdot (-0.75) + 0.25 = 0.25 - 3/16 = 1/16. \end{aligned}$$

In the first inequality, we used that fact that for any random variable V : $\mathbb{E}[V^2] \geq (\mathbb{E}[V])^2$, and in the second we used $\mathbb{E}[V_0] \leq 0.25$ and $\mathbb{E}[V_1] \geq 0.75$. This concludes the proof and establishes the optimality of deterministic rounding when no shared randomness is used.

C Proof of the biased $\{0, 1/2, 1\}$ lower bound

We recall Lemma 4:

► **Lemma 4.** *Any deterministic algorithm must incur a cost of at least $\frac{q(0) \cdot q(1/2)}{4(q(0) + q(1/2))}$.*

Proof. We denote by X_0 the set of values in $\{0, 1/2, 1\}$ that are closer to v_0 than to v_1 . We assume without loss of generality that $v_0 \leq v_1$ and $q(0) \leq q(1)$ and prove that an optimal algorithm would set $v_0 = \frac{q(1/2)}{2(q(0) + q(1/2))}$, $v_1 = 1$, which incurs a cost of $\frac{q(0) \cdot q(1/2)}{4(q(0) + q(1/2))}$. Indeed, for this choice of v_0, v_1 we have that $X_0 = \{0, 1/2\}$, and we get a cost of

$$\begin{aligned} q(0) \left(\frac{q(1/2)}{2(q(0) + q(1/2))} \right)^2 + q(1/2) \left(\frac{1}{2} - \frac{q(1/2)}{2(q(0) + q(1/2))} \right)^2 &= q(0) \left(\frac{q(1/2)}{2(q(0) + q(1/2))} \right)^2 \\ + q(1/2) \left(\frac{q(0)}{2(q(0) + q(1/2))} \right)^2 &= \frac{q(0)q(1/2)^2 + q(1/2)q(0)^2}{4(q(0) + q(1/2))^2} = \frac{q(0) \cdot q(1/2)}{4(q(0) + q(1/2))}. \end{aligned}$$

We now bound the performance of the optimal algorithm. We first notice that an optimal algorithm should have $0 \in X_0$ and $1 \notin X_0$. Next, notice that v_0 should be at most $1/2$ and v_1 should be at least $1/2$. Otherwise, one can improve the error for $x = 0$ or $x = 1$, respectively, without increasing the error at $1/2$. Further, observe that an optimal algorithm must have $v_0 = 0$ or $v_1 = 1$. That is because if $1/2 \in X_0$, we can reduce the error for $x = 1$ by setting $v_1 = 1$. Similarly, when $1/2 \notin X_0$, choosing $v_0 = 0$ decreases the error for $x = 0$. Now, we claim that there exists an optimal algorithm for which $v_1 = 1$. Consider some solution, and set $v'_0 = 1 - v_1$ and $v'_1 = 1$. This does not affect the error of $x = 1/2$, and does not increase the cost as $q(0) \leq q(1)$. We are left with choosing v_0 ; let us denote by $c(v_0) = q(0)v_0^2 + q(1/2)(1/2 - v_0)^2$ the resulting cost. This function has a minimum at $v_0 = \frac{q(1/2)}{2(q(0) + q(1/2))}$, which gives a cost of $\frac{q(0) \cdot q(1/2)}{4(q(0) + q(1/2))}$. ◀

This cost is maximized for $q(1/2) = \sqrt{2} - 1$ and $q(0) = q(1) = \frac{2 - \sqrt{2}}{2}$, giving a lower bound of $3/4 - 1/\sqrt{2} \approx 0.04289$. In fact, one can verify that this is the best attainable lower bound for *any* discrete distribution on three points. Further, in Section 6.1, we show that this is an *optimal* lower bound when x is known to be in $\{0, 1/2, 1\}$, by giving an algorithm with a matching cost.

D An optimal lower bound for the unbiased $\{0, 1/2, 1\}$ case

Assume that we have $h \in [0, 1]$. Buffy sends $X(x, h)$ to Angel which estimates $\hat{x}(X(x, h), h)$. For $x', x'' \in \{0, 1/2, 1\}$, let $p_{x', x''} = \Pr[X(x', h) = X(x'', h)]$ denote the probability (with respect to h) that the same bit is sent for x', x'' . Since we send a single bit, we have that $p_{0,1/2} + p_{1/2,1} + p_{0,1} \geq 1$. Further, any algorithm for which $p_{0,1/2} + p_{1/2,1} + p_{0,1} > 1$ can be transformed to having $p_{0,1/2} + p_{1/2,1} + p_{0,1} = 1$. For example, assume without loss of generality that for some $h \in [0, 1]$, $X(0, h) = X(1/2, h) = X(1, h)$. In this case, making the following modification still yields an algorithm with identical estimates: $X(1, h) = 1 - X(0, h)$ and $\hat{x}(X(1, h), h) = \hat{x}(X(0, h), h)$. Therefore, we can assume that:

$$p_{0,1/2} + p_{1/2,1} + p_{0,1} = 1.$$

For all $x', x'' \in \{0, 1/2, 1\}$, we define by $H_{x', x''} = \{h \in [0, 1] : X(x', h) = X(x'', h)\}$ the set of shared-randomness values that would lead Buffy to send the same bit for both x and x' . Next, denote by $G_{x', x''} = \mathbb{E}[\hat{x}|h \in H_{x', x''}, x \in \{x', x''\}]$ the expected estimate value, conditioned on the shared randomness being in $H_{x', x''}$.

We have that:

$$\begin{aligned} \text{Var}[\hat{x}|x=0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 0)^2 + p_{1/2,1} \cdot (\mathbb{E}[\hat{x}|h \in H_{1/2,1}, x=0] - 0)^2 \end{aligned} \quad (12)$$

$$\begin{aligned} \text{Var}[\hat{x}|x=1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 \\ &\quad + p_{0,1} \cdot (\mathbb{E}[\hat{x}|h \in H_{0,1}, x=1/2] - 1/2)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1/2)^2 \end{aligned} \quad (13)$$

$$\begin{aligned} \text{Var}[\hat{x}|x=1] &\geq p_{0,1/2} \cdot (\mathbb{E}[\hat{x}|h \in H_{0,1/2}, x=1] - 1)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 1)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1)^2. \end{aligned} \quad (14)$$

Due to unbiasedness, we must have

$$G_{0,1/2}p_{0,1/2} + G_{0,1}p_{0,1} + \mathbb{E}[\hat{x}|h \in H_{1/2,1}, x=0]p_{1/2,1} = 0 \quad (15)$$

$$G_{0,1/2}p_{0,1/2} + \mathbb{E}[\hat{x}|h \in H_{0,1}, x=1/2]p_{0,1} + G_{1/2,1}p_{1/2,1} = 1/2 \quad (16)$$

$$\mathbb{E}[\hat{x}|h \in H_{0,1/2}, x=1]p_{0,1/2} + G_{0,1}p_{0,1} + G_{1/2,1}p_{1/2,1} = 1. \quad (17)$$

We proceed with a case analysis based on the $p_{0,1}$, the probability that the sender would send the same bit for 0, 1.

D.1 Case $p_{0,1} = 0$

We start with the simpler case where the sender never sends the same bit for 0, 1 (and thus $p_{1/2,1} = 1 - p_{0,1/2}$). Then (12)-(14) yield:

$$\begin{aligned} \text{Var}[\hat{x}|x=0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 + (1 - p_{0,1/2}) \cdot (\mathbb{E}[\hat{x}|h \in H_{1/2,1}, x=0] - 0)^2 \\ \text{Var}[\hat{x}|x=1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 + (1 - p_{0,1/2}) \cdot (G_{1/2,1} - 1/2)^2 \\ \text{Var}[\hat{x}|x=1] &\geq p_{0,1/2} \cdot (\mathbb{E}[\hat{x}|h \in H_{0,1/2}, x=1] - 1)^2 + (1 - p_{0,1/2}) \cdot (G_{1/2,1} - 1)^2. \end{aligned}$$

Similarly, using (15)-(17) we get:

$$G_{0,1/2}p_{0,1/2} + \mathbb{E}[\hat{x}|h \in H_{1/2,1}, x=0](1 - p_{0,1/2}) = 0$$

$$\begin{aligned} G_{0,1/2}p_{0,1/2} + G_{1/2,1}(1 - p_{0,1/2}) &= 1/2 \\ \mathbb{E}[\hat{x}|h \in H_{0,1/2}, x = 1]p_{0,1/2} + G_{1/2,1}(1 - p_{0,1/2}) &= 1. \end{aligned} \quad (18)$$

This gives

$$\begin{aligned} \mathbb{E}[\hat{x}|h \in H_{1/2,1}, x = 0] &= \frac{0 - G_{0,1/2}p_{0,1/2}}{1 - p_{0,1/2}} \\ \mathbb{E}[\hat{x}|h \in H_{0,1/2}, x = 1] &= \frac{1 - G_{1/2,1}(1 - p_{0,1/2})}{p_{0,1/2}}, \end{aligned}$$

and thus:

$$\begin{aligned} \text{Var}[\hat{x}|x = 0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 + (1 - p_{0,1/2}) \cdot \left(\frac{0 - G_{0,1/2}p_{0,1/2}}{1 - p_{0,1/2}} - 0 \right)^2 \\ \text{Var}[\hat{x}|x = 1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 + (1 - p_{0,1/2}) \cdot (G_{1/2,1} - 1/2)^2 \\ \text{Var}[\hat{x}|x = 1] &\geq p_{0,1/2} \cdot \left(\frac{1 - G_{1/2,1}(1 - p_{0,1/2})}{p_{0,1/2}} - 1 \right)^2 + (1 - p_{0,1/2}) \cdot (G_{1/2,1} - 1)^2. \end{aligned} \quad (19)$$

Equation (18) gives $G_{1/2,1} = \frac{1/2 - G_{0,1/2}p_{0,1/2}}{1 - p_{0,1/2}}$ and therefore:

$$\text{Var}[\hat{x}|x = 1/2] \geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 + (1 - p_{0,1/2}) \cdot \left(\frac{1/2 - G_{0,1/2}p_{0,1/2}}{1 - p_{0,1/2}} - 1/2 \right)^2 \quad (20)$$

$$\text{Var}[\hat{x}|x = 1] \geq p_{0,1/2} \cdot \left(\frac{1 - (1/2 - G_{0,1/2}p_{0,1/2})}{p_{0,1/2}} - 1 \right)^2 + (1 - p_{0,1/2}) \cdot \left(\frac{1/2 - G_{0,1/2}p_{0,1/2}}{1 - p_{0,1/2}} - 1 \right)^2. \quad (21)$$

Minimizing $\max \{(19), (20), (21)\}$, we get a bound of $1/16$, obtained for $p_{0,1/2} = 1/2$ and $G_{0,1/2} = 1/4$.

D.2 Case $p_{0,1/2}, p_{1/2,1}, p_{0,1} > 0$

We proceed with the case where $x = 0$ and $x = 1$ may result in the same bit being sent, for some h values. For convenience, we restate equations (12)-(14):

$$\begin{aligned} \text{Var}[\hat{x}|x = 0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 0)^2 + p_{1/2,1} \cdot (\mathbb{E}[\hat{x}|h \in H_{1/2,1}, x = 0] - 0)^2 \\ \text{Var}[\hat{x}|x = 1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 \\ &\quad + p_{0,1} \cdot (\mathbb{E}[\hat{x}|h \in H_{0,1}, x = 1/2] - 1/2)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1/2)^2 \\ \text{Var}[\hat{x}|x = 1] &\geq p_{0,1/2} \cdot (\mathbb{E}[\hat{x}|h \in H_{0,1/2}, x = 1] - 1)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 1)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1)^2 \end{aligned}$$

and (15)-(17):

$$\begin{aligned} G_{0,1/2}p_{0,1/2} + G_{0,1}p_{0,1} + \mathbb{E}[\hat{x}|h \in H_{1/2,1}, x = 0]p_{1/2,1} &= 0 \\ G_{0,1/2}p_{0,1/2} + \mathbb{E}[\hat{x}|h \in H_{0,1}, x = 1/2]p_{0,1} + G_{1/2,1}p_{1/2,1} &= 1/2 \\ \mathbb{E}[\hat{x}|h \in H_{0,1/2}, x = 1]p_{0,1/2} + G_{0,1}p_{0,1} + G_{1/2,1}p_{1/2,1} &= 1. \end{aligned}$$

This gives:

$$\mathbb{E}[\hat{x}|h \in H_{1/2,1}, x = 0] = \frac{0 - G_{0,1/2}p_{0,1/2} - G_{0,1}p_{0,1}}{p_{1/2,1}}.$$

$$\mathbb{E}[\hat{x}|h \in H_{0,1}, x = 1/2] = \frac{1/2 - G_{0,1/2}p_{0,1/2} - G_{1/2,1}p_{1/2,1}}{p_{0,1}}$$

$$\mathbb{E}[\hat{x}|h \in H_{0,1/2}, x = 1] = \frac{1 - G_{0,1}p_{0,1} - G_{1/2,1}p_{1/2,1}}{p_{0,1/2}}.$$

Plugging these into (12)-(14) we have:

$$\begin{aligned} \text{Var}[\hat{x}|x = 0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 0)^2 + p_{1/2,1} \cdot \left(\frac{0 - G_{0,1/2}p_{0,1/2} - G_{0,1}p_{0,1}}{p_{1/2,1}} - 0 \right)^2 \\ \text{Var}[\hat{x}|x = 1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 \\ &\quad + p_{0,1} \cdot \left(\frac{1/2 - G_{0,1/2}p_{0,1/2} - G_{1/2,1}p_{1/2,1}}{p_{0,1}} - 1/2 \right)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1/2)^2 \\ \text{Var}[\hat{x}|x = 1] &\geq p_{0,1/2} \cdot \left(\frac{1 - G_{0,1}p_{0,1} - G_{1/2,1}p_{1/2,1}}{p_{0,1/2}} - 1 \right)^2 \\ &\quad + p_{0,1} \cdot (G_{0,1} - 1)^2 + p_{1/2,1} \cdot (G_{1/2,1} - 1)^2 \end{aligned}$$

We use the following lemma, whose proof appear in D.3.

► **Lemma 6.** *There exists an optimal unbiased solution for which $p_{0,1/2} = p_{1/2,1}$ and $G_{1/2,1} = 1 - G_{0,1/2}$.*

The lemma implies that $p_{0,1} = 1 - p_{0,1/2} - p_{1/2,1} = 1 - 2p_{0,1/2}$. Therefore, we get

$$\begin{aligned} \text{Var}[\hat{x}|x = 0] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 0)^2 \\ &\quad + (1 - 2p_{0,1/2}) \cdot (G_{0,1} - 0)^2 \\ &\quad + p_{0,1/2} \cdot \left(\frac{0 - G_{0,1/2}p_{0,1/2} - G_{0,1}(1 - 2p_{0,1/2})}{p_{0,1/2}} - 0 \right)^2 \end{aligned} \quad (22)$$

$$\begin{aligned} \text{Var}[\hat{x}|x = 1/2] &\geq p_{0,1/2} \cdot (G_{0,1/2} - 1/2)^2 \\ &\quad + (1 - 2p_{0,1/2}) \cdot \left(\frac{1/2 - G_{0,1/2}p_{0,1/2} - (1 - G_{0,1/2})p_{0,1/2}}{(1 - 2p_{0,1/2})} - 1/2 \right)^2 \\ &\quad + p_{0,1/2} \cdot (1/2 - G_{0,1/2})^2 \end{aligned} \quad (23)$$

$$\begin{aligned} \text{Var}[\hat{x}|x = 1] &\geq p_{0,1/2} \cdot \left(\frac{1 - G_{0,1}(1 - 2p_{0,1/2}) - (1 - G_{0,1/2})p_{0,1/2}}{p_{0,1/2}} - 1 \right)^2 \\ &\quad + (1 - 2p_{0,1/2}) \cdot (G_{0,1} - 1)^2 + p_{0,1/2} \cdot (-G_{0,1/2})^2. \end{aligned} \quad (24)$$

The infimum of $\max\{(22),(23),(24)\}$, over all possible $p_{0,1/2}, G_{0,1/2}, G_{0,1}$ values, we get a lower bound of $1/16$, which is obtained for $p_{0,1/2} \rightarrow 1/2, G_{0,1/2} = 1/4$ (i.e., $p_{0,1} \rightarrow 0$).

D.3 Proof of Lemma 6

Assume an optimal unbiased algorithm, defined using $X^*(x, h)$, $\hat{x}^*(X, h)$, the probabilities $p_{0,1/2}^*, p_{0,1}^*, p_{1/2,1}^*$, and $G_{0,1/2}^*, G_{0,1}^*, G_{1/2,1}^*$. We consider the following algorithm; let h_{alg} be a shared random bit that is independent of h . If $h_{\text{alg}} = 0$ Buffy sends $X = X^*(x, h)$ and otherwise (if $h_{\text{alg}} = 1$) $X = X^*(1 - x, h)$. In turn, Angel estimates $\hat{x} = \hat{x}^*(X, h)$ if $h_{\text{alg}} = 0$ or $\hat{x} = 1 - \hat{x}^*(X, h)$ otherwise.

Notice that our algorithm is unbiased:

$$\begin{aligned}\mathbb{E}[\hat{x}] &= 1/2(\mathbb{E}[\hat{x}|h_{\text{alg}} = 0] + \mathbb{E}[\hat{x}|h_{\text{alg}} = 1]) = 1/2(x + \mathbb{E}[1 - \hat{x}^*(X^*(1 - x, h), h)|h_{\text{alg}} = 1]) \\ &= 1/2(x + 1 - (1 - x)) = x.\end{aligned}$$

Next, observe that the algorithm, and thus the variance, remains unchanged for $x = 1/2$. For $x = 0$:

$$\begin{aligned}\text{Var}[Z_0] &= 1/2(\text{Var}[\hat{x}|h_{\text{alg}} = 0, x = 0] + \text{Var}[\hat{x}|h_{\text{alg}} = 1, x = 0]) \\ &= 1/2(\text{Var}[\hat{x}^*|x = 0] + \text{Var}[\hat{x}^*|x = 1]).\end{aligned}$$

Here, we used the fact that for any random variable $\text{Var}[\hat{x}^*(X^*(1, h), h)] = 1 - \text{Var}[\hat{x}^*(X^*(1, h), h)]$. Therefore, we get that $\text{Var}[\hat{x}|x = 0] \leq \max\{\text{Var}[\hat{x}^*|x = 0], \text{Var}[\hat{x}^*|x = 1]\}$ and we have not increased the cost. A symmetric analysis applies to $x = 1$.

Finally, we get that

$$\begin{aligned}p_{0,1/2} &= \Pr[X(0, h) = X(1/2, h)] \\ &= 1/2(\Pr[X(0, h) = X(1/2, h)|h_{\text{alg}} = 0] + \Pr[X(1, h) = X(1/2, h)|h_{\text{alg}} = 1]) \\ &= 1/2(p_{0,1/2}^* + p_{1/2,1}^*).\end{aligned}$$

By symmetry, we also get $p_{1/2,1} = 1/2(p_{0,1/2}^* + p_{1/2,1}^*)$ and thus $p_{0,1/2} = p_{1/2,1}$. Similarly,

$$\begin{aligned}G_{0,1/2} &= 1/2(G_{0,1/2}^* + (1 - G_{1/2,1}^*)) \\ G_{1/2,1} &= 1/2(G_{1/2,1}^* + (1 - G_{0,1/2}^*)).\end{aligned}$$

Notice that $G_{0,1/2} + G_{1/2,1} = 1$, which concludes the proof. $\blacktriangleleft$

E Analysis of the Unbiased Algorithm

In this appendix, we prove Theorem 5 which we restate here:

► **Theorem 5.** $\text{Var}[\hat{x}] \leq 1/12 \cdot (1 - 4^{-\ell}) + 1/4 \cdot 4^{-\ell} = 1/6 \cdot (1/2 + 4^{-\ell})$.

First, we state a technical lemma, whose proof appears below in Appendix E.1, that shows the periodicity of the variance in our algorithm.

► **Lemma 7.** For any $y \in [0, 1 - 2^{-\ell}]$, $\text{Var}[\hat{x}|x = y] = \text{Var}[\hat{x}|x = y + 2^{-\ell}]$.

As a result of this periodicity, we can continue the analysis, without loss of generality, under the assumption that $x \in [0, 2^{-\ell}]$. We first calculate several useful quantities:

$$\begin{aligned}\mathbb{E}[X] &= \mathbb{E}[X^2] = x \\ \mathbb{E}[h] &= (2^\ell - 1)/2\end{aligned}$$

$$\mathbb{E}[h^2] = (2^\ell - 1)(2^{\ell+1} - 1)/6$$

$$\mathbb{E}[X \cdot h | x \leq 2^{-\ell}] = 0 \quad (\text{as either } X = 0 \text{ or } h = 0, \text{ since } x \leq 2^{-\ell}).$$

We now proceed with calculating the variance.

$$\begin{aligned}
\text{Var}[\widehat{x}|x \leq 2^{-\ell}] &= \mathbb{E} \left[(X + (h - 0.5(2^\ell - 1)) \cdot 2^{-\ell})^2 \right] - x^2 \\
&= \mathbb{E} [X^2] + 2^{-2\ell} \cdot (\mathbb{E} [h^2] - \mathbb{E} [h] \cdot (2^\ell - 1) + 0.25(2^\ell - 1)^2) \\
&\quad + 2^{-\ell+1} \cdot \mathbb{E} [X \cdot h] - 2^{-\ell} \cdot (2^\ell - 1) \cdot \mathbb{E} [X] - x^2 \\
&= x + 2^{-2\ell} \cdot (\mathbb{E} [h^2] - \mathbb{E} [h] \cdot (2^\ell - 1) + 0.25(2^\ell - 1)^2) - 2^{-\ell} \cdot (2^\ell - 1) \cdot x - x^2 \\
&= x + 2^{-2\ell} \cdot (((2^\ell - 1)(2^{\ell+1} - 1)/6) - (2^\ell - 1)^2/2 + 0.25(2^\ell - 1)^2) - 2^{-\ell} \cdot (2^\ell - 1) \cdot x - x^2 \\
&= 1/12 \cdot (1 - 4^{-\ell}) + 2^{-\ell} x - x^2.
\end{aligned}$$

Finally, according to Lemma 7, we get that:

$$\text{Var}[\widehat{x}] = 1/12 \cdot (1 - 4^{-\ell}) + 2^{-\ell}(x \bmod 2^{-\ell}) - (x \bmod 2^{-\ell})^2. \quad (25)$$

This gives a worst-case bound, achieved for $x \in \{2^{-(\ell+1)} + i \cdot 2^{-\ell} \mid i \in [2^{\ell-1}]\}$, of

$$\text{Var}[\hat{x}] \leq 1/12 \cdot (1 - 4^{-\ell}) + 1/4 \cdot 4^{-\ell} = 1/6 \cdot (1/2 + 4^{-\ell}).$$

E.1 Proof of Lemma 7

Let $y \in [0, 1 - 2^{-\ell}]$ and denote $z = y + 2^{-\ell}$, $m = y \bmod 2^{-\ell}$, and $\zeta = \lfloor y \cdot 2^\ell \rfloor$. Notice that if $h < \zeta$, then Buffy will send $X = 1$ for both y and z . Similarly, if $h \geq \zeta + 1$, Buffy will send $X = 0$ for both y and z . Notice that, for any y and ℓ :

$$2^\ell y - \lfloor y \cdot 2^\ell \rfloor = 2^\ell (y \bmod 2^{-\ell}).$$

and thus

$$y \bmod 2^{-\ell} = y - \lfloor y \cdot 2^\ell \rfloor \cdot 2^{-\ell}.$$

Therefore, we can write:

$$\begin{aligned}(\hat{x}|x=y) &= (h-0.5(2^\ell-1)) \cdot 2^{-\ell} + \mathbb{1}_{h<\zeta} + \mathbb{1}_{(h=\zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \\ (\hat{x}|x=z) &= (h-0.5(2^\ell-1)) \cdot 2^{-\ell} + \mathbb{1}_{h<\zeta+1} + \mathbb{1}_{(h=\zeta+1) \wedge (r < 2^\ell \cdot (z \bmod 2^{-\ell}))}\end{aligned}$$

Denote $(h - 0.5(2^\ell - 1)) \cdot 2^{-\ell}$ by ψ . Thus, since $y \bmod 2^{-\ell} = z \bmod 2^{-\ell}$:

$$\begin{aligned} & \text{Var}[\widehat{x}|x = z] - \text{Var}[\widehat{x}|x = y] = \\ & \mathbb{E} \left[\left(\psi + \mathbb{1}_{h < \zeta + 1} + \mathbb{1}_{(h = \zeta + 1) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \right)^2 \right] - \\ & \mathbb{E} \left[\left((\psi + \mathbb{1}_{h < \zeta} + \mathbb{1}_{(h = \zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))})^2 \right) \right] - z^2 + y^2 = \\ & \mathbb{E} \left[\mathbb{1}_{h < \zeta + 1} + \mathbb{1}_{(h = \zeta + 1) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} + 2\psi \left(\mathbb{1}_{h < \zeta + 1} + \mathbb{1}_{(h = \zeta + 1) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \right) \right. \\ & \left. - \left(\mathbb{1}_{h < \zeta} + \mathbb{1}_{(h = \zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} + 2\psi \left(\mathbb{1}_{h < \zeta} + \mathbb{1}_{(h = \zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \right) \right) \right] \end{aligned}$$

$$\begin{aligned}
& -z^2 + y^2 = \\
& \mathbb{E} \left[\mathbb{1}_{h=\zeta} + \mathbb{1}_{(h=\zeta+1) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} - \left(\mathbb{1}_{(h=\zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \right) \right. \\
& \quad \left. + 2\psi \left(\mathbb{1}_{h=\zeta} + \mathbb{1}_{(h=\zeta+1) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} - \mathbb{1}_{(h=\zeta) \wedge (r < 2^\ell \cdot (y \bmod 2^{-\ell}))} \right) \right] - z^2 + y^2 = \\
& \mathbb{E} \left[\left(1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} + 2\psi \left(1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right) \right) \cdot \mathbb{1}_{h=\zeta} + \right. \\
& \quad \left. \left(\mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} + 2\psi \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right) \cdot \mathbb{1}_{h=\zeta+1} \right] - z^2 + y^2 = \\
& \mathbb{E} \left[1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} + 2\psi \left(1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right) \mid h = \zeta \right] \cdot \Pr[h = \zeta] + \\
& \mathbb{E} \left[\mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} + 2\psi \cdot \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \mid h = \zeta + 1 \right] \cdot \Pr[h = \zeta + 1] - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left(\mathbb{E} \left[2\psi \left(1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right) \mid h = \zeta \right] + \right. \\
& \quad \left. \mathbb{E} \left[2\psi \cdot \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \mid h = \zeta + 1 \right] \right) + 2^{-\ell} - z^2 + y^2 = (\text{as } h \text{ is independent of } r) \\
& 2^{-\ell} \cdot \left(\mathbb{E} \left[2 \left((\zeta - 0.5(2^\ell - 1)) \cdot 2^{-\ell} \right) \left(1 - \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right) \right] + \right. \\
& \quad \left. \mathbb{E} \left[2 \left((\zeta + 1 - 0.5(2^\ell - 1)) \cdot 2^{-\ell} \right) \cdot \mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right] \right) + 2^{-\ell} - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left(2 \left((\zeta - 0.5(2^\ell - 1)) \cdot 2^{-\ell} \right) + 2^{1-\ell} \cdot \mathbb{E} \left[\mathbb{1}_{r < 2^\ell \cdot (y \bmod 2^{-\ell})} \right] \right) + 2^{-\ell} - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left(2 \left((\zeta - 0.5(2^\ell - 1)) \cdot 2^{-\ell} \right) + 2^{1-\ell} \cdot (2^\ell \cdot (y \bmod 2^{-\ell})) \right) + 2^{-\ell} - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left((2\zeta - (2^\ell - 1)) \cdot 2^{-\ell} + 2 \cdot (y \bmod 2^{-\ell}) \right) + 2^{-\ell} - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left((2\zeta + 1) \cdot 2^{-\ell} + 2 \cdot (y \bmod 2^{-\ell}) \right) - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left((2 \lfloor y \cdot 2^\ell \rfloor + 1) \cdot 2^{-\ell} + 2 \cdot (y - \lfloor y \cdot 2^\ell \rfloor \cdot 2^{-\ell}) \right) - z^2 + y^2 = \\
& 2^{-\ell} \cdot \left(2^{-\ell} + 2y \right) - (z - y)(z + y) = 0. \quad \blacktriangleleft
\end{aligned}$$

F Generalization to k Bits

F.1 General Quantized Algorithm

We use a hash function h such that $h \in \{0, 1\}^\ell$ is uniformly distributed. Let $A \sim U[0, 1]$ be independent of h .

$$C = \lfloor (2^k - 1) \cdot x \rfloor$$

$$p = (2^k - 1) \cdot x - \lfloor (2^k - 1) \cdot x \rfloor$$

$$R = 2^k - 1$$

We then set

$$X \triangleq \begin{cases} C + 1 & \text{if } p \geq (A + h)2^{-\ell} \\ C & \text{otherwise} \end{cases}$$

We send X to the Angel which estimates

$$\hat{x} = \frac{X + (h - 0.5(2^\ell - 1)) \cdot 2^{-\ell}}{R}.$$

To show that our protocol is unbiased, notice that: $\mathbb{E}[X] = R \cdot x$ and that $\mathbb{E}[h] = 0.5(2^\ell - 1)$.

F.2 Lower Bounds

Similarly to the 1-bit case, we consider the discrete distribution over

$$\left\{ i \cdot \left(\frac{1}{3 \cdot 2^{k-1} - 1} \right) \mid i \in \{0, 1, \dots, 3 \cdot 2^{k-1} - 1\} \right\}.$$

We set $a_{1/2} = \frac{\sqrt{2}-1}{2^{k-1}}$ and $a_0 = a_1 = \frac{1-a_{1/2}}{2^k}$ and

$$\forall i : q \left(i \cdot \left(\frac{1}{3 \cdot 2^{k-1} - 1} \right) \right) = a_{(i \bmod 3)/2}.$$

When each consecutive set of three points has the same probability, one can derive an optimal algorithm with precisely two values between each such triplet. The optimal choice of locations of the values in each triplet is similar to our single-bit analysis of the previous subsection, i.e., one should have a values at

$$\left\{ \frac{\sqrt{2}-1+i}{2^{k-1}} \mid i \in \{0, 1, \dots, 2^{k-1} - 1\} \right\} \cup \left\{ \frac{i}{2^{k-1}} \mid i \in \{1, \dots, 2^{k-1}\} \right\}.$$

We turn into calculating the cost. Notice that every triplet has a width of $\frac{2}{3 \cdot 2^{k-1} - 1}$. Therefore, the cost now reduces, compared to the 1-bit analysis, by a factor of $\left(\frac{2}{3 \cdot 2^{k-1} - 1} \right)^2$. That is, we get a lower bound of $\frac{3-2\sqrt{2}}{(3 \cdot 2^{k-1} - 1)^2} = \frac{3-2\sqrt{2}}{2.25(2^k - 2/3)^2}$.

We note that, for large values of k and ℓ , our variance is within 10% of the lower bound, as

$$\lim_{k, \ell \rightarrow \infty} \frac{\frac{1}{12(2^k - 1)^2}}{\frac{3-2\sqrt{2}}{2.25(2^k - 2/3)^2}} = \frac{9 + 6\sqrt{2}}{16} \approx 1.093.$$

G Limiting Algorithm Uniformness Proof

Recall that the algorithm uses $h \sim U[0, 1]$ where Buffy sends

$$X \triangleq \begin{cases} 1 & \text{if } x \geq h \\ 0 & \text{otherwise} \end{cases}$$

and Angel estimates $\hat{x} = X + h - 0.5$.

► **Lemma 8.** *For a fixed value of x , it holds that $\hat{x} \sim U[x - \frac{1}{2}, x + \frac{1}{2}]$.*

Proof. Let $Z = \mathbb{1}_{h \leq x} + h$. We have that $Z \sim U[x, 1 + x]$, i.e.,

$$f_Z(z) = \begin{cases} 1 & \text{if } z \in [x, 1 + x] \\ 0 & \text{Otherwise} \end{cases}.$$

This is because

$$\Pr[Z \leq z] = \begin{cases} 1 & \text{if } z \geq 1 + x \\ z - x & \text{if } z \in (x, 1 + x) \\ 0 & \text{if } z \leq x \end{cases}.$$

Therefore,

$$X + h - 1/2 = Z - 1/2 \sim U[x - 1/2, (1 + x) - 1/2].$$

This concludes the proof. ◀

► **Corollary 9.** *Our estimator is unbiased, i.e., $\mathbb{E}[\hat{x}] = x$.*

► **Corollary 10.** *Our variance is constant for all $x \in [0, 1]$ and satisfies $\text{Var}[\hat{x}] = \frac{1}{12}$.*

H Reducing Algorithms to Monotone T, Z_0 Functions

We now show how an algorithm can be represented as described in Section 6.2. Fixing the shared randomness value h , Angel estimates $\hat{x}$ solely based on the sent bit X ; denote these values by $A_X(h)$. Without loss of generality, assume that $\forall h : A_1(h) \geq A_0(h)$.² This means that, in an optimal algorithm, Buffy should send $X = 1$ if $x \geq \frac{A_0(h) + A_1(h)}{2}$, as otherwise the error would be suboptimal for any x not satisfying the condition. In particular, this means that we can express Buffy's algorithm using a threshold function T .

Next, we claim that the threshold function can be considered monotone, without increasing the cost. To that end, we first consider a finite shared randomness $h \in [2^\ell]$. In such a case, if there exists some $h_1 > h_2 \in [2^\ell]$ such that $T(h_1) < T(h_2)$, we can modify the algorithm as follows: For all $h \notin \{h_1, h_2\}$, no modification is made. If $h = h_1$, then the modified algorithm works as if $h = h_2$, and vice versa. Following this process, we can sort T until it becomes monotone. A similar argument can be made for the continuous ($h \in [0, 1]$) case (possibly with an additional ϵ discretization cost).

We proceed with showing that there exists an optimal algorithm in which $Z_1(h) = 1 - Z_0(1 - h)$. This is achieved using a symmetry observation. Specifically, if an algorithm does not satisfy the above, consider its “dual algorithm”: instead of sending x using $T(h)$, we send $x' = 1 - x$ using $T'(h) = 1 - T(1 - h)$; similarly, Angel estimates $\hat{x}' = 1 - \hat{x}$. Then, if both Buffy and Angel use the shared randomness to implicitly agree on whether to run the original or dual algorithms, the cost can only decrease. The proof follows similarly to that of Lemma 6 (Appendix D.3).

² If for some h , $A_0(h) > A_1(h)$, there exists an equivalent algorithm that replaces the role of $X = 0$ and $X = 1$ for this specific h .

I Truncated Dithering

We now analyze the cost attainable by truncating the subtractive dithering algorithm to some interval $[z, 1 - z]$. Let $h \sim U[0, 1]$ be a shared uniform random variable. Consider sending

$$X = \begin{cases} 1 & \text{if } x \geq h \\ 0 & \text{otherwise} \end{cases}$$

similarly to our algorithm for $\ell \rightarrow \infty$ (see Section 5). However, unlike our algorithm, for a parameter $z \in [0, 1/2]$, Angel estimates x as

$$\hat{x} = \min \{ \max \{ X + h - 1/2, z \}, 1 - z \}.$$

That is, we truncate the estimation to the interval $[z, 1 - z]$, for some parameter $z \in [0, 0.5]$ that we determine later.

► **Lemma 11.** *For the $z \in [0, 1/2]$ that satisfies $1/24 + z^2/2 + (2z^3)/3 = 0$ ($z \approx 0.17349$), the cost of the above algorithm is $2/3 \cdot z^3 + 1/2 \cdot z^2 + 1/24 \approx 0.0602$.*

Proof. Assume, without loss of generality, that $x \in [0, 1/2]$. According to Lemma 8, we have that

$$X + h - 1/2 \sim U[x - 1/2, x + 1/2].$$

Therefore, Angel will estimate x as follows: With probability $1/2 + z - x$, $\hat{x} = z$; with probability $\max \{0, x + 1/2 - (1 - z)\}$, $\hat{x} = 1 - z$; and otherwise $\hat{x} \sim U[z, \min \{x + 1/2, 1 - z\}]$. We proceed with a case analysis. First, let us consider the $(x < 1/2 - z)$ case. This yields

$$\hat{x} = \begin{cases} \text{uniform on } [z, x + 1/2] & \text{with probability } 1/2 + z - x \\ z & \text{otherwise} \end{cases}.$$

Therefore, the cost would be

$$\begin{aligned} & (1/2 + z - x) \cdot (z - x)^2 + (1/2 + x - z) \cdot \int_z^{x+1/2} \frac{(t - x)^2}{1/2 + x - z} dt \\ &= (1/2 + z - x) \cdot (z - x)^2 + 1/24 + x^3/3 - x^2z + xz^2 - z^3/3 \\ &= 1/24 + x^2/2 - (2x^3)/3 - xz + 2x^2z + z^2/2 - 2xz^2 + (2z^3)/3. \end{aligned}$$

We have that the derivative with respect to x is:

$$-2x^2 + x(4z + 1) - z(1 + 2z).$$

Therefore, the potential extrema are $x \in \{0, 1/2 - z\}$ and when the derivative vanishes, which gives $x = z$ (the other extreme point is not in $[0, 0.5]$). We then get

$$\text{cost} = \begin{cases} 2/3 \cdot z^3 + 1/2 \cdot z^2 + 1/24 & \text{if } x = 0 \\ 1/24 & \text{if } x = z \\ 1/12 - 2z^2 + 16/3 \cdot z^3 & \text{if } x = 1/2 - z \end{cases}$$

Next, we consider the $x \geq 1/2 - z$ case. In such a case, we get that

$$\hat{x} = \begin{cases} \text{uniform on } [z, 1 - z] & \text{with probability } 1 - 2z \\ z & \text{with probability } 1/2 + z - x \\ 1 - z & \text{otherwise (w.p. } z + x - 1/2) \end{cases}$$

Then, our cost is:

$$\begin{aligned} (1/2 + z - x) \cdot (z - x)^2 + (z + x - 1/2) \cdot (1 - z)^2 + (1 - 2z) \cdot \int_z^{1-z} \frac{(t - x)^2}{1 - 2z} dt \\ = -1/6 + (3x^2)/2 - x^3 + z - xz + x^2z - z^2 - 2xz^2 + (4z^3)/3. \end{aligned}$$

We then get

$$\text{cost} = -3x^2 - z(1 + 2z) + x(3 + 2z).$$

Therefore, the potential extrema are $x \in \{1/2 - z, 1/2\}$ and where $\mathbb{E}[\text{cost}] = 0$, which gives $x = \frac{3+2z-\sqrt{9-20z^2}}{6}$. This yields

$$\text{cost} = \begin{cases} 1/12 - 2z^2 + (16z^3)/3 & \text{if } x = 1/2 - z \\ 4/3z^3 - 2z^2 + 3/4z + 1/12 & \text{if } x = 1/2 \\ 1/12 + z - 2z^2 + 20/27 \cdot z^3 + \sqrt{9 - 20z^2} \cdot (-1/12 + 5/27 \cdot z^2) & \text{if } x = \frac{3+2z-\sqrt{9-20z^2}}{6} \end{cases}$$

It follows that

$$1/12 + z - 2z^2 + 20/27 \cdot z^3 + \sqrt{9 - 20z^2} \cdot (-1/12 + 5/27 \cdot z^2) \leq 4/3z^3 - 2z^2 + 3/4z + 1/12$$

for all $z \in [0, 0.5]$, and therefore we focus on $x \in \{0, 1/2 - z, 1/2\}$. Notice that $\min_{z \in [0, 1/2]} 1/12 - 2z^2 + (16z^3)/3 = 1/24$, which is achieved for $z = 1/24$.

Finally, by choosing the z value which minimizes

$$\max \left\{ 2z(0.5 - z)^2 + \int_z^{1-z} (0.5 - t)^2 dt, (0.5 + z) \cdot z^2 + \int_z^{0.5} t^2 dt \right\},$$

which is obtained for the z value that satisfies $1/24 + z^2/2 + (2z^3)/3 = 0$ ($z \approx 0.17349$), we get an expected worst-case squared error of $\approx 0.0602 \approx 1/16.6$.

As a side note, one can obtain a slightly stronger bound by further truncating the estimations to $\epsilon \pm 1/2$, where $\epsilon = X + h - 1/2$. For example, if $\epsilon = -0.4$ then the algorithm should not estimate $\hat{x} \approx 0.173$ as x is guaranteed to be at most 0.1. Instead, the algorithm would estimate:

$$\hat{x} = \max(\min(\epsilon, \max(1 - z, \epsilon - 1/2)), \min(z, \epsilon + 1/2)).$$

In such a case, we can choose $z \approx 0.182$ and get a cost of ≈ 0.05824 . For simplicity, we omit the technical details. $\blacktriangleleft$

J Convex-combination Biased Adaptation for Subtractive Dithering

Here, we analyze the algorithm in which Buffy sends (for a shared $h \sim U[0, 1]$)

$$X = \begin{cases} 1 & \text{if } x \geq h \\ 0 & \text{otherwise} \end{cases},$$

and Angel estimates, for some $\alpha \in [0, 1]$,

$$\hat{x} = \alpha \cdot h + (1 - \alpha) \cdot X.$$

Notice that

$$\mathbb{E}[\hat{x}] = \alpha/2 + (1 - \alpha) \cdot x.$$

$$\mathbb{E}[h^2] = 1/3.$$

$$\mathbb{E}[h \cdot X] = \int_0^x t dt = x^2/2.$$

$$\mathbb{E}[\hat{x}^2] = \alpha^2 \cdot \mathbb{E}[h^2] + (1 - \alpha)^2 \cdot x + 2\alpha(1 - \alpha)\mathbb{E}[h \cdot X] = \alpha^2/3 + (1 - \alpha)^2 \cdot x + \alpha(1 - \alpha) \cdot x^2.$$

We compute the expected squared error:

$$\begin{aligned} \mathbb{E}[(\hat{x} - x)^2] &= \mathbb{E}[\hat{x}^2] - 2x\mathbb{E}[\hat{x}] + x^2 \\ &= \alpha^2/3 + (1 - \alpha)^2 \cdot x + \alpha(1 - \alpha) \cdot x^2 - 2x(\alpha/2 + (1 - \alpha) \cdot x) + x^2 \\ &= x - x^2 - 3x\alpha + 3x^2\alpha + \alpha^2/3 + x\alpha^2 - x^2\alpha^2. \end{aligned} \tag{26}$$

We then get

$$\frac{\partial \mathbb{E}[(\hat{x} - x)^2]}{\partial x} = (-1 + 2x)(-1 + 3\alpha - \alpha^2).$$

Therefore, the possible extrema are $\{0, 1/2, 1\}$. Minimizing (26), we get that the optimal choice is $\alpha = 2 - \phi \approx 0.382$, which gives $\mathbb{E}[(\hat{x} - x)^2] \leq 5/3 - \phi \approx 0.04863$.

K Analysis of the algorithm of Section 6.2.2

We first derive several quantities that will be useful for calculating the cost.

$$\begin{aligned} \text{--- } \mathbb{E}[X] &= \mathbb{E}[X^2] = \begin{cases} 0 & \text{if } x < (1 - \alpha)/2 \\ (i + 1)/2^\ell & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ 1 & \text{if } x \geq (1 + \alpha)/2 \end{cases} \\ \text{--- } \mathbb{E}[h/(2^\ell - 1)] &= 1/2. \\ \text{--- } \mathbb{E}[\hat{x}] &= \alpha \cdot \mathbb{E}[h/(2^\ell - 1)] + (1 - \alpha) \cdot \mathbb{E}[X] = \begin{cases} \alpha/2 & \text{if } x < (1 - \alpha)/2 \\ \alpha/2 + (1 - \alpha) \cdot (i + 1)/2^\ell & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ 1 - \alpha/2 & \text{if } x \geq (1 + \alpha)/2 \end{cases} \\ \text{--- } \mathbb{E}[h/(2^\ell - 1) \cdot X] &= \begin{cases} 0 & \text{if } x < (1 - \alpha)/2 \\ i \cdot (i + 1)/(2^{\ell+1} \cdot (2^\ell - 1)) & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ 1/2 & \text{if } x \geq (1 + \alpha)/2 \end{cases} \end{aligned}$$

$$\blacksquare \mathbb{E}[(h/(2^\ell - 1))^2] = (2^{\ell+1} - 1)/(6 \cdot (2^\ell - 1)).$$

Next, we calculate the second moment of the estimate:

$$\mathbb{E}[\hat{x}^2] = \alpha^2 \cdot \mathbb{E}[(h/(2^\ell - 1))^2] + 2\alpha(1 - \alpha) \cdot \mathbb{E}[h/(2^\ell - 1) \cdot X] + (1 - \alpha)^2 \cdot \mathbb{E}[X^2] =$$

$$\begin{cases} \Psi & \text{if } x < (1 - \alpha)/2 \\ \Psi + \Gamma_i & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ \Psi + \alpha(1 - \alpha) + (1 - \alpha)^2 & \text{if } x \geq (1 + \alpha)/2 \end{cases},$$

where

$$\Psi = \alpha^2 \cdot (2^{\ell+1} - 1)/(6 \cdot (2^\ell - 1))$$

and

$$\Gamma_i = \alpha(1 - \alpha) \cdot i \cdot (i + 1)/(2^\ell \cdot (2^\ell - 1)) + (1 - \alpha)^2 \cdot (i + 1)/2^\ell.$$

Finally, we are ready to express the expected squared error:

$$\mathbb{E}[(\hat{x} - x)^2] = \mathbb{E}[\hat{x}^2] - 2x\mathbb{E}[\hat{x}] + x^2 =$$

$$\begin{cases} \Psi - x\alpha + x^2 & \text{if } x < (1 - \alpha)/2 \\ \Psi + \Gamma_i - x\alpha - x \cdot (1 - \alpha) \cdot (i + 1)/2^{\ell-1} + x^2 & \text{if } x \in \mathcal{I}_i, i \in [2^\ell - 1] \\ \Psi - (1 - x)\alpha + (1 - x)^2 & \text{if } x \geq (1 + \alpha)/2 \end{cases}.$$

Solving $\text{cost} = \min_\alpha \max_x \mathbb{E}[(\hat{x} - x)^2]$ yields

$$\text{cost} = \frac{(2^\ell - 1) \cdot (2^{\ell+1} - 1) \cdot \left(2 - 3 \cdot 2^\ell + 2^{\ell/2} \cdot \sqrt{5 \cdot 2^\ell - 8}\right)^2}{24 \cdot (4^\ell - 2^\ell + 1)^2},$$

which is obtained for

$$x \in \left\{0, (1 - \alpha)/2 + (2^{\ell-1} - 1) \cdot \frac{\alpha}{2^\ell - 1}, (1 - \alpha)/2 + 2^{\ell-1} \cdot \frac{\alpha}{2^\ell - 1}, 1\right\}$$

and

$$\alpha = \frac{1 - 5 \cdot 2^{\ell-1} + 3/2 \cdot 4^\ell - \sqrt{2^{\ell-2} \cdot (2^\ell - 1)^2 \cdot (5 \cdot 2^\ell - 8)}}{4^\ell - 2^\ell + 1}.$$

L Limiting Biased Algorithm

In the limiting algorithm, Buffy sends:

$$X = \begin{cases} 0 & \text{if } x < (1 - \alpha)/2 \\ \mathbb{1}_{h \leq \frac{x - (1 - \alpha)/2}{\alpha}} & \text{if } x \in [(1 - \alpha)/2, (1 + \alpha)/2] \\ 1 & \text{if } x \geq (1 + \alpha)/2 \end{cases}.$$

In turn, Angel estimates:

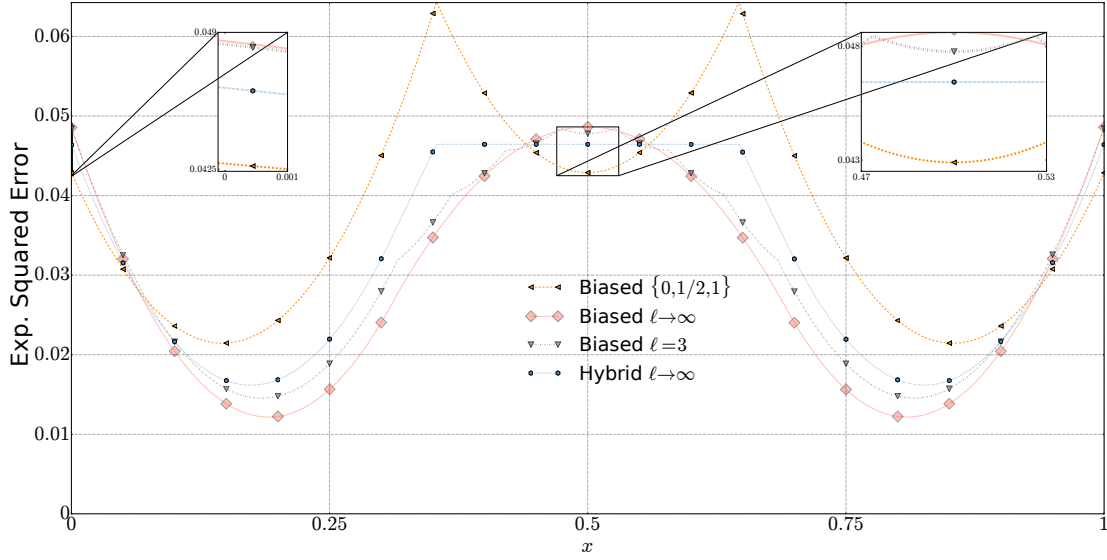
$$\hat{x} = \alpha \cdot h + (1 - \alpha) \cdot X$$

Let us calculate several useful quantities:

$$\begin{aligned}
\mathbb{E}[X] = \mathbb{E}[X^2] &= \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{x-(1-\alpha)/2}{\alpha} & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
\mathbb{E}[\hat{x}] &= \alpha/2 + (1-\alpha) \cdot \mathbb{E}[X] = \alpha/2 + (1-\alpha) \cdot \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{x-(1-\alpha)/2}{\alpha} & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
\mathbb{E}[h^2] &= 1/3. \\
\mathbb{E}[h \cdot X] &= \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \int_0^{\frac{x-(1-\alpha)/2}{\alpha}} t dt & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1/2 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
&= \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{1}{2} \left(\frac{x-(1-\alpha)/2}{\alpha} \right)^2 & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1/2 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
\mathbb{E}[\hat{x}^2] &= \alpha^2 \cdot \mathbb{E}[h^2] + (1-\alpha)^2 \cdot \mathbb{E}[X^2] + 2\alpha(1-\alpha)\mathbb{E}[h \cdot X] \\
&= \alpha^2/3 + (1-\alpha)^2 \cdot \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{x-(1-\alpha)/2}{\alpha} & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
&\quad + 2\alpha(1-\alpha) \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{1}{2} \left(\frac{x-(1-\alpha)/2}{\alpha} \right)^2 & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1/2 & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
&= \begin{cases} \alpha^2/3 & \text{if } x < (1-\alpha)/2 \\ \alpha^2/3 + (1-\alpha)^2 \cdot \frac{x-(1-\alpha)/2}{\alpha} + \alpha(1-\alpha) \cdot \left(\frac{x-(1-\alpha)/2}{\alpha} \right)^2 & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ \alpha^2/3 + 1 - \alpha & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
&= \begin{cases} \alpha^2/3 & \text{if } x < (1-\alpha)/2 \\ 3/4 - x^2 - 1/(4\alpha) + x^2/\alpha - (3\alpha)/4 + (7\alpha^2)/12 & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ \alpha^2/3 + 1 - \alpha & \text{if } x \geq (1+\alpha)/2 \end{cases}
\end{aligned}$$

We proceed with analyzing the expected squared error:

$$\begin{aligned}
\mathbb{E}[(\hat{x} - x)^2] &= \mathbb{E}[\hat{x}^2] - 2x\mathbb{E}[\hat{x}] + x^2 \\
&= x^2 + \begin{cases} \alpha^2/3 & \text{if } x < (1-\alpha)/2 \\ 3/4 - x^2 - 1/(4\alpha) + x^2/\alpha - (3\alpha)/4 + (7\alpha^2)/12 & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ \alpha^2/3 + 1 - \alpha & \text{if } x \geq (1+\alpha)/2 \end{cases} \\
&\quad - 2x \cdot \left(\alpha/2 + (1-\alpha) \cdot \begin{cases} 0 & \text{if } x < (1-\alpha)/2 \\ \frac{x-(1-\alpha)/2}{\alpha} & \text{if } x \in [(1-\alpha)/2, (1+\alpha)/2] \\ 1 & \text{if } x \geq (1+\alpha)/2 \end{cases} \right)
\end{aligned}$$



■ **Figure 4** An illustration of the expected squared errors that motivate our choice of creating a hybrid of the optimal $\{0, 1/2, 1\}$ and the biased $\ell \rightarrow \infty$ algorithms.

$$= \begin{cases} x^2 + \alpha^2/3 - x \cdot \alpha & \text{if } x < (1 - \alpha)/2 \\ \frac{3\alpha-1}{4\alpha} + x^2/\alpha + (7\alpha^2 - 9\alpha)/12 - x \cdot \alpha - 2x(1 - \alpha) \cdot \frac{x-(1-\alpha)/2}{\alpha} & \text{if } x \in [(1 - \alpha)/2, (1 + \alpha)/2] \\ x^2 + \alpha^2/3 + 1 - \alpha - 2x(1 - \alpha/2) & \text{if } x \geq (1 + \alpha)/2 \end{cases}$$

Choosing $\alpha = 2 - \phi$ as before, which minimizes the worst-case expected squared error, we get that

$$\mathbb{E}[(\hat{x} - x)^2] = \begin{cases} x^2 + (\phi - 2) \cdot x + (5/3 - \phi) & \text{if } x < (\phi - 1)/2 \\ \frac{2\phi-3}{\phi-2} \cdot x^2 + (\phi - 1) \cdot x + \frac{23-15\phi}{12} & \text{if } x \in [(1 - \alpha)/2, (1 + \alpha)/2] \\ x^2 - \phi \cdot x + 2/3 & \text{if } x > (3 - \phi)/2 \end{cases}$$

Our analysis indicates that the cost becomes $5/3 - \phi \approx 0.04863$, which is reached for $x \in \{0, 1/2, 1\}$.

M Improving the cost further using a hybrid algorithm

As mentioned, for the case $x \in [0, 1]$, our algorithms above do not improve when given access to more than $\ell = 3$ bits. This suggests that an optimal algorithm, unlike the solution from Section 6.2.2, may need to use a non-uniform or *probabilistic* partitioning of the $[0, 1]$ interval using the α parameter.

We now explore how to reduce the cost using randomized thresholding, achieved through probabilistic multiplexing of the above algorithms. That is, Buffy and Angel will randomly select the executed protocol using the shared randomness, thus achieving implicit coordination.

To simplify the notation, we use $A_{[0,1]}$ to denote our general (i.e., $x \in [0, 1]$) with $\ell \rightarrow \infty$ (also given explicitly in Appendix L) algorithm, and $A_{\{0,1/2,1\}}$ to denote our algorithm for when x is guaranteed to be in $\{0, 1/2, 1\}$. Our observation is that $A_{[0,1]}$ behaves differently than $A_{\{0,1/2,1\}}$. Specifically, for the first, the expected squared error is *maximized* at $\{0, 1/2, 1\}$, while for the latter, the expected squared error is lower at these points. This suggests that by randomly choosing which of these algorithms to execute one can lower the cost.

In particular, we propose to multiplex between $A_{[0,1]}$ and $A_{\{0,1/2,1\}}$ as follows. With probability p , to be determined later, both Buffy and Angel use $A_{[0,1]}$ and otherwise $A_{\{0,1/2,1\}}$, using the shared randomness to implicitly decide on the protocol. This means that the expected squared error becomes:

$$\mathbb{E}[(\hat{x} - x)^2] = \mathbb{E}[(\hat{x} - x)^2 | \text{running } A_{[0,1]}] \cdot p + \mathbb{E}[(\hat{x} - x)^2 | \text{running } A_{\{0,1/2,1\}}] \cdot (1 - p).$$

We get that the cost, optimized for $p = \phi - 1$, is $\frac{6\sqrt{10}+11\sqrt{5}-18\sqrt{2}-17}{24} \approx 0.04644$ (obtained for $x \in \left(\{0,1\} \cup \left[\frac{1}{2\sqrt{2}}, 1 - \frac{1}{2\sqrt{2}}\right]\right)$). Notice that the cost of the algorithm is within 3.01% from the lower bound in Section 4.

Figure 4 illustrates how the non-hybrid (Section 6.2.2) $\ell = 3$ algorithm has a lower cost than that of $\ell \rightarrow \infty$. However, as its worst-case expected squared error is not at $x = 1/2$, it does not multiplex as well with $A_{\{0,1/2,1\}}$. Specifically, a hybrid algorithm that uses $\ell = 3$ bits instead of the limit ($\ell \rightarrow \infty$) algorithm results in a *higher* cost of $102/361 - 1/(3\sqrt{2}) \approx 0.04685$ (which is obtained for $p = 2/3$).

The above hybrid algorithms use unbounded shared randomness. In cases where we wish to use a small number of shared bits, we can approximate the better algorithm (that uses $\ell \rightarrow \infty$ for $A_{[0,1]}$); below we give a couple of examples.

Example I: Consider using $\ell = 4$ bits. In this case, we use $p = 3/4$ and use the 2-bit algorithm from Section 6.2.2 as $A_{[0,1]}$. The cost is then $\frac{1049-169\sqrt{2}-430\sqrt{3}}{1352} \approx 0.0482$ (obtained for $x \in \left\{\frac{4+\sqrt{3}}{13}, \frac{9-\sqrt{3}}{13}\right\}$), which improves over the 3-bit algorithm.

Example II: Consider using one random byte ($\ell = 8$). In that case, we use $p = 11/16$ together with the 4-bit algorithm. The cost then becomes $\frac{1830635-1232945\sqrt{2}}{1858592} \approx 0.04680$ (obtained for $x \in \left\{\frac{109+6\sqrt{2}}{241}, \frac{132-6\sqrt{2}}{241}\right\}$). This further improves over the cost of Example I, over the best hybrid solution with $\ell = 3$ (that uses unbounded randomness to represent the $p = 2/3$ value), and is within 1% of the unbounded shared randomness algorithm.

N Analysis of the Linear Sigmoid (Section 6.2.3)

First, we have (see Section 6.2) that the estimate function for $X = 1$ is:

$$Z_1(h) = 1 - Z_0(1 - h) = \begin{cases} 2\alpha & \text{if } h < h_0 \\ 2\alpha + (1 - 2\alpha) \cdot \frac{h-h_0}{1-2h_0} & \text{if } h \in [h_0, 1-h_0] \\ 1 & \text{otherwise} \end{cases}.$$

For $x \in [\alpha, 1 - \alpha]$, denote by $T^{-1}(x) = \frac{(1-2h_0)(x-\alpha)}{1-2\alpha} + h_0$ the value such that $T(T^{-1}(x)) = x$. We proceed with computing the expectation:

$$\mathbb{E}[\hat{x}|x < \alpha](\implies X = 0) = h_0 \cdot (1 - 2\alpha) + \int_{h_0}^{1-h_0} \left((1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} \right) dh = 1/2 - \alpha$$

$$\mathbb{E}[\hat{x}|x > 1 - \alpha](\implies X = 1) = h_0 \cdot (1 + 2\alpha) + \int_{h_0}^{1-h_0} \left(2\alpha + (1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} \right) dh = 1/2 + \alpha$$

$$\begin{aligned} \mathbb{E}[\hat{x}|x \in [\alpha, 1 - \alpha]] &= h_0 \cdot (2\alpha) + h_0(1 - 2\alpha) \\ &\quad + \int_{h_0}^{T^{-1}(x)} \left(2\alpha + (1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} \right) dh + \int_{T^{-1}(x)}^{1-h_0} \left((1 - 2\alpha) \cdot \frac{h - h_0}{1 - 2h_0} \right) dh \end{aligned}$$

$$\begin{aligned}
&= h_0 + 2\alpha(T^{-1}(x) - h_0) + \int_{h_0}^{1-h_0} \left((1-2\alpha) \cdot \frac{h-h_0}{1-2h_0} \right) dh \\
&= 1/2 + (-1+2h_0)\alpha + 2\alpha(T^{-1}(x) - h_0) = 1/2 - \alpha + 2\alpha T^{-1}(x)
\end{aligned}$$

Therefore, we have:

$$\mathbb{E}[\hat{x}] = \begin{cases} 1/2 - \alpha & \text{if } x < \alpha \\ 1/2 - \alpha + 2\alpha \left(\frac{(1-2h_0)(x-\alpha)}{1-2\alpha} + h_0 \right) & \text{if } x \in [\alpha, 1-\alpha] \\ 1/2 + \alpha & \text{otherwise} \end{cases}$$

Next, we calculate the second moment of the estimate:

$$\begin{aligned}
\mathbb{E}[(\hat{x})^2 | x < \alpha] (\implies X = 0) &= h_0 \cdot (1-2\alpha)^2 + \int_{h_0}^{1-h_0} \left((1-2\alpha) \cdot \frac{h-h_0}{1-2h_0} \right)^2 dh \\
&= 1/3 + h_0/3 - 4\alpha/3 - 4h_0\alpha/3 + 4\alpha^2/3 + 4h_0\alpha^2/3
\end{aligned}$$

$$\begin{aligned}
\mathbb{E}[(\hat{x})^2 | x > 1-\alpha] (\implies X = 1) &= h_0 \cdot (1+(2\alpha)^2) + \int_{h_0}^{1-h_0} \left(2\alpha + (1-2\alpha) \cdot \frac{h-h_0}{1-2h_0} \right)^2 dh \\
&= 1/3 + h_0/3 + 2\alpha/3 - 4h_0\alpha/3 + 4\alpha^2/3 + 4h_0\alpha^2/3
\end{aligned}$$

$$\begin{aligned}
\mathbb{E}[(\hat{x})^2 | x \in [\alpha, 1-\alpha]] &= h_0 \cdot ((1-2\alpha)^2 + (2\alpha)^2) \\
&\quad + \int_{h_0}^{T^{-1}(x)} \left(2\alpha + (1-2\alpha) \cdot \frac{h-h_0}{1-2h_0} \right)^2 dh + \int_{T^{-1}(x)}^{1-h_0} \left((1-2\alpha) \cdot \frac{h-h_0}{1-2h_0} \right)^2 dh \\
&= h_0 \cdot ((1-2\alpha)^2 + (2\alpha)^2) + \frac{(1-2h_0)(x-\alpha)(x^2+4x\alpha+7\alpha^2)}{3-6\alpha} \\
&\quad + \frac{(2h_0-1)(-1+x+\alpha)(1+x+x^2-4x\alpha+\alpha(-5+7\alpha))}{3-6\alpha} \\
&= \frac{1-6\alpha-6\alpha x^2(-1+2h_0)+12\alpha^2-14\alpha^3+h_0(1-6\alpha+24\alpha^2-20\alpha^3)}{3-6\alpha}
\end{aligned}$$

Putting it together, we get:

$$\mathbb{E}[\hat{x}^2] = \begin{cases} 1/3 + h_0/3 - 4\alpha/3 - 4h_0\alpha/3 + 4\alpha^2/3 + 4h_0\alpha^2/3 & \text{if } x < \alpha \\ \frac{1-6\alpha-6\alpha x^2(-1+2h_0)+12\alpha^2-14\alpha^3+h_0(1-6\alpha+24\alpha^2-20\alpha^3)}{3-6\alpha} & \text{if } x \in [\alpha, 1-\alpha] \\ 1/3 + h_0/3 + 2\alpha/3 - 4h_0\alpha/3 + 4\alpha^2/3 + 4h_0\alpha^2/3 & \text{otherwise} \end{cases}$$

By solving

$$\min_{h_0 \in [0, 1/2], \alpha \in [0, 1/2]} \max_{x \in [0, 1]} \mathbb{E}[(\hat{x} - x)^2] = \min_{h_0 \in [0, 1/2], \alpha \in [0, 1/2]} \max_{x \in [0, 1]} \mathbb{E}[\hat{x}^2] - 2x\mathbb{E}[\hat{x}] + x^2,$$

we get that the algorithm is optimized for $\alpha = 1/3$ and $h_0 = 1/4$, where the resulting cost is:

$$\begin{aligned}
\mathbb{E}[(\hat{x} - x)^2] &= \mathbb{E}[(\hat{x})^2] - 2x\mathbb{E}[\hat{x}] + x^2 \\
&= \begin{cases} 5/108 - x/3 + x^2 & \text{if } x < 1/3 \\ 5/108 & \text{if } x \in [1/3, 2/3] \\ 77/108 - 5x/3 + x^2 & \text{otherwise} \end{cases}
\end{aligned}$$

Therefore, the cost is $5/108 \approx 0.0463$, which is less than 0.9% higher than the 0.0459 lower bound (Section 4.1.2).