

Article

Asymptotically Optimal Adversarial Strategies for the Probability Estimation Framework

Soumyadip Patra *  and Peter Bierhorst 

Department of Mathematics, University of New Orleans, New Orleans, LA 70148, USA; plbierho@uno.edu

* Correspondence: spatra@uno.edu

Abstract: The probability estimation framework involves direct estimation of the probability of occurrences of outcomes conditioned on measurement settings and side information. It is a powerful tool for certifying randomness in quantum nonlocality experiments. In this paper, we present a self-contained proof of the asymptotic optimality of the method. Our approach refines earlier results to allow a better characterisation of optimal adversarial attacks on the protocol. We apply these results to the (2,2,2) Bell scenario, obtaining an analytic characterisation of the optimal adversarial attacks bound by no-signalling principles, while also demonstrating the asymptotic robustness of the PEF method to deviations from expected experimental behaviour. We also study extensions of the analysis to quantum-limited adversaries in the (2,2,2) Bell scenario and no-signalling adversaries in higher (n, m, k) Bell scenarios.

Keywords: device-independent quantum random number generation; quantum nonlocality; Bell inequalities; asymptotic equipartition property; min-entropy



Citation: Patra, S.; Bierhorst, P. Asymptotically Optimal Adversarial Strategies for the Probability Estimation Framework. *Entropy* **2023**, *25*, 1291. <https://doi.org/10.3390/e25091291>

Academic Editors: Hector Zenil and Rosario Lo Franco

Received: 4 July 2023

Revised: 16 August 2023

Accepted: 28 August 2023

Published: 2 September 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Randomness has proven to be a valuable resource for a multitude of tasks, be it computation or communication. In cryptography, access to reliable random bits is essential, since the security of various cryptographic primitives is known to be compromised if the incorporated randomness is of poor quality [1–3]. In the study of random network modelling, being able to sample random graphs uniformly and (reliably) at random is crucial [4]. And, for some problems, randomised algorithms are known to vastly outperform their deterministic counterparts [5].

A distinction between two notions of randomness, those of *process* and *product*, is discussed in [6] (chapter 8). Although both notions are tightly connected, randomness of a process refers to its unpredictability, while that of a product refers to a lack of pattern in it. An unpredictable process will, with high probability, produce a sequence (a string of bits, say) that is patternless; on the other hand, a seemingly irregular string of bits might not be unpredictable and instead be a probabilistic mixture of pre-recorded information. While product randomness suffices for tasks like Monte Carlo simulations, sampling and those involving randomised algorithms, cryptographic applications involving an adversary necessitate process randomness.

Process randomness, while being non-existent in the strictest interpretation of any classical theory, is permissible in quantum mechanics; an important example of this is quantum nonlocality as manifested in a *Bell experiment*. Quintessentially, the setup of a Bell experiment constitutes an entangled quantum system shared between two spatially separated stations A and B receiving inputs x and y , and recording outcomes a and b , respectively. If after n successive trials the observed correlations between the outcomes conditioned on the settings violate a *Bell inequality* then it can be ruled out that the outcomes were pre-assigned by some probabilistic mixture of deterministic processes. Also, the outcomes are (unpredictably) random, not only to the respective users of the devices at

the two stations but also to an adversary, even to one having a complete understanding of the Bell experiment. This relationship between nonlocality in quantum mechanics and its random nature is at the foundation of various device-independent random number generation protocols.

Device independence is considered a gold standard in cryptographic tasks such as quantum random number generation and quantum key distribution, in which the respective users are not required to know or trust the inner machinery of their devices, thus treating them as mere black boxes to which they can provide inputs and record outcomes. The only assumption that the experimental setup must satisfy is that the measurement choices of the devices must be uncorrelated with their inner workings. This is the measurement independence assumption, which is ultimately untestable but is tacitly assumed, arguably, in almost all scientific experiments. The no-signalling condition that the outcome recorded at each station is not influenced by the choice of measurement at the other station holds throughout the experiment because of the space-like separation between the stations and the impossibility of superluminal signalling in accordance with the special theory of relativity. Furthermore, the adversary trying to simulate the observed statistics may be considered computationally unbounded, a standard that falls under the paradigm of information-theoretic security. Over the years, technological advancement has facilitated loophole-free Bell nonlocality experiments, which have not only provided experimental validation to rule out a classical description of nature [7–10], but have also found practical applications in device-independent quantum randomness generation and device-independent quantum key distribution [11–13].

The probability estimation framework is a broadly applicable framework for performing device-independent quantum randomness generation (DIQRNG) upon a finite sequence of loophole-free Bell experiment data and involves direct estimation of the amount of certifiable randomness by obtaining high-confidence bounds on the conditional probability of the observed measurement outcomes conditioned on the measurement settings in the presence of classical side information [14–16]. Advantageous primarily for its demonstrated applicability to Bell tests with small Bell violations and high efficiency for a finite number of trials, it can also accommodate changing experimental conditions and allows early stoppage upon meeting certain criteria. Also, it can be extended to randomness generation with quantum devices beyond the device-independent scenario.

The probability estimation framework for DIQRNG is provably secure against adversaries who do not possess entanglement with the sources. Security against more general adversaries, with quantum entanglement with the sources, is possible with the quantum estimation framework [17], for which the constructions of the probability estimation framework can often be translated to the quantum estimation framework (as was carried out in [18]), so that progress with the former framework can often be used for the more general latter framework.

The asymptotic optimality of the probability estimation framework was discussed in [15]. The specific result of asymptotic optimality is as follows: given a sufficiently large number of trials sampling from a fixed behaviour (i.e., a set of quantum statistics), the amount of certified randomness per trial is arbitrarily close to a certain upper limit. Then [15] argues, appealing to convex geometry and the asymptotic equipartition property (AEP), that an adversary can always implement a probabilistic mixture of conditional probability distributions, independent and identically distributed across successive experimental trials, that generates observed statistics consistent with the fixed behaviour while not needing to generate more than that same upper limit of randomness per trial that is certified by the probability estimation framework. This is important in the sense that the framework certifies all the randomness conceded by the adversary in that particular attack, while also showing that there is no advantage to be gained for the adversary by resorting to (more sophisticated) memory attacks.

In this paper, we provide a full derivation of the asymptotic optimality of the probability estimation framework, filling in some steps omitted by [15], along the way obtaining a

better characterisation of the adversary's optimal probabilistic mixture for generating the observed statistics. Making precise the arguments from convex geometry, we explicitly describe the optimal attack that an adversary can employ with the minimum required number of different conditional distributions in convex mixture to simulate the observed statistics. Our improvement, with a more self-contained approach, upon the result in [15] is to reduce by one the cardinality of the adversary's (finite-cardinality) set from which the auxiliary random variable takes values. This random variable serves as her side-information and records which conditional distribution occurs in which trial. Specifically, we prove that the number of possible conditional distributions in her optimal probabilistic mixture attack *need not be more than one plus the dimension of the set of admissible distributions of a trial* (Theorem 4). (We assume the set of admissible probability distributions of a given trial to be closed and convex, where we can take the convex closure when this assumption is not met; then the dimension $\dim(C)$ of a non-empty convex subset C of X is the dimension of the smallest affine subset containing C .) An earlier result (Theorem 43 in [15] under the same assumptions) proved only that the cardinality of the value space of the adversary's side-information need not be more than two plus the dimension of the set of admissible distributions of a trial. Besides contributing to a methodological improvement, we have thus improved the result itself: a better understanding of the optimal attack in the asymptotic regime will establish a benchmark that will enable the implementer of the protocol to defend against these attack modes.

The central results on asymptotic optimality of the method of probability estimation comprise establishing an upper bound on the randomness per trial more than which the adversary need not concede (Theorems 3 and 4) and which is certified by the method of probability estimation (Theorems 5 and 6). Our derivation in Theorem 3 elucidates how only the classical form of the asymptotic equipartition property is needed for the probability estimation framework, allowing a simplified treatment. In addition to strengthening the result in Theorem 4, we have presented proofs for Theorems 5 and 6 (which have appeared previously in [15]), including more details and specifications where we deemed fit. For instance, in the proof for Theorem 6, enlisting the extreme value theorem we avoid an explicit analytic construction as presented in [15] (see Theorem 41 therein). We also consider the question of *robustness* of the probability estimation framework, not considered in [14,15]; we derive a sufficient condition (Theorem 7) for a probability estimation factor (optimised at a particular distribution) to certify randomness at a positive rate at a *statistically different* distribution.

We apply our results to the (2,2,2) Bell scenario (the scenario of two parties, two measurement settings and two outcomes), obtaining an analytic characterisation of the optimal attack of an adversary (restricted only by the no-signalling condition) holding classical side information. We show that the optimal adversarial attack involves a decomposition of the observed statistics in terms of a single extremal no-signalling (super-quantum) correlation and eight local deterministic correlations. The proof of optimality relies upon the fact that equal mixtures of two extremal no-signalling nonlocal super-quantum correlations are expressible as an equal mixture of four local deterministic correlations. We show that this result does not generalise to higher scenarios such as the (3,2,2), (2,3,2) and (2,2,3) Bell scenarios, thereby indicating that the possibility of an optimal attack involving only a single extremal strategy is only ensured in the minimal (2,2,2) Bell scenario. Furthermore, we considered the possibility of an adversary holding classical side information (and, hence, restricted to probabilistic attack strategies) but trying to simulate the observed statistics using quantum-achievable probability distributions, while conceding as little randomness as possible. Assuming uniform settings distribution, numerical studies restricted to a two-dimensional slice of the set of quantum-achievable distributions provided some initial evidence that the optimal quantum-achievable attack strategy involves only one extremal quantum correlation, but we were not able to settle this and have phrased it as a conjecture.

The rest of the article is organised as follows: In Section 2, we review the probability estimation framework where Theorem 1 formalises the central idea and Theorem 2 estab-

lishes a lower bound on the smooth conditional min-entropy of the sequence of outcomes conditioned on the settings and side-information. We also present a simplified proof of Lemma 1, an important result that allows the algorithm to execute the PEF method, compared to the proofs in [14,15]. In Section 3, we present our complete proof of asymptotic optimality, study the implications for finding an optimal adversarial attack strategy and derive a robustness result. In Section 4, we apply our results to the (2,2,2) Bell scenario obtaining an analytic characterisation of the optimal attack strategy for an adversary restricted only by the no-signalling condition. The optimal attack comprises a decomposition of the observed statistics in terms of a single Popescu–Rohrlich (PR) correlation and (up to) eight local deterministic correlations. We show that, for a higher number of parties, settings and/or outcomes, a crucial result from the (2,2,2) Bell scenario concerning equal mixtures of extremal nonlocal no-signalling correlations does not hold, and infer that the optimal attack may require more than one nonlocal distribution in general. Returning to the (2,2,2) scenario, we discuss a conjecture that the optimal strategy to mimic the observed statistics by means of a probabilistic mixture of quantum-achievable correlations constitutes only a single extremal quantum correlation and (up to) eight local deterministic correlations.

2. The Probability Estimation Framework

The probability estimation method relies on the probability estimation factor (PEF), which is a function that assigns a score to the results of a single trial of a quantum experiment, with higher scores corresponding to more randomness. The paradigmatic application is to a Bell nonlocality experiment comprising multiple spatially separated parties providing inputs (measurement settings) to measuring devices and recording outputs (observed outcomes); an experimental trial's results then consist of both the choice of inputs and the recorded outputs for that trial. Figure 1 below shows a schematic two-party representation of such an experimental setting. After many repeated trials the product of the PEFs from all the trials is used to estimate the probability of outcomes conditioned on the settings.

For the examples considered in Section 4, we will consider the canonical scenario of two measuring parties Alice and Bob each selecting respective binary measurement settings X and Y and recording respective binary outcomes A and B , which we refer to as the (2,2,2) Bell scenario. For now, we treat things in a general manner as is carried out in [14,15], modelling the trial settings for all parties and outcomes for all parties with single random variables Z and C , respectively, taking values from respective finite-cardinality sets $\mathcal{Z}$ and $\mathcal{C}$. When applied to the (2,2,2) Bell scenario, C comprises the ordered pair (A, B) and Z comprises the ordered pair (X, Y) .

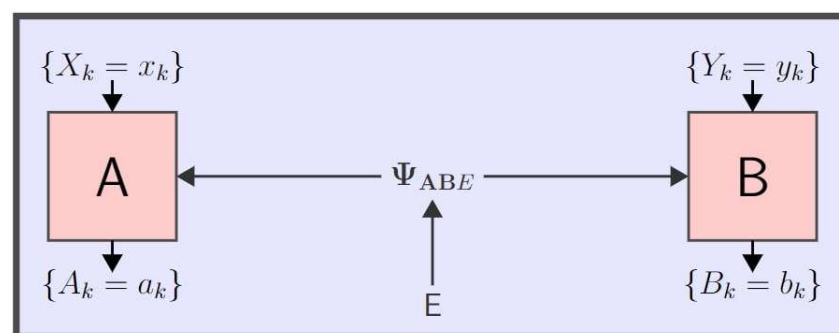


Figure 1. A schematic representation of the set-up for device-independent randomness generation in a two-party experiment. The outer rectangular box represents a secure location. The adversary E has perfect knowledge of the processes inside the secure location but cannot tamper with them. The state Ψ_{ABE} represents the resource shared between the two parties. X_k, Y_k are the trial inputs and A_k, B_k are the trial outcomes for the k th trial.

The results of a sequence of n time-ordered trials are represented by the sequences $\mathbf{C} = \{C_i\}_{i=1}^n, \mathbf{Z} = \{Z_i\}_{i=1}^n$; and, so, $(\mathbf{C}, \mathbf{Z})$ realises values $(\mathbf{c}, \mathbf{z}) \in \mathcal{C}^n \times \mathcal{Z}^n$, where $\mathcal{C}^n, \mathcal{Z}^n$

are the n -fold Cartesian products of $\mathcal{C}$, $\mathcal{Z}$. A PEF is then a real-valued function of $\mathbf{C}$ and $\mathbf{Z}$ satisfying certain conditions, while the product of PEFs from all trials will be a function of $\mathbf{C}$ and $\mathbf{Z}$. High values of the PEF product will correlate with low values of $\mathbb{P}(\mathbf{C}|\mathbf{Z})$, the conditional probability of the outcomes given the settings.

To define PEFs, we introduce the notion of a *trial model*: a set Π encompassing all joint probability distributions of settings and outcomes which are compatible with basic assumptions about the experiment. One important trial model that we consider is Π_Q , consisting of joint distributions of (C, Z) for which the conditional distribution of C conditioned on Z can be realised by a measurement on a quantum system. Here, we introduce the convention, used throughout, of using lower case Greek letters with random variables as arguments to denote distributions, i.e., $\mu(C, Z)$ and $\mu(C|Z)$ denote the joint distribution of (C, Z) and the conditional distribution of C given Z , respectively. Another important trial model is Π_{NS} (NS stands for “no-signalling”), consisting of distributions for which probabilities of measurement outcomes at one location are independent of measurement settings at the other distant locations. (This is more clearly understood in considering the Alice–Bob example, where one of the no-signalling conditions is that $\sum_b \mu(A = a, B = b|X = x, Y = y) = \sum_b \mu(A = a, B = b|X = x, Y = y')$ for all a, b, x and $y \neq y'$.) A third important trial model is the set Π_L of distributions for which the conditional distribution of outcomes conditioned on settings are *local*, which means they can be expressed as convex mixtures of local deterministic behaviours. In the bipartite setting, the conditional distribution $\mu_{LD,\lambda}(A, B|X, Y)$, also referred to as a behaviour, is local deterministic if $\mu_{LD,\lambda}(A = a, B = b|X = x, Y = y) = \llbracket a = f(x, \lambda) \rrbracket \llbracket b = g(y, \lambda) \rrbracket$ (where the notation $\llbracket \cdot \rrbracket$ represents the function that evaluates to 1 if the condition within holds, 0 otherwise). In words, the outcomes are functions of the local settings and the local hidden variable λ which can be understood to be a list of outcomes for all possible settings. A formal definition involving more parties and an arbitrary (albeit same) number of outcomes and settings for each party can be found in (48). The sets Π_L , Π_Q and Π_{NS} satisfy the following strict inclusions:

$$\Pi_L \subsetneq \Pi_Q \subsetneq \Pi_{NS}.$$

Certain distributions in Π_Q and Π_{NS} violate a Bell inequality and are known to contain randomness; they are contained in $\Pi_Q \setminus \Pi_L$ and $\Pi_{NS} \setminus \Pi_L$, respectively. It is precisely the inability to decompose such distributions into deterministic ones, as in Π_L , that implies the presence of randomness. The objective of the PEF approach is to quantify the randomness contained in such distributions. As trial models specify the joint distribution $\mu(C, Z)$, and for the above examples of trial models we gave only the conditional distributions $\mu(C|Z)$, one must also specify the marginal distribution of the settings $\mu(Z)$. For the discussions of Π_Q and Π_{NS} in subsequent sections, any fixed distribution satisfying $\mu(Z = z) > 0$ for all $z \in \mathcal{Z}$ is permitted. An example of a fixed settings distribution is the equiprobable distribution $\text{Unif}(Z)$ defined as $\text{Unif}(z) = 1/|\mathcal{Z}|$ for all $z \in \mathcal{Z}$.

As a discrete probability distribution is effectively an ordered list of numbers in $[0, 1]$ (the probabilities), trial models are always subsets of $\mathbb{R}^N$, where N is fixed by the cardinality of $\mathcal{C}$ and $\mathcal{Z}$. This enables us to use a geometric approach to study these sets, which prove to be invaluable for some arguments.

We can now define PEFs. We use the notation $\mathbb{E}_\mu[\dots]$ and $\mathbb{P}_\mu(\dots)$ to denote expectation and probability, respectively, with respect to a distribution μ ; and for the sake of notational concision we omit commas in distributions or functions of more than one random variable, for instance, $\mu(CZ)$ and $f(CZ)$ must be understood to mean $\mu(C, Z)$ and $f(C, Z)$.

Definition 1 (Probability Estimation Factor). *A probability estimation factor (PEF) with power $\beta > 0$ for the model of distributions Π is a function $F: \mathcal{C} \times \mathcal{Z} \rightarrow \mathbb{R}^+$ of the random variables (C, Z) such that for all $\sigma(CZ) \in \Pi$, $\mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\beta] \leq 1$ holds.*

In the expression above, $\sigma(C|Z)$ denotes a random variable that is a function of the random variables C and Z : $\sigma(C|Z)$ is the random variable that assumes the standard conditional probability (according to σ) of C taking the value c conditioned on Z taking the value z ; it is assigned the value zero if the probability $\sigma(Z = z)$ is zero. The parameter β can be any positive real value. We then note that the constant PEF $F(cz) = 1$ for all $(c, z) \in \mathcal{C} \times \mathcal{Z}$ is a valid PEF for any choice of $\beta > 0$. We will notice in the subsequent sections, however, that the parameter does have an effect on the method employed for choosing useful PEFs for the purpose of randomness certification; and in practice we choose the value of β that corresponds to the maximum randomness certification.

Prior to defining a PEF we introduced the notion of a trial model. For the application of probability estimation to the outcomes of an experiment, which is a sequence of n time-ordered trials, we introduce the notion of an *experiment model*: it is a set Θ constraining the joint distribution of $\mathbf{C}, \mathbf{Z}$ and E , constructed as a chain of individual trial models Π ; it consists of joint distributions $\mu(\mathbf{CZ}|E = e)$ conditioned on the event $\{E = e\}$, where E is the random variable denoting the adversary's side information and realising values e from the finite set $\mathcal{E}$. It satisfies the following two assumptions:

$$\begin{aligned} \mu(C_{i+1}Z_{i+1}|\mathbf{C}_{\leq i} = \mathbf{c}_{\leq i}, \mathbf{Z}_{\leq i} = \mathbf{z}_{\leq i}, E = e) &\in \Pi, \forall \mathbf{c}_{\leq i} \in \mathcal{C}^i, \mathbf{z}_{\leq i} \in \mathcal{Z}^i, e \in \mathcal{E}, \\ \mu(Z_{i+1}, \mathbf{C}_{\leq i}|\mathbf{Z}_{\leq i} = \mathbf{z}_{\leq i}, E = e) &= \mu(Z_{i+1}|E = e)\mu(\mathbf{C}_{\leq i}|\mathbf{Z}_{\leq i} = \mathbf{z}_{\leq i}, E = e), \forall e \in \mathcal{E}. \end{aligned} \quad (1)$$

In (1), $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$ denote the outcomes and measurement settings for the first $i \in [n]$ trials, where $[n] := \{1, 2, \dots, n\}$, with $\mathbf{c}_{\leq i}, \mathbf{z}_{\leq i}$ denoting their respective realisations. The random variables C_{i+1}, Z_{i+1} are the outcomes and settings for the $(i + 1)$ 'th trial. The first condition in (1) formalises the assumption that the (joint) probability of the $(i + 1)$ 'th outcome and setting, conditioned on the outcomes and settings for the first i trials and each realised value $E = e$ of the adversary's side information, belongs to the $(i + 1)$ 'th trial model, i.e., it is compatible with the conditions dictated by the trial model. The second condition states that for each $E = e$ the setting for the *next* trial is independent of the outcomes and settings of the *past and present* trials. Our second condition is a stronger assumption than the corresponding assumption given in [14], which is as follows: the joint distribution μ of $\mathbf{CZ}E$ is such that Z_{i+1} is independent of $\mathbf{C}_{\leq i}$ conditionally on both $\mathbf{Z}_{\leq i}$ and E . It is a straightforward exercise to check that our stronger assumption implies the one stated in [14]. While the weaker assumption is sufficient for the following result, we find the stronger assumption operationally clearer as an assumption that the future settings are independent of "everything in the past" for each realisation of e .

For the rest of the paper we adopt the abbreviated notation of $\mu_y(X)$ for $\mu(X|Y = y)$. The following theorem, appearing as Theorem 9 in Appendix C in [14], formalises the central idea behind the framework of probability estimation. We include a proof for this theorem in Appendix A.1.1 for completeness.

Theorem 1. Suppose $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ is a distribution of $\mathbf{CZ}E$ such that $\mu_e(\mathbf{CZ}) \in \Theta$ for each $e \in \mathcal{E}$. Then, for fixed $\beta, \epsilon > 0$

$$\mathbb{P}_{\mu_e} \left(\mu_e(\mathbf{C}|\mathbf{Z}) \geq \left(\epsilon \prod_{i=1}^n F_i(C_i Z_i) \right)^{-1/\beta} \right) \leq \epsilon \quad (2)$$

holds for each $e \in \mathcal{E}$, where $F_i(C_i Z_i)$ is the probability estimation factor for the i 'th trial.

Proof. See Appendix A.1.1. $\square$

The distinguishing feature of the framework of probability estimation is the direct estimation of $\mu_e(\mathbf{C}|\mathbf{Z})$ for each $e \in \mathcal{E}$ by constructing PEFs $F_i(C_i Z_i)$ and accumulating them trial-wise in a multiplicative fashion. For a fixed error bound $\epsilon > 0$ and the power parameter $\beta > 0$, the term $(\epsilon \prod_{i=1}^n F_i(C_i Z_i))^{-1/\beta}$ serves as an estimate for $\mu_e(\mathbf{C}|\mathbf{Z})$. It is

important to note that PEFs are functions of only the measurement outcomes and settings and not of the side information held by the adversary to which we do not have access. For a large value of n —the number of trials—the trial-wise product $\prod_{i=1}^n F_i(C_i Z_i)$ will be large if the experiment is well-calibrated and run properly. For the purpose of randomness generation the inequality (2) in Theorem 1 can then be understood, intuitively, as follows: Since the trial-wise product $\prod_{i=1}^n F_i(C_i Z_i)$ of the PEFs is large and so, for fixed $\epsilon, \beta > 0$, the quantity $(\epsilon \prod_{i=1}^n F_i(C_i Z_i))^{-1/\beta}$ is small, for each $e \in \mathcal{E}$ there is a very small probability (denoted by the outer probability $\mathbb{P}_{\mu_e}(\cdot)$) that the conditional probability of the sequence of outcomes $\mathbf{C}$ conditioned on the sequence of settings $\mathbf{Z}$ (denoted by $\mu_e(\mathbf{C}|\mathbf{Z})$) is more than a small value. This translates to the measurement outcomes $\mathbf{C}$ being unpredictably random for a given $\mathbf{z}e$. Since this string of experimental outcomes is unpredictable even given the adversary's side information, it can be used as a source of certifiable randomness. We stress that in the method of probability estimation the estimates on the conditional probability of measurement outcomes given the settings choices and side-information depend solely on the experimental data.

Conventional methods of randomness extraction, however, involve obtaining a lower bound on the smooth conditional min-entropy which quantifies the amount of raw randomness from a source. The lower bound then goes as one of the parameters in extractor functions to extract near-uniform random bits. It is therefore useful to translate the bound in (2) into a statement about the smooth conditional min-entropy with respect to an adversary.

We motivate and introduce conditional min-entropy as follows. An adversary's goal is to predict C . Conditioned on a particular realisation of the settings sequence $\mathbf{z} \in \mathcal{Z}^n$ and side information $e \in \mathcal{E}$, one can measure the “predictability” of the sequence of outcomes $\mathbf{C}$ with the following maximum probability:

$$\max_{\mathbf{c} \in \mathcal{C}^n} \mu(\mathbf{c}|\mathbf{z}e).$$

It quantifies the best guess of the adversary. The $\mathbf{z}e$ -conditional min-entropy of $\mathbf{C}$, corresponding to that particular realisation $\mathbf{z}e \in \mathcal{Z}^n \times \mathcal{E}$, is the following negative logarithm:

$$\mathbb{H}_{\infty, \mu}(\mathbf{C}|\mathbf{z}e) := -\log_2 \left(\max_{\mathbf{c} \in \mathcal{C}^n} \mu(\mathbf{c}|\mathbf{z}e) \right).$$

The subscript μ in the notation $\mathbb{H}_{\infty, \mu}(\cdot \cdot \cdot)$ refers to the distribution $\mu(\mathbf{C}|\mathbf{Z}E)$. The average $\mathbf{Z}E$ -conditional min-entropy is then defined as follows:

$$\mathbb{H}_{\infty, \mu}^{\text{avg}}(\mathbf{C}|\mathbf{Z}E) := -\log_2 \left[\sum_{\mathbf{z}e \in \mathcal{Z}^n \times \mathcal{E}} \left(\max_{\mathbf{c} \in \mathcal{C}^n} \mu(\mathbf{c}|\mathbf{z}e) \right) \mu(\mathbf{z}e) \right].$$

But, information-theoretic security of cryptographic protocols take into account a more realistic measure of average $\mathbf{Z}E$ -conditional min-entropy which involves a smoothing parameter ϵ , a type of error bound, and is known as the ϵ -smooth average $\mathbf{Z}E$ -conditional min-entropy. This quantity is useful for our scenario in which the probability distribution is not known exactly and its characteristics can only be inferred from observed data, which introduces the possibility of error. It is defined as follows.

Definition 2 (Smooth Average Conditional Min-Entropy). *For a distribution $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ of $\mathbf{C}, \mathbf{Z}, E$ the set $\mathcal{B}^\epsilon(\mu)$ of distributions of $\mathbf{C}, \mathbf{Z}, E$ is defined as*

$$\mathcal{B}^\epsilon(\mu) := \{\sigma: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1] \mid d_{\text{TV}}(\sigma, \mu) \leq \epsilon\}, \quad (3)$$

where $\epsilon \in (0, 1)$ and $d_{\text{TV}}(\sigma, \mu)$ is the total variation distance between σ and μ defined as

$$d_{\text{TV}}(\sigma, \mu) := \frac{1}{2} \sum_{\mathbf{c} \mathbf{z} e \in \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E}} |\mu(\mathbf{c} \mathbf{z} e) - \sigma(\mathbf{c} \mathbf{z} e)|. \quad (4)$$

The ϵ -smooth average $\mathbf{ZE}$ -conditional min-entropy is then defined as follows.

$$\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(\mathbf{C}|\mathbf{ZE}) := \max_{\sigma \in \mathcal{B}^{\epsilon}(\mu)} \left[-\log_2 \left[\sum_{\mathbf{ze} \in \mathcal{Z}^n \times \mathcal{E}} \left(\max_{\mathbf{c} \in \mathcal{C}^n} \sigma(\mathbf{c}|\mathbf{ze}) \right) \sigma(\mathbf{ze}) \right] \right]. \quad (5)$$

The lower bound obtained on this quantity goes as one of the inputs to extractor functions in randomness extraction, whose purpose is to convert random functions with uneven distributions into shorter, close to uniformly distributed bit strings. We note that alternative definitions of ϵ -smooth conditional min-entropy can be used, for instance, the ϵ -smooth *worst-case* conditional min-entropy of [19]. A known result from the literature, proven in Proposition A1 in Appendix E, justifies our usage of the ϵ -smooth average conditional min-entropy without having to be concerned with the stricter ϵ -smooth worst-case conditional min-entropy (defined in (A30)): specifically, the two quantities converge to one another in the asymptotic limit.

The result obtained from Theorem 1 can be translated into a result on smooth average conditional min-entropy formalised in Theorem 2 below. This theorem appears as Theorem 1 in [14]. We include a proof for this theorem in Appendix A.1.2 for completeness. In the notation of ϵ -smooth average $\mathbf{ZE}$ -conditional min-entropy in (7), the semicolon followed by S denotes that this information-quantity is assessed with respect to the distribution μ after conditioning on the occurrence of the event S defined in the statement of Theorem 2. It pertains to an abort criterion. The protocol succeeds only if the product of the trial-wise PEFs exceeds some threshold value, otherwise it is aborted. So we want to establish the lower bound for smooth conditional min-entropy conditioned on the event that the protocol succeeds, because it is precisely this scenario in which we extract randomness. Since a completely predictable local distribution can always have a chance of passing the protocol, however minuscule (in the order of $(3/4)^n$, where the number of trials n often goes up to millions)—and $\mu(\mathbf{c}|\mathbf{z})$ will equal 1 in this case—it is necessary to assume a small but positive lower bound on the probability of not aborting to derive a useful min-entropy bound. This can be thought of as another type of error parameter. The assumed lower bound for the probability of success of the protocol is κ .

Theorem 2. Let μ be a distribution $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ of $\mathbf{C}, \mathbf{Z}, E$ such that, for each $e \in \mathcal{E}$, the following holds for every $\epsilon \in (0, 1)$:

$$\mathbb{P}_{\mu_e} \left(\mu_e(\mathbf{C}|\mathbf{Z}) \leq \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \right) \geq 1 - \epsilon, \quad (6)$$

where F_i is a PEF with power β for the i 'th trial. For a fixed choice of $\epsilon \in (0, 1)$ and $p \geq |\mathcal{C}|^{-n}$, define the event $S := \left\{ \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \leq p \right\}$. Then, if κ satisfies $0 < \kappa \leq \mathbb{P}_{\mu}(S)$, the following holds:

$$\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon/\kappa}(\mathbf{C}|\mathbf{ZE}; S) \geq \log_2(\kappa) - \log_2(p) \quad (7)$$

Proof. See Appendix A.1.2. $\square$

Under the same conditions of Theorem 2, the main result (7) admits a minor reformulation as follows. This is the formulation that aligns with the statement of Theorem 1 in [14].

Corollary 1. Let $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ be a distribution of **CZE** and F be a PEF with power β such that (6) holds for each $e \in \mathcal{E}$. For a fixed choice of $\epsilon \in (0, 1)$, $p \geq |\mathcal{C}|^{-n}$ and positive $\kappa \leq \mathbb{P}_\mu(S)$ where $S = \{(\epsilon \prod_{i=1}^n F_i)^{-1/\beta} \leq p\}$, we have

$$\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(\mathbf{C}|\mathbf{Z}\mathbf{E}; S) \geq \left(1 + \frac{1}{\beta}\right) \log_2(\kappa) - \log_2(p). \quad (8)$$

Proof. Use Theorem 2 with $\epsilon' = \kappa\epsilon$, $p' = p/\kappa^{1/\beta}$ and $\kappa' = \kappa$, noting that, since $0 < \kappa \leq 1$ and $\beta > 0$ hold, we have $\epsilon' \in (0, 1)$ and $p' \geq |\mathcal{C}|^{-n}$ as required for invoking the theorem. Then, notice the corresponding event $S' = \{(\epsilon' \prod_{i=1}^n F_i)^{-1/\beta} \leq p'\}$ aligns with the event S . $\square$

The above results hold when we consider distributions $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E}^n \rightarrow [0, 1]$ of **CZE**, i.e., where the side information is structured as a sequence of random variables. The proof remains the same with the exception that we condition on an arbitrary sequence of realisation $\mathbf{e} \in \mathcal{E}^n$ of $\mathbf{E}$. We consider this scenario in Section 3 where we define an IID attack from the adversary.

Theorem 1 does not indicate how to find PEFs. One way to find useful PEFs is to first notice that the success criterion of the protocol is the event S that the inequality $(\epsilon \prod_{i=1}^n F_i)^{-1/\beta} \leq p$ holds, which can be equivalently expressed as

$$\sum_{i=1}^n \log_2(F_i)/\beta + \log_2(\epsilon)/\beta \geq -\log_2(p), \quad (9)$$

where ϵ, β and p are pre-determined quantities to be chosen in advance of running the protocol. Then, considering an anticipated trial distribution $\rho(CZ)$ based on observed results and calibrations from previous trials, in the limit of sufficiently large n the difference between the term on the left hand side of (9) (which consists of the trial-wise sum of (base-2) logarithm of PEFs) and $n\mathbb{E}_\rho[\log_2(F(CZ))/\beta]$ will be either greater or less than zero with roughly equal probability. This follows from the Central Limit Theorem if the distribution remains roughly stable from trial to trial. Since it is desirable to have the largest value of $-\log_2(p)$ possible, one can then perform the following constrained maximisation using any convex programming software owing to the concavity of the objective function and the linearity of the constraints.

$$\begin{aligned} &\text{Maximise: } \mathbb{E}_\rho[(n \log_2(F(CZ)) + \log_2(\epsilon))/\beta] \\ &\text{Subject to: } \mathbb{E}_\nu[F(CZ)\nu(C|Z)^\beta] \leq 1, \text{ for all } \nu(CZ) \in \Pi, \\ &\quad F(cz) \geq 0, \text{ for all } (c, z) \in \mathcal{C} \times \mathcal{Z} \end{aligned} \quad (10)$$

Since n, ϵ and β are fixed, it is sufficient to maximise $\mathbb{E}_\rho[\log_2(F(CZ))]$ subject to the same constraints. In practice, one can consider a range of values of β and perform the constrained maximisation with the objective $\mathbb{E}_\rho[\log_2(F(CZ))]$, then plug in the maximum value in the expression $\mathbb{E}_\rho[(n \log_2(F(CZ)) + \log_2(\epsilon))/\beta]$ and obtain a plot with respect to the considered range of β values (see, for example, Figure 2 in [16]; a similar pattern is observed in Figure 2 in Section 2).

The following lemma (from [14], see Lemma 15)—for which we provide a more direct proof—enables us to restrict the satisfiability constraints of the optimisation routine in (10) to the extremal distributions of the model Π under the condition that the model is convex and closed. So, the first line of constraints in (10) can be replaced with $\mathbb{E}_\nu[F(CZ)\nu(C|Z)^\beta] \leq 1, \forall \nu(CZ) \in \Pi_{\text{extr}}$, where Π_{extr} is the set of extremal distributions of Π . If the model Π is not convex and closed, we take its convex closure. In words, the lemma states that, if $F(CZ)$ is a PEF with power $\beta > 0$ for the distributions $\sigma_1(CZ)$ and $\sigma_2(CZ)$, then it is a PEF with the same power for all distributions that can be expressed as a convex combination of σ_1 and σ_2 .

Lemma 1. For distributions $\sigma_i(CZ) \in \Pi$ satisfying $\mathbb{E}_{\sigma_i}[F(CZ)\sigma_i(C|Z)^\beta] \leq 1$, for $i = 1, 2$, if $\sigma(CZ) \in \Pi$ is expressible as $\sigma(CZ) = \lambda\sigma_1(CZ) + (1 - \lambda)\sigma_2(CZ)$ for $\lambda \in [0, 1]$, then it satisfies $\mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\beta] \leq 1$.

Proof. For z such that $\sigma_1(z), \sigma_2(z) > 0$, we have $\sigma(z) > 0$ as well and, from $\sigma(CZ) = \lambda\sigma_1(CZ) + (1 - \lambda)\sigma_2(CZ)$, straightforward algebra shows that $\sigma(c|z) = \delta\sigma_1(c|z) + (1 - \delta)\sigma_2(c|z)$ for any $(c, z) \in \mathcal{C} \times \mathcal{Z}$, where $\delta = \lambda\sigma_1(z)/\sigma(z) \in [0, 1]$. Since, for $\alpha > 1$, x^α is convex for $x \geq 0$, we can write

$$\begin{aligned} \sigma(c|z)^{1+\beta} &\leq \delta\sigma_1(c|z)^{1+\beta} + (1 - \delta)\sigma_2(c|z)^{1+\beta} \\ \Rightarrow \sigma(c|z)^{1+\beta}\sigma(z) &\leq \lambda\sigma_1(c|z)^{1+\beta}\sigma_1(z) + (1 - \lambda)\sigma_2(c|z)^{1+\beta}\sigma_2(z). \end{aligned} \quad (11)$$

Turning to cases where $\sigma_1(z)$ and/or $\sigma_2(z)$ may equal zero, we can also demonstrate (11) under the convention of taking $\sigma_i(c|z)$ to be zero when $\sigma_i(z) = 0$. Then, the inequality holds as an equality when $\sigma_1(z) = \sigma_2(z) = 0$ (which implies $\sigma(z) = 0$ as well); for $0 = \sigma_2(z) < \sigma_1(z)$ one can verify (11) after noting $\sigma(cz) = \lambda\sigma_1(cz)$ and $\sigma(z) = \lambda\sigma_1(z)$, and the $0 = \sigma_1(z) < \sigma_2(z)$ case follows symmetrically. Now, multiplying both sides of (11) by $F(cz)$ and summing over $(c, z) \in \mathcal{C} \times \mathcal{Z}$ gives

$$\begin{aligned} \sum_{c,z} F(cz)\sigma(c|z)^{1+\beta}\sigma(z) &\leq \lambda \sum_{c,z} F(cz)\sigma_1(c|z)^{1+\beta}\sigma_1(z) + (1 - \lambda) \sum_{c,z} F(cz)\sigma_2(c|z)^{1+\beta}\sigma_2(z) \\ \Rightarrow \mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\beta] &\leq \lambda \mathbb{E}_{\sigma_1}[F(CZ)\sigma_1(C|Z)^\beta] + (1 - \lambda) \mathbb{E}_{\sigma_2}[F(CZ)\sigma_2(C|Z)^\beta] \\ &\leq \lambda + (1 - \lambda) = 1. \end{aligned}$$

□

We remark that the result of Lemma 1 can also be obtained through specialisation of known quantum results to classical distributions; however, this requires a more technical argument with additional machinery. To elaborate, the proof for Lemma 1 involves showing the joint convexity of $\sigma(C|Z)^{1+\beta}\sigma(Z)$ which can be seen as a special case of the joint convexity of sandwiched Rényi powers. To be more specific, it arises as a special case of the joint convexity of $e^{\beta D_{1+\beta}(\sigma||\omega)}$ for $\beta > 0$ when the distribution $\omega(CZ)$ is taken to be $\omega(cz) = \sigma(z)/|\mathcal{C}|$, $\forall (c, z) \in \mathcal{C} \times \mathcal{Z}$. Notice that $D_{1+\beta}(\sigma||\omega)$ is the (classical) Rényi divergence of order $(1 + \beta) \in (1, \infty)$ of $\sigma(CZ)$ with respect to $\omega(CZ)$. The functional $e^{D_{1+\beta}(\sigma||\omega)}$ can also be seen as a specialisation (to classical states) of the same functional, defined in terms of (quantum) density states σ and ω , whose joint convexity was proven in proposition 3 of [20] with an extended technical argument.

3. Asymptotic Performance

The results of the previous section give us a method for certifying randomness. In this section, we assess the asymptotic performance of the method. Our figure of merit is the amount of randomness certified per trial, as measured by the average conditional min-entropy divided by the number of trials n . We will see in this section that the PEF method is asymptotically optimal, in the following sense: given a fixed observed distribution, the PEF method can asymptotically certify an amount of per-trial conditional min-entropy that is equal to the actual per-trial conditional min-entropy generated by an adversary replicating the observed distribution with as little randomness as possible.

To elaborate on this, consider that the adversary's goal is to minimise the following quantity:

$$\frac{1}{n} \mathbb{H}_{\infty, \mu}^{\text{avg}}(\mathbf{C}|\mathbf{Z}\mathbf{E}).$$

We assume that the adversary has complete knowledge of the distribution μ , and can have access to not just the realised value of E but also the realised value of $\mathbf{Z}$ in guessing $\mathbf{C}$. This access to $\mathbf{Z}$ aligns with the paradigm, as discussed in [11], of “using public (settings)

randomness to generate private (outcome) randomness”. The adversary is constrained, however, in that the statistics when marginalised over E must appear to be consistent with an expected observed trial distribution $\rho(CZ)$ for the protocol to not abort. Technically, all that is necessary for the protocol to pass is that the observed product of the PEFs must exceed some threshold value chosen by the experimenter—which could be possible with high probability with many different distributions μ —but, as the experimenter’s threshold value will likely be chosen based on a full behaviour that they expect to observe, we study attacks that match the expected observed trial distribution exactly. We will find attacks meeting this criterion that are asymptotically optimal for minimising the conditional min-entropy.

Given an expected observed distribution, how can the adversary generate observed statistics consistent with it while yielding as little randomness as possible? She can employ a strategy of preparing multiple different states to be measured that will yield different distributions, each one consistent with the trial model Π , whose convex mixture is equal to the observed distribution. If she has an auxiliary random variable E realising values from the finite-cardinality set $\mathcal{E}$ and recording which state was prepared on which trial, she can predict better the outcome conditioned on her side information $E = e$, in conjunction with the settings Z . Indeed, some of her e -conditional distributions could be deterministic—specifically, the product of a fixed settings distribution and a deterministic behaviour (conditional distribution of the outcomes conditioned on settings), in which case she does not yield any randomness to Alice and Bob on a trial where E takes that value. But, if the overall observed statistics are nonlocal, then she is forced to prepare at least some states that contain randomness even conditioned on e ; this, in essence, is because the information that she possesses with E is a local hidden variable.

3.1. I.I.D. Attacks

Given a convex decomposition of the observed distribution, the adversary’s simplest form of an attack is to select e from some finite-cardinality set $\mathcal{E}$ in an i.i.d manner on each trial according to the distribution that recovers the observed distribution $\rho(CZ)$. A more general attack would allow her to use memory of earlier trials but we will see later that, asymptotically, this does not yield meaningful improvement.

Operationally, we do not like to think of the adversary accessing the devices in between trials to provide a choice of e_i for each trial. Instead, one can imagine her randomly sampling from the distribution of $\mathbf{E}$ for all trials, coming up with a choice $\mathbf{e}$ that encodes all the choices of e_i for each trial and then supplying this choice to the measured system, in advance, to determine its behaviour in each trial. She keeps a record of $\mathbf{e}$ to help her predict C later. Through this sampling process there is a small chance that she will sample an atypical “bad” $\mathbf{e}$ that results in statistics deviating from the observed distribution but the probability that her $\mathbf{e}$ is typical is asymptotically high. Our figure of merit for the adversary now is:

$$\frac{1}{n} \mathbb{H}_{\infty, \mu}^{\text{avg}}(\mathbf{C} | \mathbf{Z} \mathbf{E}),$$

which she wants to minimise with a distribution that, marginalised over $\mathbf{E}$, is consistent with i.i.d sampling from an expected observed distribution ρ . We formally define the set $\Sigma_{\mathbf{E}}^{\rho}$ of distributions $\omega: \mathcal{C} \times \mathcal{Z} \times \mathcal{E} \rightarrow [0, 1]$ of C, Z, E mimicking ρ through such a convex decomposition as follows, where e is shorthand for the event $\{E = e\}$:

$$\Sigma_{\mathbf{E}}^{\rho} := \left\{ \omega(CZE): \omega(CZ|e) \in \Pi \forall e \in \mathcal{E}, \sum_{e \in \mathcal{E}} \omega(CZ|e) \omega(e) = \rho(CZ) \right\}. \quad (12)$$

Then, an IID attack can be defined as follows.

Definition 3 (IID Attack). Given a distribution $\omega(\mathbf{CZE}) \in \Sigma_E^0$, we define an IID attack (with ω) to be the distribution ϕ consisting of n independent and identical realisations of random variables C_i, Z_i, E_i distributed according to ω ; i.e., the joint distribution of the sequence of random variables $\mathbf{C}, \mathbf{Z}, \mathbf{E}$ is $\phi: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E}^n \rightarrow [0, 1]$ such that $\phi(\mathbf{CZE}) = \prod_{i=1}^n \omega(C_i Z_i E_i)$.

As mentioned earlier, the adversary randomly samples from the distribution of $\mathbf{E}$ which represents their knowledge of all trials; $\mathbf{e} \equiv (e_1, e_2, \dots, e_n) \in \mathcal{E}^n$ encodes the individual choices e_i for trial $i \in \{1, 2, \dots, n\}$. The IID attack satisfies the two assumptions of the experiment model discussed earlier (see (1) and the short discussion that follows immediately). Namely, the (joint) probability of the $(i+1)$ 'th trial outcome and input setting, conditioned on each realisation of the outcomes and settings for the first i trials and each realisation $\mathbf{e} \in \mathcal{E}^n$ of the side information, satisfies the conditions of the trial model; and, conditioned on each $\mathbf{e} \in \mathcal{E}^n$, the settings for the $(i+1)$ 'th trial are (unconditionally) independent of the outcomes and settings of the past and present trials (i.e., the first i trials). This is formally stated and proved in Lemma 2 below.

Lemma 2. The IID attack as defined in Definition 3 satisfies the following conditions.

$$\phi(C_{i+1} Z_{i+1} | \mathbf{c}_{\leq i} \mathbf{z}_{\leq i} \mathbf{e}) \in \Pi, \forall \mathbf{c}_{\leq i} \in \mathcal{C}^i, \mathbf{z}_{\leq i} \in \mathcal{Z}^i, \mathbf{e} \in \mathcal{E}^n \quad (13)$$

$$\phi(Z_{i+1} \mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}) = \phi(Z_{i+1} | \mathbf{e}) \phi(\mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}), \forall \mathbf{e} \in \mathcal{E}^n \quad (14)$$

Proof. Consider the distribution $\phi(\mathbf{CZ} | \mathbf{e})$ conditioned on a realisation $\mathbf{E} = \mathbf{e}$, where $\phi(\mathbf{CZE}) = \prod_{i=1}^n \omega(C_i Z_i E_i)$. Notice that $\phi(\mathbf{CZ} | \mathbf{e}) = \prod_{i=1}^n \omega(C_i Z_i | e_i)$. Marginalising over the random variables $C_{i+2}, C_{i+3}, \dots, C_n, Z_{i+2}, Z_{i+3}, \dots, Z_n$ we obtain:

$$\phi(C_{i+1} Z_{i+1} \mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}) = \prod_{j=1}^{i+1} \omega(C_j Z_j | e_j) \quad (15)$$

Corresponding to a particular realisation $\mathbf{c}_{\leq i} \in \mathcal{C}^i, \mathbf{z}_{\leq i} \in \mathcal{Z}^i$, we then have $\phi(C_{i+1} Z_{i+1} \mathbf{c}_{\leq i} \mathbf{z}_{\leq i} | \mathbf{e}) = \omega(C_{i+1} Z_{i+1} | e_{i+1}) \prod_{j=1}^i \omega(c_j z_j | e_j)$; and, since $\phi(\mathbf{c}_{\leq i} \mathbf{z}_{\leq i} | \mathbf{e}) = \prod_{j=1}^i \omega(c_j z_j | e_j)$, we have

$$\frac{\phi(C_{i+1} Z_{i+1} \mathbf{c}_{\leq i} \mathbf{z}_{\leq i} | \mathbf{e})}{\phi(\mathbf{c}_{\leq i} \mathbf{z}_{\leq i} | \mathbf{e})} = \phi(C_{i+1} Z_{i+1} | \mathbf{c}_{\leq i} \mathbf{z}_{\leq i} \mathbf{e}) = \omega(C_{i+1} Z_{i+1} | e_{i+1}). \quad (16)$$

$\omega(C_{i+1} Z_{i+1} | e_{i+1})$ belongs to the set Π for all values of $e_{i+1} \in \mathcal{E}$ (by construction of the set Σ_E , see (12)). Since (16) is true for all realisations $\mathbf{c}_{\leq i} \in \mathcal{C}^i, \mathbf{z}_{\leq i} \in \mathcal{Z}^i, \mathbf{e} \in \mathcal{E}^n$ we conclude (13) holds. Next, marginalising (15) over C_{i+1} we have:

$$\phi(Z_{i+1} \mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}) = \omega(Z_{i+1} | e_{i+1}) \prod_{j=1}^i \omega(C_j Z_j | e_j) = \phi(Z_{i+1} | \mathbf{e}) \phi(\mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}) \quad (17)$$

In (17), $\omega(Z_{i+1} | e_{i+1}) = \phi(Z_{i+1} | \mathbf{e})$ can be observed by marginalising (15) over the random variables $C_1, \dots, C_i, C_{i+1}, Z_1, \dots, Z_i$ and $\phi(\mathbf{C}_{\leq i} \mathbf{Z}_{\leq i} | \mathbf{e}) = \prod_{j=1}^i \omega(C_j Z_j | e_j)$ (from marginalising (15) over C_{i+1}, Z_{i+1}); (17) is true for all $\mathbf{e} \in \mathcal{E}^n$; hence, we conclude (14). $\square$

Next, the adversary would like to implement an attack that “generates as little randomness as possible”. One measure of the randomness is the conditional Shannon entropy of the outcomes C conditioned on the inputs Z and the side information E .

Definition 4 (Conditional Shannon Entropy). For a distribution $\mu: \mathcal{C} \times \mathcal{Z} \times \mathcal{E} \rightarrow [0, 1]$ of C, Z, E the conditional Shannon entropy of the outcomes C conditioned on the settings Z and the side information E is defined as

$$\mathbb{H}_\mu(C|ZE) = - \sum_{cze} \log_2 \mu(c|ze) \mu(cze) = \mathbb{E}_\mu[-\log_2 \mu(C|ZE)]. \quad (18)$$

The Greek letter μ in the subscript of $\mathbb{H}_\mu(\cdot|\cdot)$ refers to the distribution $\mu(CZE)$ with respect to which the conditional Shannon entropy is defined.

Theorem 3 below shows that $H_\omega(C|ZE)$ is an asymptotic upper bound on the per-trial conditional min-entropy that the adversary generates with an IID attack employing a trial distribution ω that is consistent with the observed distribution ρ . This result was discussed but not demonstrated explicitly in [15]. The proof of Theorem 3 involves one of the fundamental technical tools from information theory, the (classical) asymptotic equipartition property (AEP), or equivalently the notion of typical sequences which has the weak law of large numbers at its core.

Suppose μ , the distribution of all trials, is obtained as n i.i.d. copies of a single-trial distribution ω . Then, for $\epsilon_a \in (0, 1)$, $\delta \geq 0$ there exists $N(\epsilon_a, \delta)$ such that $n \geq N(\epsilon_a, \delta)$ ensures $\mathbb{E}_{\mu(ZE)}[\mathbb{P}_{\mu(C|ZE)}(\mu(C|ZE) \geq \gamma)] \geq 1 - \epsilon_a$, where $\gamma = 2^{-nH_\omega(C|ZE) - n\delta}$ and $H_\omega(C|ZE)$ is the conditional Shannon entropy. We refer to this as the AEP condition; it holds by a conditional form of the classical AEP (see, for instance, Section 14.6 in [21]). The set $B^{\epsilon_s}(\mu)$ of distributions of C, Z, E that are within a TV distance of ϵ_s from μ and the sets $\mathcal{A}_{ze}$ are as defined below:

$$\mathcal{B}^{\epsilon_s}(\mu) := \{\sigma: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E}^n \rightarrow [0, 1] \mid d_{TV}(\mu, \sigma) \leq \epsilon_s\}, \quad (19)$$

$$\mathcal{A}_{ze} := \{c \in \mathcal{C}^n \mid \mu(c|ze) \geq \gamma\}, \quad (20)$$

where $\mathcal{A}_{ze}$ is defined for any ze for which $\mu(ze) > 0$. Note that the case $\epsilon_s = 0$ reduces to a bound on the standard (non-smooth) average conditional min-entropy. We now state the result as follows.

Theorem 3. Let μ be an IID attack with ω . For $\epsilon_s \geq 0$, $\epsilon_a, \delta > 0$ and $\epsilon_a + 2\epsilon_s < 1$, there exists $N(\epsilon_a, \epsilon_s, \delta)$ such that for $n \geq N(\epsilon_a, \epsilon_s, \delta)$

$$\frac{1}{n} \mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon_s}(C|ZE) \leq \mathbb{H}_\mu(C|ZE) + \frac{1}{n} \log_2 \frac{1}{1 - \epsilon_a - 2\epsilon_s} + \delta. \quad (21)$$

Proof. Throughout, we follow the convention that $\sigma(c|ze) = 0$ for all $c \in \mathcal{C}^n$ for any $ze \in \mathcal{Z}^n \times \mathcal{E}^n$ with $\sigma(ze) = 0$. We begin with the inequality $d_{TV}(\sigma, \mu) \leq \epsilon_s$ that any $\sigma \in \mathcal{B}^{\epsilon_s}(\mu)$ must satisfy and proceed as follows:

$$\begin{aligned} 2\epsilon_s &\geq \|\mu - \sigma\|_1 = \sum_{cze \in \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E}^n} |\mu(cze) - \sigma(cze)| \\ &\geq \sum_{ze: \mu(ze) > 0} \sum_{c \in \mathcal{A}_{ze}} |\mu(cze) - \sigma(cze)| \end{aligned} \quad (22)$$

$$\begin{aligned} &\geq \left| \sum_{ze: \mu(ze) > 0} \sum_{c \in \mathcal{A}_{ze}} (\mu(cze) - \sigma(cze)) \right| \\ &= \left| \mathbb{E}_{\mu(ZE)}[\mathbb{P}_{\mu(C|ZE)}(\mu(C|ZE) \geq \gamma)] - \sum_{ze: \mu(ze) > 0} \sum_{c \in \mathcal{A}_{ze}} \sigma(cze) \right| \\ &\geq \mathbb{E}_{\mu(ZE)}[\mathbb{P}_{\mu(C|ZE)}(\mu(C|ZE) \geq \gamma)] - \sum_{ze: \mu(ze) > 0} \sum_{c \in \mathcal{A}_{ze}} \sigma(cze). \end{aligned} \quad (23)$$

The inequality in (22) follows as a result of the sum containing fewer terms; the inequality in (23) follows from the triangle inequality. Now from the AEP condition mentioned above we have the following:

$$\sum_{ze: \mu(ze) > 0} \sum_{c \in \mathcal{A}_{ze}} \sigma(cze) \geq \mathbb{E}_{\mu(ZE)}[\mathbb{P}_{\mu(C|ZE)}(\mu(C|ZE) \geq \gamma)] - 2\epsilon_s \geq 1 - \epsilon_a - 2\epsilon_s. \quad (24)$$

For any $\sigma \in \mathcal{B}^{\epsilon_s}(\mu)$, we define $M_{\mathbf{ze}}^\sigma$ for any $\mathbf{ze} \in \mathcal{Z}^n \times \mathcal{E}^n$ as $M_{\mathbf{ze}}^\sigma := \max_{\mathbf{c} \in \mathcal{C}^n} \sigma(\mathbf{c}|\mathbf{ze})$. The average conditional maximum probability is then expressed as $\bar{M}_\sigma := \sum_{\mathbf{ze}} M_{\mathbf{ze}}^\sigma \sigma(\mathbf{ze})$. Because $1 \leq \sum_{\mathbf{c} \in \mathcal{A}_{\mathbf{ze}}} \mu(\mathbf{c}|\mathbf{ze}) \leq \gamma |\mathcal{A}_{\mathbf{ze}}|$, we have $|\mathcal{A}_{\mathbf{ze}}| \leq 1/\gamma$ for each $\mathbf{ze}$ and we can write:

$$\begin{aligned} \sum_{\mathbf{ze}: \mu(\mathbf{ze}) > 0} \sum_{\mathbf{c} \in \mathcal{A}_{\mathbf{ze}}} \sigma(\mathbf{c}|\mathbf{ze}) &= \sum_{\mathbf{ze}: \mu(\mathbf{ze}) > 0} \sum_{\mathbf{c} \in \mathcal{A}_{\mathbf{ze}}} \sigma(\mathbf{c}|\mathbf{ze}) \sigma(\mathbf{ze}) \leq \sum_{\mathbf{ze}: \mu(\mathbf{ze}) > 0} \sum_{\mathbf{c} \in \mathcal{A}_{\mathbf{ze}}} M_{\mathbf{ze}}^\sigma \sigma(\mathbf{ze}) \\ &= \sum_{\mathbf{ze}: \mu(\mathbf{ze}) > 0} |\mathcal{A}_{\mathbf{ze}}| M_{\mathbf{ze}}^\sigma \sigma(\mathbf{ze}) \leq \frac{1}{\gamma} \sum_{\mathbf{ze}} M_{\mathbf{ze}}^\sigma \sigma(\mathbf{ze}) = \frac{\bar{M}_\sigma}{\gamma}. \end{aligned} \quad (25)$$

Using (24) and (25) we obtain $\bar{M}_\sigma \geq \gamma(1 - \epsilon_a - 2\epsilon_s)$ from which (21) follows using the definition of smooth average conditional min-entropy. $\square$

Having shown that the per-trial min-entropy generated by an IID attack is asymptotically bounded by the conditional Shannon entropy, we give the following definition of an *optimal* attack.

Definition 5 (Optimal IID Attack). *The distribution $\mu(\mathbf{CZE})$ of the sequence of random variables $\mathbf{C}, \mathbf{Z}, \mathbf{E}$ is an optimal IID attack if μ is obtained through an IID attack based on a single-trial distribution ω whose conditional Shannon entropy achieves the infimum defined below:*

$$h_{\min}(\rho) := \inf_{\omega(\mathbf{CZE}) \in \Sigma_{\mathbf{E}}^\rho} \mathbb{H}_\omega(\mathbf{C}|\mathbf{ZE}) \quad (26)$$

Additional motivation for naming the attack of Definition 5 *optimal* is provided by later results in this section, which show that the adversary must generate *at least* $h_{\min}(\rho)$ of per-trial conditional min-entropy asymptotically with any attack that replicates the observed distribution ρ .

In the theorem that follows, we formalise the claim that the infimum in (26) is achieved. This theorem corresponds to Theorem 43 in [15]; in comparison, the comprehensive proof provided here explicitly works out more of the steps. Crucially, this explicit approach also allowed us to provide an improvement upon the result of Theorem 43 in [15], decreasing the required value of $|\mathcal{E}|$ by one, thereby better characterising the adversary's optimal attack. Results in Section 4.2 will illustrate that no further improvement, i.e., a decrease in $|\mathcal{E}|$, is possible.

Theorem 4. *Suppose Π is closed and equal to the convex hull of its extreme points. Then, there is a distribution $\mu(\mathbf{CZE}) \in \Sigma_{\mathbf{E}}^\rho$ with $|\mathcal{E}| = 1 + \dim \Pi$ such that $\mathbb{H}_\mu(\mathbf{C}|\mathbf{ZE}) = h_{\min}(\rho)$.*

Proof. See Appendix B.1.1. $\square$

Theorem 4, in conjunction with the bound in Theorem 3, sets a benchmark for how well the adversary can perform with an IID attack that replicates the observed distribution $\rho(\mathbf{CZ})$. Specifically, the adversary's goal is to minimise the amount of per trial conditional min-entropy and this shows there exists a strategy to replicate the observed statistics while conceding no more min-entropy per trial than $h_{\min}(\rho)$, asymptotically.

3.2. Optimal PEFs

We now show that PEFs can asymptotically certify a min-entropy of $h_{\min}(\rho)$ per trial from an observed distribution ρ . This is notable since it shows that an IID attack can be asymptotically optimal: since the PEF method certifies the presence of $h_{\min}(\rho)$ min-entropy per trial against *any* attack, this means no attack can generate observed statistics consistent with ρ while conceding a smaller amount of randomness. This furthermore demonstrates that there is nothing to be gained (asymptotically) by the adversary employing a more sophisticated memory-based attack, since the PEF method allows for the possibility of memory attacks. Conversely, the below results show that the PEF method is asymptotically

optimal: no (correct) method can certify more min-entropy per trial from ρ than the amount that is present in an explicit attack.

To formalise and prove these claims, we use the following technical tool, called an “entropy estimator” as in [15].

Definition 6 (Entropy Estimator). *An entropy estimator of the model Π is a function $K(CZ)$ of the random variables C, Z such that $\mathbb{E}_\sigma[K(CZ)] \leq \mathbb{E}_\sigma[-\log_2(\sigma(C|Z))]$ holds for all $\sigma(CZ) \in \Pi$.*

Given an entropy estimator $K(CZ)$, we say that its *entropy estimate* at a distribution $\sigma(CZ)$ is $\mathbb{E}_\sigma[K(CZ)]$. We will see below that an entropy estimator can be used to construct PEFs certifying per-trial min-entropy arbitrarily close to its entropy estimate, underlying the significance of the following result:

Theorem 5. *Suppose Π satisfies the conditions of Theorem 4 and ρ is in the interior of Π . Then, there exists an entropy estimator whose entropy estimate at ρ is equal to $h_{\min}(\rho)$.*

Proof. See Appendix B.1.2. $\square$

The assumption that ρ is in the interior of Π will generally hold if ρ is estimated from real data, as the boundary of Π is a measure zero set. If the assumption is removed, a weaker version of the theorem can still be obtained, which is discussed in the proof in Appendix B.1.

The entropy estimator $K(CZ)$ whose existence is guaranteed by the above theorem can be used to show the existence of a family of PEFs that can become arbitrarily close to certifying $h_{\min}(\rho)$ amount of per-trial min-entropy. However, for a precise formulation of this claim we need a way to measure the asymptotic rate of min-entropy using PEFs. Recall from (8) that we can lower-bound the per-trial min-entropy certified by a PEF as:

$$\frac{1}{n} \mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(\mathbf{C}|\mathbf{Z}; \mathbf{S}) \geq \frac{1}{n} \left(1 + \frac{1}{\beta}\right) \log_2(\kappa) - \frac{1}{n} \log_2(p). \quad (27)$$

As in [15], we ignore the $\log_2(\kappa)$ term in the asymptotic regime, as the completeness parameter κ can be thought of as a “reasonable” lower bound on the probability that the protocol does not abort, a type of error parameter that one might try to decrease somewhat for longer experiments but not at the exponential decay rate required to make this term asymptotically significant. Focusing then on the $-(1/n) \log_2(p)$ term, recall that success of the protocol is determined by the occurrence of the event $S := \left\{ \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \leq p \right\}$, the inequality in which can be expressed equivalently as:

$$\frac{1}{n\beta} \sum_{i=1}^n \log_2(F_i) + \frac{1}{n\beta} \log_2(\epsilon) \geq -\frac{1}{n} \log_2(p).$$

The expression on the left hand side of the above inequality is the negative base-2 logarithm of the upper bound on $\mu_{\mathbf{e}}(\mathbf{C}|\mathbf{Z})$ for each $\mathbf{e} \in \mathcal{E}^n$ (refer to (2) and the comments following Corollary 1) and so is a rough measure of the amount of randomness, up to an error probability of ϵ , present in the outcome data. More concretely, since p will be chosen to make $-(1/n) \log(p)$ as large as reasonably possible to optimise min-entropy certified by (27), the anticipated value of the left hand side quantity can be used as a measure of certifiable randomness. For a stable experiment (i.e., one with each trial having the same distribution σ belonging to the same model Π), the quantity $(1/n) \sum_{i=1}^n \log_2(F_i) / \beta$ approaches $\mathbb{E}_\sigma[\log_2(F(CZ))] / \beta$ in the limit $n \rightarrow \infty$, while the term $(1/n\beta) \log_2(\epsilon)$ goes to zero for any fixed value of β and ϵ . Hence, we introduce the following quantity as a measure of per-trial min-entropy certified by a PEF.

Definition 7 (Log-Prob Rate). The log-prob rate of a PEF $F(CZ)$ with power β at a distribution $\rho(CZ)$ is defined as $\mathcal{O}_\rho(F; \beta) = \mathbb{E}_\rho[\log_2(F(CZ))]/\beta$.

We say that a PEF certifies randomness at a distribution ρ if the quantity $\mathcal{O}_\rho(F; \beta)$ is positive. We note that this definition is consistent with our expectation that only non-local distributions allow the certification of randomness, as the log-prob rate for a local distribution is a non-positive number, i.e., $\mathcal{O}_{\sigma_L}(F; \beta) \leq 0$: a local behaviour is a convex mixture of (finitely many) local deterministic behaviours $\sigma_{LD}(C|Z)$. Hence, with a fixed settings distribution $\pi(z) > 0$, the defining condition $\mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\beta] \leq 1$ of a PEF for a distribution defined as $\sigma(cz) = \sigma_{LD}(c|z)\pi(z)$, for all c, z , is equivalently expressed as $\mathbb{E}_\sigma[F(CZ)] \leq 1$, since $\sigma_{LD}(c|z)$ is either 0 or 1 for all c, z . Due to the concavity of log function, we then have $\mathbb{E}_\sigma[\log_2(F(CZ))] \leq \log_2(\mathbb{E}_\sigma[F(CZ)]) \leq 0$ using Jensen's inequality. Hence, no device-independent randomness can be certified at a local-realistic distribution.

Theorem 6. Given an entropy estimator $K(CZ)$ and an observed distribution $\rho(CZ)$, for any $\epsilon \in (0, 1/2)$ there is a PEF whose log-prob rate at ρ is greater than $\mathbb{E}_\rho[K(CZ)] - \epsilon$.

Our proof follows the general approach of Theorem 41 in [15], though we are able to shorten the argument.

Proof. Given an entropy estimator $K(CZ)$ and $\epsilon \in (0, 1/2)$ from the statement of the theorem, for any $\gamma > 0$ we can define a function

$$F(CZ) = 2^{(K(CZ) - \epsilon)\gamma} \quad (28)$$

We will show that there exists a (small) positive value of γ for which $F(CZ)$ is a PEF with power $\beta = \gamma$; the asymptotic log-prob rate of this PEF at ρ will then be $\mathbb{E}_\rho[\log_2(F(CZ))]/\beta = \mathbb{E}_\rho[K(CZ)] - \epsilon$ as desired. So, our task is to find a value of γ such that the following inequality holds for all $\sigma \in \Pi$:

$$\mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\gamma] \leq 1$$

We study the left side of the above expression as a function of γ ; specifically, define a function

$$f_\sigma(\gamma) = \mathbb{E}_\sigma[F(CZ)\sigma(C|Z)^\gamma] = \sum_{c,z:\sigma(cz)>0} \left[2^{K(cz)-\epsilon} \sigma(c|z) \right]^\gamma \sigma(cz)$$

which is, for any fixed choice of σ and $K(CZ)$, a convex combination of positive constants raised to the power of γ and so is infinitely differentiable at all $\gamma \in \mathbb{R}$. (Note that we never encounter the problematic form 0^0 because the argument of $[\cdot]^\gamma$ will always be strictly positive, as the sum defining f_σ extends only over values of c, z for which $\sigma(cz)$ is positive, and hence $\sigma(c|z) > 0$.) We can thus Taylor-expand f_σ about $\gamma = 0$, obtaining via the Lagrange remainder theorem that, for any positive γ , there exists a $k \in (0, \gamma)$ making the following equality hold:

$$f_\sigma(\gamma) = f_\sigma(0) + f'_\sigma(0)\gamma + \frac{f''_\sigma(k)}{2}\gamma^2 \quad (29)$$

The first term in the expansion satisfies $f_\sigma(0) = \sum_{cz} 1 \cdot \sigma(cz) = 1$. The coefficient of γ in (29) satisfies:

$$\begin{aligned} f'_\sigma(0) &= \sum_{c,z:\sigma(cz)>0} \left[2^{K(cz)-\epsilon} \sigma(c|z) \right]^0 \sigma(cz) \ln \left[2^{K(cz)-\epsilon} \sigma(c|z) \right] \\ &= \sum_{c,z:\sigma(cz)>0} \sigma(cz) [K(cz) - \epsilon + \log_2(\sigma(c|z))] \ln(2) \\ &= \ln(2) (\mathbb{E}_\sigma[K(CZ)] - \mathbb{E}[-\log_2(\sigma(c|z))] - \epsilon) \leq -\epsilon \ln(2) \end{aligned}$$

where the inequality follows from the condition $\mathbb{E}_\sigma[K(CZ)] \leq \mathbb{E}_\sigma[-\log_2(\sigma(C|Z))]$ in the definition of an entropy estimator. Hence, (29) yields

$$f_\sigma(\gamma) \leq 1 - \epsilon\gamma \ln(2) + \frac{f''_\sigma(k)}{2}\gamma^2 \quad (30)$$

for some $k \in (0, \gamma)$. Now, it is given that a fixed γ, k may be different in (30) for different choices of σ ; however, it must always lie in the interval $(0, \gamma)$, so if we can show that there is a choice of γ such that for *any* σ the following inequality holds for all $k \in (0, \gamma)$

$$\frac{f''_\sigma(k)}{2}\gamma^2 \leq \epsilon\gamma \ln(2) \quad (31)$$

then, for that value of γ , we will know that $F(CZ)$ as defined in (28) is a valid PEF satisfying the conditions of the theorem. To find the needed value of γ making (31) hold and complete the proof, we calculate

$$\begin{aligned} f''_\sigma(k) &= \ln^2(2) \sum_{c,z:\sigma(cz)>0} \left[2^{K(cz)-\epsilon}\sigma(c|z) \right]^k \left[\log_2 \left(2^{K(cz)-\epsilon}\sigma(c|z) \right) \right]^2 \sigma(cz) \\ &\leq \ln^2(2) M^k \sum_{cz:\sigma(cz)>0} \sigma(c|z)^{k+1} [K(cz) - \epsilon + \log_2(\sigma(c|z))]^2 \sigma(z) \end{aligned}$$

where $M = \max_{cz} 2^{K(cz)}$. We now assert that each quantity $\sigma(c|z)^{k+1} [K(cz) - \epsilon + \log_2(\sigma(c|z))]^2$ is bounded above by a constant N_{cz} for all $k > 0$ and N_{cz} is independent of σ . This follows because, for any fixed choice of c and z , this quantity is strictly smaller than the expression $g_{cz}(x) = x[K(cz) - \epsilon + \log_2(x)]^2$ for the choice of $x = \sigma(c|z) \in (0, 1]$ (note that since $\sigma(c|z) \in (0, 1]$, $\sigma(c|z)^{k+1} \leq \sigma(c|z)$ holds for any $k > 0$). Then, two applications of l'Hôpital's rule demonstrate that $\lim_{x \rightarrow 0} g_{cz}(x)$ exists and so g_{cz} can be extended to a continuous function on $[0, 1]$ where it has a maximum by the extreme value theorem. Invocation of the extreme value theorem, rather than computing an explicit bound, is what primarily allows us to shorten the proof compared to the argument proving Theorem 41 in [15]. Referring to this maximum as N_{cz} and letting $N = \max_{cz} N_{cz}$, we obtain the desired bound as shown below.

$$f''_\sigma(k) \leq \ln^2(2) M^k \sum_{z:\sigma(z)>0} \sigma(z) \sum_{c:\sigma(c,z)>0} N \leq \ln(2) M^k \sum_{z:\sigma(z)>0} \sigma(z) |\mathcal{C}| N = \ln(2) M^k |\mathcal{C}| N. \quad (32)$$

This shows that, if $M^k \gamma \leq 2\epsilon/|\mathcal{C}|N$ holds, then (31) holds, from which it follows that a sufficiently small choice of $\gamma > 0$ makes (31) hold for all $k \in (0, \gamma)$. $\square$

The combination of Theorem 5, which shows the existence of an entropy estimator with entropy estimate $h_{\min}(\rho)$, and Theorem 6, which enables the construction of a family of PEFs with log-prob rate arbitrarily close to this entropy estimate, demonstrates the asymptotic optimality of the PEF method.

3.3. Robustness of PEFs

We want to consider a question not considered in the previous PEF papers: can a PEF optimised for $\rho(CZ)$ certify randomness for a distribution different from ρ , where the difference is measured in terms of the total variation distance between them; in other words, how *robust* is the PEF? We will see in the next section that, in the (2,2,2) Bell scenario, for any behaviour corresponding to ρ violating the CHSH–Bell inequality, PEFs can be (up to any desired ϵ -tolerance) asymptotically optimal in terms of log-prob rate at ρ while also generating randomness at a positive rate for any behaviour (corresponding to a distribution of outcomes and settings) that violates the CHSH–Bell inequality by a fixed positive amount, which can be chosen to be as small as desired.

The following theorem gives a useful sufficient condition for a distribution different from ρ to have a positive log-prob rate and demonstrates that any nontrivial (i.e., non-constant) PEF will have at least some degree of robustness.

Theorem 7. Let $F(CZ) = G(CZ)^\beta$ be a non-constant positive PEF with power $\beta > 0$ for Π . The log-prob rate $\mathcal{O}_\sigma(F; \beta)$ at a distribution $\sigma(CZ) \in \Pi$ is related to the log-prob rate $\mathcal{O}_\rho(F; \beta)$ at $\rho(CZ) \in \Pi$ and the total variation distance between ρ and σ as

$$|\mathcal{O}_\rho(F; \beta) - \mathcal{O}_\sigma(F; \beta)| \leq (L - l)d_{TV}(\rho, \sigma), \quad (33)$$

where $L = \max_{cz} \log_2(G(cz))$ and $l = \min_{cz} \log_2(G(cz))$. Consequently, assuming that $\mathcal{O}_\rho(F; \beta)$ is positive, the following upper bound on the total variation distance between $\rho(CZ)$ and $\sigma(CZ)$ is a sufficient condition for F to have a positive log-prob rate at $\sigma(CZ)$

$$d_{TV}(\rho, \sigma) < \mathbb{E}_\rho[\log_2(G)] / (L - l). \quad (34)$$

Proof. Using the definition of log-prob rate at a given distribution we have

$$\begin{aligned} |\mathcal{O}_\rho(F; \beta) - \mathcal{O}_\sigma(F; \beta)| &= \left| \sum_{cz} \frac{1}{\beta} \left[\log_2(G(cz)^\beta) (\rho(cz) - \sigma(cz)) \right] \right| \\ &= \left| \sum_{cz} \left(\log_2(G(cz)) + \frac{L+l}{2} - \frac{L+l}{2} \right) (\rho(cz) - \sigma(cz)) \right| \\ &= \left| \sum_{cz} \left(\log_2(G(cz)) - \frac{L+l}{2} \right) (\rho(cz) - \sigma(cz)) + \frac{L+l}{2} \sum_{cz} (\rho(cz) - \sigma(cz)) \right| \\ &\leq \sum_{cz} \left| \log_2(G(cz)) - \frac{L+l}{2} \right| |\rho(cz) - \sigma(cz)| \\ &\leq (L - l) \frac{1}{2} \sum_{cz} |\rho(cz) - \sigma(cz)| = (L - l)d_{TV}(\rho, \sigma) \end{aligned}$$

Hence, we have

$$\mathcal{O}_\rho(F; \beta) - (L - l)d_{TV}(\rho, \sigma) \leq \mathcal{O}_\sigma(F; \beta) \leq \mathcal{O}_\rho(F; \beta) + (L - l)d_{TV}(\rho, \sigma).$$

Assuming that $\mathcal{O}_\rho(F; \beta)$ is positive, a sufficient condition for $\mathcal{O}_\sigma(F; \beta)$ to be positive is $\mathcal{O}_\rho(F; \beta) > |L - l|d_{TV}(\rho, \sigma)$ or, equivalently, the following bound on $d_{TV}(\rho, \sigma)$:

$$d_{TV}(\rho, \sigma) < \mathcal{O}_\rho(F; \beta) / (L - l) = \mathbb{E}_\rho[\log_2(G)] / (L - l).$$

□

We will see in Section 4.2 that the bound (33) can be saturated and so is tight.

4. Application to the (2,2,2) Bell Scenario

Here, we explore the application of the results of the previous section to the (2,2,2) Bell scenario (that of two parties, two measurement settings and two outcomes). First, working within the trial model of no-signalling distributions Π_{NS} , we show that PEFs can be simultaneously asymptotically optimal and robust by means of an explicit construction of a sequence of PEFs that approaches the optimal log-prob rate for the target distribution while simultaneously generating randomness at a positive rate for any other distribution violating the CHSH inequality.

In the course of this exercise, we will observe that the optimal adversarial attack—one generating the observed statistics (consistent with an expected trial distribution ρ) while asymptotically yielding $h_{\min}(\rho)$ amount of per-trial randomness—is always achieved through a single-trial distribution that marginalises to ρ through a convex combination of a

single extremal no-signalling nonlocal distribution and a local realistic distribution (which itself consists of a convex mixture of up to eight extremal local deterministic distributions). This is a notable feature, revealing that the adversary never needs to prepare more than one nonlocal distribution to simulate the observed distribution with as little min-entropy as possible. Later in this section, we explore the potential for generalisation of this feature to the (2,2,2) scenario restricted to quantum distributions (Π_Q); if true, this would be an important finding, outlining the optimal approach of a (more realistic) quantum-limited adversary attacking the PEF protocol. The general observation that preparing a single nonlocal state is preferable to preparing multiple ones underlies the significance of the answer to this question. We find some evidence that the feature—only requiring one extremal nonlocal distribution in the convex combination attack—may hold for Π_Q in the (2,2,2) Bell scenario but this may be a difficult question to resolve due to the complicated geometry of the quantum set. We also explore possible generalisations of this feature to no-signalling trial models for (n, m, k) Bell scenarios where n, m or k exceed 2, and find that it *does not* hold in any of these cases—so the question of whether this holds in a given Bell scenario and trial model is non-trivial in general.

We begin with a brief review of the (2,2,2) Bell scenario and some features of the set Π_{NS} of no-signalling distributions in this scenario.

4.1. A Brief Review of the (2,2,2) Bell Scenario

The (2,2,2) Bell scenario is the minimal Bell scenario, comprising two spatially separated parties Alice and Bob, each having two measurement settings and two possible outcomes corresponding to each setting. The measurement settings for Alice and Bob are represented by the RVs X, Y realising values $x, y \in \{0, 1\}$ and the measurement outcomes are represented by the RVs A, B realising values $a, b \in \{0, 1\}$. With $\sigma_s(XY)$ representing a fixed settings distribution, we refer to the sets Π_{NS}, Π_Q and Π_L as no-signalling, quantum and local models, respectively, when they comprise of distributions $\mu(ABXY) := \mu(AB|XY)\sigma_s(XY)$, where the conditional probabilities $\mu(AB|XY)$, referred to as behaviours, are constrained by the no-signalling, quantum and local realism principle, respectively. Henceforth, all distributions $\mu(ABXY)$ belonging to a model are defined as $\mu(ABXY) := \mu(AB|XY)\sigma_s(XY)$, and we associate a model with its constituent behaviour $\mu(AB|XY)$ or distribution $\mu(ABXY)$, indistinctively, since the settings distribution is fixed. Recall that the model Π_{NS} is a polytope, the extremal points of which consist of the behaviours $\mu_{\text{extr}}(AB|XY) \equiv \{\mu_{\text{extr}}(ab|xy) : a, b, x, y \in \{0, 1\}\}$ defined below.

$$\mu_{\text{PR}}^{\alpha\beta\gamma}(ab|xy) := \begin{cases} \frac{1}{2} & : a \oplus b = xy \oplus \alpha x \oplus \beta y \oplus \gamma \\ 0 & : \text{otherwise} \end{cases} \quad (35)$$

$$\mu_{\text{LD}}^{\alpha\beta\gamma\delta}(ab|xy) := \begin{cases} 1 & : a = \alpha x \oplus \beta, b = \gamma y \oplus \delta \\ 0 & : \text{otherwise} \end{cases} \quad (36)$$

where $\alpha, \beta, \gamma, \delta \in \{0, 1\}$ and $\oplus$ denotes addition modulo 2; (35) and (36) are known as the Popescu–Rohrlich (PR) behaviours [22] and the local deterministic (LD) behaviours, respectively. The CHSH–Bell inequalities shown below are known to be the only non-trivial facet inequalities delimiting the local polytope which is the convex hull of the LD behaviours [23]. Corresponding to each choice of $\alpha, \beta, \gamma \in \{0, 1\}$, the inequalities represent a version of the canonical CHSH–Bell inequality.

$$B^{\alpha\beta\gamma} := (-1)^\gamma E_{00} + (-1)^{\beta+\gamma} E_{01} + (-1)^{\alpha+\gamma} E_{10} + (-1)^{\alpha+\beta+\gamma+1} E_{11} \leq 2, \quad (37)$$

where $E_{xy} := \sum_{a,b=0}^1 (-1)^{a+b} \mu(ab|xy)$ for $x, y \in \{0, 1\}$. The nonlocal algebraic maximum for the expression $B^{\alpha\beta\gamma}$ is 4. The local maximum is obtained by eight $\mu_{\text{LD}}^{\alpha\beta\gamma\delta}(AB|XY)$ behaviours for each $B^{\alpha\beta\gamma}$. The sets LD_i , $i \in \{1, 2, \dots, 8\}$, each comprise of eight LD behaviours that saturate—i.e., achieve a value of 2—exactly one $B^{\alpha\beta\gamma}$. A result proven in [24] (see Theorems 2.1 and 2.2 therein) states that any behaviour violating (37) can be represented as

a convex combination of one PR box achieving the nonlocal maximum for $B^{\alpha\beta\gamma}$ and (up to) eight LD behaviours of the corresponding LD_i set saturating it. In fact, the geometry of the no-signalling polytope in this Bell scenario is such that there is a one-to-one correspondence between the nonlocal no-signalling extremal points, the PR boxes, in (35) and the non-trivial facets of the local polytope described by (37), with exactly one extremal point violating it up to the algebraic maximum of four for each choice of $(\alpha, \beta, \gamma) \in \{0, 1\}^3$. Hence, any nonlocal behaviour—that violates a given version of the CHSH–Bell inequality—is contained in a *nonlocal 8-simplex* whose vertices are the one PR box that maximally violates that particular version and the eight LD behaviours that saturate it. Recall that a p -simplex is a p -dimensional polytope which is the convex hull of its $p + 1$ vertices. More formally, if the set $C := \{\vec{a}_0, \vec{a}_1, \dots, \vec{a}_p\} \subset \mathbb{R}^n$ of $p + 1$ points are affinely independent, then the p -simplex determined by them is the following set of points:

$$\Delta^p := \left\{ \sum_{k=0}^p \theta_k \vec{a}_k \mid \sum_{k=0}^p \theta_k = 1, \theta_k \geq 0 \text{ for } k = 0, 1, \dots, p \right\}.$$

The affine independence condition means that the only admissible choice of $\theta_k \in \mathbb{R}$ such that $\sum_{k=0}^p \theta_k \vec{a}_k = \vec{0}$ and $\sum_{k=0}^p \theta_k = 0$ are satisfied is $\theta_k = 0$ for all k ; this holds if and only if the vectors $\vec{a}_k - \vec{a}_0$ are linearly independent for $k = 1, 2, \dots, p$.

One can check that the PR box that achieves the nonlocal maximum for a given version of the CHSH–Bell expression $B^{\alpha\beta\gamma}$ and the eight LD behaviours that achieve the local maximum for it are affinely independent. Since $a, b, x, y \in \{0, 1\}$ and $|\{0, 1\}^4| = 16$, we can represent the behaviours $\mu(ab|xy)$ in this Bell scenario as vectors $\vec{\mu} \in \mathbb{R}^{16}$ as shown in Table 1. Then, the affine independence is apparent: letting the PR box behaviour be $\vec{a}_0$ and the LD behaviours be the other $\vec{a}_k$, each $\vec{a}_k - \vec{a}_0$ term has a unique column where it contains a “1” while all of the other terms contain “0”, ensuring linear independence.

Table 1. These probability vectors in $\mathbb{R}^{16}$ are the PR box $\vec{\mu}_{PR,1} \equiv \mu_{PR}^{000}$ that achieves the nonlocal maximum of 4 and the eight LD behaviours $\vec{\mu}_{LD,1}, \dots, \vec{\mu}_{LD,8}$ that achieve the local maximum of 2 for the standard CHSH–Bell expression B^{000} , with the LD behaviours corresponding to the eight probability tables numbered 1, 4, 5, 8, 9, 12, 14 and 15 in Table A2 of [24], and also given in the first row of Table 2. One can verify the affine independence of the nine vectors above by verifying that the eight vectors obtained by subtracting the first vector from the remaining eight are linearly independent.

	00				01				xy				10				11			
	ab				ab				ab				ab				ab			
	00	01	10	11	00	01	10	11	00	01	10	11	00	01	10	11	00	01	10	11
$\vec{\mu}_{PR,1}$	1/2	0	0	1/2	1/2	0	0	1/2	1/2	0	0	1/2	0	1/2	1/2	0	0	0	0	0
$\vec{\mu}_{LD,1}$	1	0	0	0	1	0	0	0	1	0	0	0	1	0	0	0	0	0	0	0
$\vec{\mu}_{LD,2}$	0	0	0	1	0	0	0	1	0	0	0	1	0	0	0	0	0	0	1	0
$\vec{\mu}_{LD,3}$	1	0	0	0	0	1	0	0	1	0	0	0	0	1	0	0	1	0	0	0
$\vec{\mu}_{LD,4}$	0	0	0	1	0	0	1	0	0	0	1	0	0	0	1	0	0	1	0	0
$\vec{\mu}_{LD,5}$	1	0	0	0	0	1	0	0	0	0	1	0	0	0	0	1	0	1	0	0
$\vec{\mu}_{LD,6}$	0	0	0	1	0	0	0	1	0	1	0	0	0	1	0	0	0	1	0	0
$\vec{\mu}_{LD,7}$	0	1	0	0	1	0	0	0	0	0	0	1	0	0	0	1	0	0	1	0
$\vec{\mu}_{LD,8}$	0	0	1	0	0	0	0	1	1	0	0	0	0	0	1	0	0	0	0	0

It is known that a behaviour belonging to $\Pi_{NS} \setminus \Pi_L$ violates exactly one of the eight CHSH–Bell inequalities. The impossibility of simultaneously violating a specific pair of CHSH–Bell inequalities can be seen as presented in [25]: suppose a behaviour in $\Pi_{NS} \setminus \Pi_L$ violates both inequalities corresponding to $(\alpha, \beta, \gamma) = (0, 0, 0)$ and $(\alpha, \beta, \gamma) = (1, 0, 0)$, then $E_{00} + E_{01} + E_{10} - E_{11} > 2$ and $E_{00} + E_{01} - E_{10} + E_{11} > 2$ holds for the same behaviour. Adding these two inequalities we have $2(E_{00} + E_{01}) > 4$, i.e., $E_{00} + E_{01} > 2$, which is not possible to satisfy since the correlations E_{xy} satisfy $|E_{xy}| \leq 1$.

Table 2 lists the eight versions of the Bell expression $B^{\alpha\beta\gamma}$ and the eight nonlocal 8-simplices $\Delta_{\text{PR},i}^8$ containing points that violate the corresponding CHSH–Bell inequality. Any nonlocal no-signalling behaviour ultimately belongs to exactly one such simplex.

4.2. Robust PEFs and Optimal Adversarial Attacks in the (2,2,2) Bell Scenario

We now examine the *robustness* of PEFs that are optimal for an anticipated distribution ρ and a fixed number of planned trials n . We first review how we find optimal PEFs in this scenario. The constrained maximisation routine in (10) provides a method to find useful PEFs with respect to an anticipated trial distribution, with Lemma 1 showing that the feasibility constraints in (10) can be restricted to only the distributions corresponding to the eight PR and sixteen LD behaviours (with a fixed settings distribution $\sigma_s(XY) > 0$).

In practice, the number of trials n will affect the choice of β and the PEF that optimises the quantity $\mathbb{E}_\rho[(n \log_2(F(CZ)) + \log_2(\epsilon))/\beta]$, a quantity which (per the discussion surrounding (10)) can be thought of as the anticipated amount of raw randomness from running the experiment whose trial distribution is expected to be ρ . If we divide this quantity by n , we arrive at a measure of expected randomness per trial for the optimal PEF at a given value of β , called the *net log-prob rate*: the function $(\max_F \mathcal{O}_\rho(F; \beta)) + \log_2(\epsilon)/n\beta$. Figure 2 shows a plot of the net log-prob rates corresponding to two different values of n , as well as the supremum of the log-prob rate, for β varying from 0.001 to 0.1 and ϵ fixed at the value 10^{-4} . The value of β , and the corresponding PEF that maximises the curve, is then the best choice for the given planned number of trials n .

Table 2. The eight nonlocal 8-simplices containing behaviours that violate the corresponding version of the CHSH–Bell inequality. We identify each 8-simplex $\Delta_{\text{PR},i}^8$ with a PR box which solely contributes to the nonlocality of the behaviour violating the CHSH–Bell inequality.

$B^{\alpha\beta\gamma}$	$\Delta_{\text{PR},i}^8$
B^{000}	$\Delta_{\text{PR},1}^8 := \text{conv}\{\mu_{\text{PR}}^{000}, \mu_{\text{LD}}^{0000}, \mu_{\text{LD}}^{0101}, \mu_{\text{LD}}^{0010}, \mu_{\text{LD}}^{0111}, \mu_{\text{LD}}^{1000}, \mu_{\text{LD}}^{1101}, \mu_{\text{LD}}^{1011}, \mu_{\text{LD}}^{1110}\}$
B^{001}	$\Delta_{\text{PR},2}^8 := \text{conv}\{\mu_{\text{PR}}^{001}, \mu_{\text{LD}}^{0001}, \mu_{\text{LD}}^{0011}, \mu_{\text{LD}}^{0100}, \mu_{\text{LD}}^{0110}, \mu_{\text{LD}}^{1001}, \mu_{\text{LD}}^{1010}, \mu_{\text{LD}}^{1100}, \mu_{\text{LD}}^{1111}\}$
B^{010}	$\Delta_{\text{PR},3}^8 := \text{conv}\{\mu_{\text{PR}}^{010}, \mu_{\text{LD}}^{0000}, \mu_{\text{LD}}^{0010}, \mu_{\text{LD}}^{0101}, \mu_{\text{LD}}^{0111}, \mu_{\text{LD}}^{1001}, \mu_{\text{LD}}^{1010}, \mu_{\text{LD}}^{1100}, \mu_{\text{LD}}^{1111}\}$
B^{011}	$\Delta_{\text{PR},4}^8 := \text{conv}\{\mu_{\text{PR}}^{011}, \mu_{\text{LD}}^{0001}, \mu_{\text{LD}}^{0011}, \mu_{\text{LD}}^{0100}, \mu_{\text{LD}}^{0110}, \mu_{\text{LD}}^{1000}, \mu_{\text{LD}}^{1011}, \mu_{\text{LD}}^{1101}, \mu_{\text{LD}}^{1110}\}$
B^{100}	$\Delta_{\text{PR},5}^8 := \text{conv}\{\mu_{\text{PR}}^{100}, \mu_{\text{LD}}^{0000}, \mu_{\text{LD}}^{0011}, \mu_{\text{LD}}^{0101}, \mu_{\text{LD}}^{0110}, \mu_{\text{LD}}^{1000}, \mu_{\text{LD}}^{1010}, \mu_{\text{LD}}^{1101}, \mu_{\text{LD}}^{1111}\}$
B^{101}	$\Delta_{\text{PR},6}^8 := \text{conv}\{\mu_{\text{PR}}^{101}, \mu_{\text{LD}}^{0001}, \mu_{\text{LD}}^{0010}, \mu_{\text{LD}}^{0100}, \mu_{\text{LD}}^{0111}, \mu_{\text{LD}}^{1001}, \mu_{\text{LD}}^{1011}, \mu_{\text{LD}}^{1100}, \mu_{\text{LD}}^{1110}\}$
B^{110}	$\Delta_{\text{PR},7}^8 := \text{conv}\{\mu_{\text{PR}}^{110}, \mu_{\text{LD}}^{0001}, \mu_{\text{LD}}^{0010}, \mu_{\text{LD}}^{0100}, \mu_{\text{LD}}^{0111}, \mu_{\text{LD}}^{1000}, \mu_{\text{LD}}^{1010}, \mu_{\text{LD}}^{1101}, \mu_{\text{LD}}^{1111}\}$
B^{111}	$\Delta_{\text{PR},8}^8 := \text{conv}\{\mu_{\text{PR}}^{111}, \mu_{\text{LD}}^{0000}, \mu_{\text{LD}}^{0011}, \mu_{\text{LD}}^{0101}, \mu_{\text{LD}}^{0110}, \mu_{\text{LD}}^{1001}, \mu_{\text{LD}}^{1011}, \mu_{\text{LD}}^{1100}, \mu_{\text{LD}}^{1110}\}$

The plot illustrates some notable features of PEFs. First, it was proved in Appendix D of [14] that assuming a stable experiment (with each trial distribution ρ) the function $\sup_F \mathcal{O}_\rho(F; \beta)$ is monotonically non-increasing in $\beta > 0$ which implies that the global supremum of the log-prob rates $\sup_{\beta > 0} \sup_F \mathcal{O}_\rho(F; \beta)$, for all PEFs with positive powers, is achieved in the limit $\beta \rightarrow 0$. We observe this with the top curve. For a fixed ϵ , the net log-prob rate converges upwards to $\sup_F \mathcal{O}_\rho(F; \beta)$ for each β as $n \rightarrow 0$ but, for any fixed value of n , $\log_2(\epsilon)/n\beta$ diverges to $-\infty$ as $\beta \rightarrow 0$. Hence, in a finite trial regime the supremum of the log-prob rates (attainable by PEFs with positive powers) is not achieved—the maximum value of the *net log-prob rate* is achieved at a β away from 0. The general trend is that for a value of n the net log-prob rate achieves a higher value corresponding to a lower value of β ; the net log-prob rate is improved by a reduction in power and an increase in the number of trials. This is observed in Figure 2 for the two choices of $n = 1.5 \times 10^5$ and $n = 2.4 \times 10^5$. As a side note, the proof that $\beta' < \beta$ implies $\sup_F \mathcal{O}_\rho(F; \beta') \leq \sup_F \mathcal{O}_\rho(F; \beta)$ is straightforward: write $\beta' = \gamma\beta$ with $0 < \gamma < 1$; then, for any F in the scope of $\sup_F \mathcal{O}_\rho(F; \beta)$, it turns out F^γ is a PEF with power β' , for which the equality $\mathcal{O}_\rho(F^\gamma; \beta') = \mathcal{O}_\rho(F; \beta)$ follows immediately from the definition of log-prob rate—hence, the supremum of log-prob rates cannot be smaller at β' . F^γ is a PEF with power β' as $\mathbb{E}_\rho(F^\gamma \sigma(c|z)^{\beta\gamma}) \leq \mathbb{E}_\rho(F \sigma(c|z)^\beta)^\gamma \leq 1^\gamma = 1$,

with the first inequality holding by Jensen's inequality ($f(x) = x^\gamma$ is concave) and the second because F is a PEF with power β .

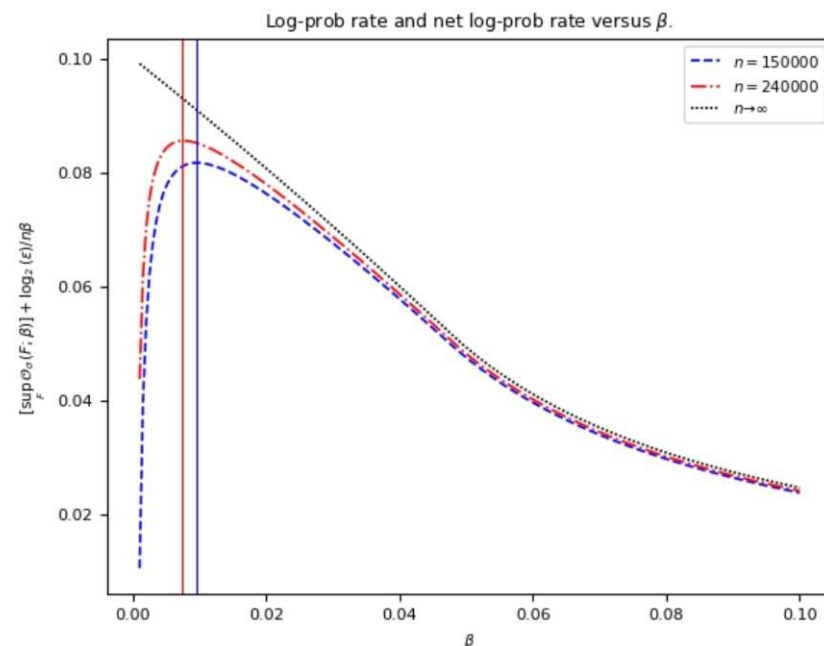


Figure 2. A plot showing the net log-prob rates for $n = 1.5 \times 10^5$ (the dashed curve) and $n = 2.4 \times 10^5$ (the dash-dotted curve) with $\epsilon = 10^{-4}$ and β varying in the interval $(0.001, 0.1)$. The dotted curve is the log-prob rate $\sup_F \mathcal{O}_\rho(F; \beta)$, an upper bound for the net log-prob rate in the limit as $n \rightarrow \infty$. We selected 200 equally spaced points in the interval $(0.001, 0.1)$ for β and performed the maximisation $\max_F \mathbb{E}_\rho[\log_2(F(ABXY))]$ constrained by: (1) the non-negativity of PEFs and (2) the defining condition $\mathbb{E}_\mu[F(ABXY)\mu(AB|XY)^\beta] \leq 1$ at all distributions μ corresponding to the eight PR and sixteen LD behaviours with a fixed uniform settings distribution $\mu(xy) = 1/4$ for all $x, y \in \{0, 1\}$. The anticipated distribution ρ used here was the one corresponding to the behaviour given in Table I in [15]. We observe that the maximum value for the net log-prob rate—indicated by the solid vertical lines—is achieved at a lower value of β for a higher value of n .

The arguments above illustrate how it is necessary to consider a range of β values to find the optimal choice. We remark there is an upper limit to the range of β values that must be considered: it was noted in [14] (see Appendix F therein) that there exists a certain threshold value $\beta_{\text{th}}^{\text{NS}}$ such that, for all $\beta \geq \beta_{\text{th}}^{\text{NS}}$, the optimisation problem in (10) will return the same PEF independent of the choice of β and [14] cites numerical evidence that this bound is $\beta_{\text{th}}^{\text{NS}} \simeq 0.4151$. The following result, whose proof we give in the appendix, derives this threshold analytically, finding it to have the exact value $\log_2(4/3)$.

Proposition 1. For the set of behaviours Π_{NS} , the PEF optimisation in (10) is independent of the power β for $\beta \geq \log_2(4/3)$.

Proof. See Appendix F. $\square$

We now ask how optimal PEFs for lower and lower values of β (and correspondingly higher values of n) compare on the question of *robustness*, in the following sense: can a PEF optimised with respect to a distribution ρ violating the standard CHSH–Bell inequality be used to certify randomness of distributions that are different from ρ , provided they violate the same CHSH–Bell inequality? This question is relevant because, in practice, the observed experimental distribution will never be exactly the same as the anticipated one and may be somewhat different depending on many potential factors. Figure 3 gives an illustration of the matter of robustness. Comparing the two plots of the log-prob rate for

quantum-realizable distributions on the two-dimensional slice (shown in Figure 4b) above the standard CHSH–Bell facet, we observe that the level set denoting a zero amount of certified randomness in the right hand plot (which corresponds to a lower value of β than that on the left) is pushed further down to (almost touching) the standard CHSH–Bell facet.

This suggests that the asymptotic optimality of a PEF need not entail a trade-off with its robustness; indeed, we observed that, in many cases, as $\beta > 0$ assumes smaller and smaller values, the PEF optimised for a fixed ρ violating the standard CHSH–Bell inequality becomes more and more robust in the sense that it certifies randomness at a positive rate (asymptotically) for increasingly statistically different σ .

We show that this is a general feature. To this end, we define a sequence of PEFs that is both asymptotically optimal with respect to the log-prob rate and is asymptotically robust in the sense that, given any distribution violating the standard CHSH–Bell inequality, all the PEFs beyond a point in the sequence certify randomness at a positive rate. To construct this PEF sequence, we first define the function $K_*(ABXY)$ as shown below:

$$K_*(abxy) := 4[a \oplus b = xy] - 3, \quad (38)$$

where $a, b, x, y \in \{0, 1\}$ and the function $[\cdot]$ evaluates to 1 if the condition within holds, 0 otherwise. The function defined in (38) is an entropy estimator for the distributions in the no-signalling polytope when the settings are equiprobable, i.e., $\sigma_s(xy) = 1/4$ for all choices of x and y . To see this, recalling Definition 6 we can check—by direct evaluation—whether K_* satisfies the inequality $\mathbb{E}_\sigma[K(CZ)] \leq \mathbb{E}_\sigma[-\log_2(\sigma(C|Z))]$ when σ is each of the extremal points of the no-signalling polytope. It is sufficient to check this condition for the extremal points of the no-signalling set, i.e., the PR behaviours and the LD behaviours. This is because if σ is expressible as $\sigma = \lambda\sigma_1 + (1 - \lambda)\sigma_2$ then, for any function K satisfying $\mathbb{E}_{\sigma_i}[K(ABXY)] \leq \mathbb{H}_{\sigma_i}(AB|XY)$, we have $\mathbb{E}_\sigma[K] = \lambda\mathbb{E}_{\sigma_1}[K] + (1 - \lambda)\mathbb{E}_{\sigma_2}[K] \leq \lambda\mathbb{H}_{\sigma_1}(AB|XY) + (1 - \lambda)\mathbb{H}_{\sigma_2}(AB|XY) \leq \mathbb{H}_\sigma(AB|XY)$. Hence, if the condition holds for the extremal points, it will hold for all points in the set. To see that it does, we confirm by inspection that $\mathbb{E}_\sigma[K_*]$ attains the value 1 for the PR behaviour achieving the no-signalling maximum for the standard CHSH function, the value -3 for the PR behaviour achieving -4 and the value -1 for each of the PR behaviours that achieve the value 0, which are all less than or equal to the conditional Shannon entropy of the respective PR behaviours, which is 1. Likewise, we can check that K_* is a valid entropy estimator for all the LD behaviours; it takes the value zero for the eight local deterministic distributions appearing in Table 1 and -2 for the other eight, while $\mathbb{H}(AB|XY) = 0$ for these distributions. Hence, we have verified that K_* satisfies the entropy estimator condition for all the extremal behaviours and by extension all behaviours in the no-signalling polytope.

Having shown K_* is an entropy estimator, we next consider a sequence of functions $\{F_k\}_{k=1}^\infty$ where F_k is defined according to the construction in Theorem 6:

$$F_k(ABXY) = 2^{(K_*(ABXY) - e^{-k})\beta_k}, \quad (39)$$

where we choose a positive β_k making F_k a PEF for each k , whose existence is guaranteed by the theorem. By construction, for each k the function F_k is a valid PEF with power $\beta_k > 0$ for the set of no-signalling distributions. The log-prob rate of F_k at σ is:

$$\mathcal{O}_\sigma(F_k; \beta_k) = \frac{1}{\beta_k} \mathbb{E}_\sigma \left[\log_2 \left(2^{(K_* - e^{-k})\beta_k} \right) \right] = \mathbb{E}_\sigma[K_*] - e^{-k}. \quad (40)$$

We show robustness of the sequence in the following sense: for any $\sigma \in \Pi_{\text{NS}}$ violating the standard CHSH–Bell inequality, the log-prob rate of the sequence of PEFs $\{F_k\}_{k=1}^\infty$ is eventually positive. To see this, recall that, as discussed in our brief review of the (2,2,2) Bell scenario, behaviours violating the standard CHSH–Bell inequality are contained in the nonlocal 8-simplex $\Delta_{\text{PR},1}^8$ (see Table 2). Hence, σ is expressible as a convex combination of the vertices of $\Delta_{\text{PR},1}^8$:

$$\sigma(ab|xy)\sigma_s(xy) = \lambda_{\text{PR},1}\mu_{\text{PR},1}(ab|xy)\sigma_s(xy) + \sum_{i=1}^8 \alpha_i \mu_{\text{LD},i}(ab|xy)\sigma_s(xy), \quad (41)$$

where $\lambda_{\text{PR},1} + \sum_{i=1}^8 \alpha_i = 1$. This decomposition allows us to express the log-prob rate in terms of the standard CHSH–Bell function, which we define as

$$S(ABXY) := (-1)^{XY}(-1)^{A+B}/\sigma_s(XY),$$

where $\sigma_s(XY)$ is the fixed settings distribution. We see that $\lambda_{\text{PR},1} = (S_\sigma - 2)/2$ in (41), where S_σ is the expected standard CHSH–Bell value according to the distribution $\sigma(ABXY) = \sigma(ab|xy)\sigma_s(xy)$. This follows by computing the expectation of S according to the PR box distribution $\mu_{\text{PR},1} = \mu_{\text{PR},1}(abxy) = \mu_{\text{PR},1}(ab|xy)\sigma_s(xy)$, which is 4, and the expectation of S according to the local distribution $\mu_{\text{L},i} = \mu_{\text{L},i}(abxy) = \mu_{\text{LD},i}(ab|xy)\sigma_s(xy)$, which is 2. The log-prob rate $\mathcal{O}_\sigma(F_k; \beta_k)$ for F_k at σ is then expressed as:

$$\mathcal{O}_\sigma(F_k; \beta_k) = \left(\frac{S_\sigma - 2}{2} \right) \mathbb{E}_{\mu_{\text{PR},1}}[K_*] + \sum_{i=1}^8 \alpha_i \mathbb{E}_{\mu_{\text{LD},i}}[K_*] - e^{-k}. \quad (42)$$

Since $\mathbb{E}_{\mu_{\text{LD},i}}[K_*]$ evaluates to zero for each $\mu_{\text{LD},i}$ and $\mathbb{E}_{\mu_{\text{PR},1}}[K_*]$ evaluates to 1, the expression for $\mathcal{O}_\sigma(F_k; \beta_k)$ reduces to $\mathcal{O}_\sigma(F_k; \beta_k) = \frac{S_\sigma - 2}{2} - e^{-k}$. As $k \rightarrow \infty$, $\mathcal{O}_\sigma(F_k; \beta_k) = (S_\sigma - 2)/2$ and so the quantity is eventually strictly positive provided $S_\sigma > 2$, i.e., provided σ violates the standard CHSH–Bell inequality.

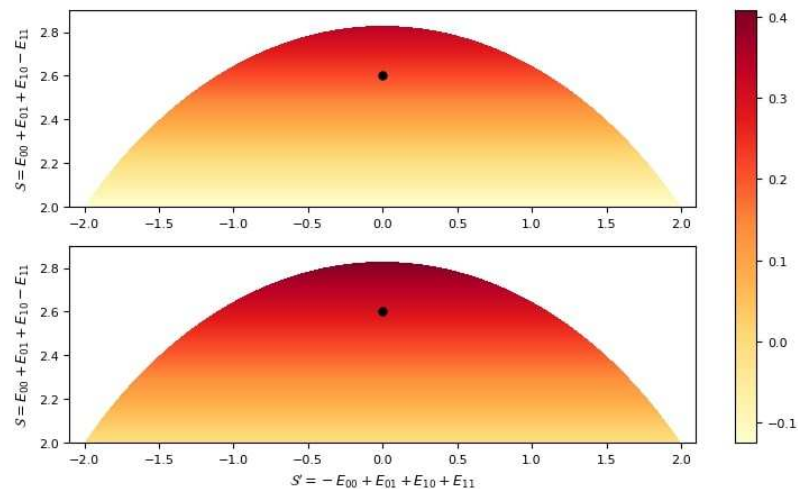


Figure 3. A heat map illustrating the robustness of PEF with log-prob rate as the figure of merit, evaluated for behaviours $\sigma(ab|xy)$ on the two-dimensional slice of the set of quantum behaviours (shown in Figure 4b) above the standard CHSH–Bell facet. The behaviours on the two-dimensional slice shown above are parameterised by S and S' as shown in (46) with the added restrictions $S^2 + (S')^2 \leq 8$ and $2 \leq S \leq 2\sqrt{2}$, $-2 \leq S' \leq 2$ (see also Table 3). Assuming a uniform distribution for the settings, $\sigma_s(xy) = 1/4$ for all x, y , we plot the log-prob rate $\sum_{abxy} [\log_2 F_*(abxy)\sigma(ab|xy)\sigma_s(xy)]/\beta$ for all distributions in the slice. The black dot corresponds to the behaviour (and hence the distribution) with respect to which we perform the PEF optimisation for a fixed n and ϵ to obtain F_* . The coordinates for the black dot are $(S', S) \equiv (0, 2.6)$. (a) *Top figure*: Heat map with F_* obtained from the PEF optimisation in (10) with respect to the fixed distribution (corresponding to the black dot in the figures), fixed n , ϵ and $\beta = 0.1$. Below $S \simeq 2.22145$ no device-independent randomness can be certified. (b) *Bottom figure*: Heat map with F_* obtained from the PEF optimisation in (10) with respect to the fixed distribution (corresponding to the black dot in the figures), fixed n , ϵ and $\beta = 0.01$. Below $S \simeq 2.02072$ no device-independent randomness can be certified.

Continuing our discussion on robustness, a different perspective on it would be to ask: given a PEF F with power $\beta > 0$ optimised with respect to the distribution ρ , how far in terms of total-variation distance can another distribution σ be such that the same PEF (with the same power) can be used to certify randomness? Theorem 7 provides a sufficient condition for the robustness of a positive, non-constant PEF $F = G^\beta$ with power β in the following sense: assuming the log-prob rate of F at ρ is positive, the log-prob rate of F at a different distribution σ is positive if $d_{TV}(\rho, \sigma)$ is within a certain bound (as given in (34)). For the sequence $\{F_k\}_{k=1}^\infty$ of PEFs the upper bound on $d_{TV}(\rho, \sigma)$ is computed as follows: Notice that in the sequence $\{F_k\}_{k=1}^\infty$ of PEFs, F_k is of the form $F_k = G_k^{\beta_k}$, where $G_k = 2^{K_* - e^{-k}}$. The upper bound on $d_{TV}(\rho, \sigma)$ (as given in (34)) is then $\mathbb{E}_\rho[G_k] / (L - 1) = \frac{1}{4} \left(\frac{S_\rho - 2}{2} - e^{-k} \right)$. It is worthwhile to observe that, given a standard CHSH–Bell inequality violating distribution ρ , this upper bound approaches the *strength of nonlocality* for ρ which is expressed as $(S_\rho - 2)/8$. The strength of nonlocality is defined in terms of how far the nonlocal no-signalling distribution ρ is from the local set Π_L [26]. It is defined as follows:

$$d_{NL}(\rho) := \frac{1}{|\mathcal{X}||\mathcal{Y}|} \frac{1}{2} \min_{\tau \in \Pi_L} \sum_{abxy} |\rho(ab|xy) - \tau(ab|xy)|, \quad (43)$$

where the minimum is over all distributions τ belonging to the local set Π_L . In the definition of $d_{NL}(\rho)$ in (43) we have assumed a uniform settings distribution as is evident from the factor $1/|\mathcal{X}||\mathcal{Y}|$, where $|\mathcal{X}|$ and $|\mathcal{Y}|$ denote the number of the measurement settings choices for Alice and Bob, respectively (which for the (2,2,2) Bell scenario is 2 for Alice and 2 for Bob). A theorem in [24] (see Theorem 3.1) provides a condition for the local distribution τ such that the minimum $(1/2) \min_{\tau \in \Pi_L} \sum_{abxy} |\rho(ab|xy) - \tau(ab|xy)|$ in (43) is achieved and that the minimum comes out to be the weight $(S_\rho - 2)/2$ on the PR box in the expression of ρ as the convex combination of the vertices of $\Delta_{PR,1}^8$; and so per the definition in (43) $d_{NL}(\rho) = (S_\rho - 2)/8$. Thus, the bound $\frac{1}{4} \left(\frac{S_\rho - 2}{2} - e^{-k} \right)$ from Theorem 7 approaches $\frac{S_\rho - 2}{8}$ which is the strength of nonlocality $d_{NL}(\rho)$ for ρ . This illustrates that a bound of this form cannot be improved, in the sense that increasing the total variation distance from ρ by any positive amount will encompass local distributions which cannot certify randomness.

Thus, $\{F_k\}_{k=1}^\infty$ is fully robust as $k \rightarrow \infty$. Next, we confirm that $\{F_k\}_{k=1}^\infty$ is asymptotically optimal in terms of min-entropy per trial (i.e., log-prob rate), for any distribution σ violating the standard CHSH inequality. Since Π_{NS} is closed and equal to the convex hull of its extremal points, Theorem 4 implies that, given such a σ , the adversary has a strategy obtained through an IID attack based on a single-trial distribution whose conditional Shannon entropy is equal to the infimum defined in (26). We can identify this attack. The optimisation in (26) can be expressed as follows:

$$h_{\min}(\sigma) = \min \left\{ H_\nu(AB|XYE) : \nu_e \in \Pi_{NS}, \sum_e \nu(e) \nu_e = \sigma \right\}, \quad (44)$$

where $\nu_e = \nu(ABXY|e)$. We compute $H_\nu(AB|XYE)$ for the decomposition of σ given in (41), where we have noted $\lambda_{PR,1} = (S_\sigma - 2)/2$. Since the conditional Shannon entropy is one for PR boxes and zero for LD behaviours, we obtain $H_\nu(AB|XYE) = (S_\sigma - 2)/2$ and, hence, $h_{\min}(\sigma)$ is no larger than this value. But since this expression is the same as that of the asymptotic log-prob rate of the sequence $\{F_k\}_{k=1}^\infty$ of valid PEFs, we can say $h_{\min}(\sigma)$ is also no smaller than this value and so $h_{\min}(\sigma) = (S_\sigma - 2)/2$. This demonstrates the asymptotic optimality of the sequence $\{F_k\}_{k=1}^\infty$ in the sense that the PEFs in the sequence become arbitrarily close to certifying an asymptotic randomness rate of $h_{\min}(\sigma)$.

In our proof of the asymptotic optimality of the sequence $\{F_k\}_{k=1}^\infty$, we identified the optimal attack by an adversary: it is to prepare the decomposition in (41) with each e corresponding to one of the (up to) nine extremal behaviours, with respective $\nu(e)$ weights of $\lambda_{PR,1}$ and α_i . This can be seen to be the *unique* attack achieving $h_{\min}(\sigma)$, through an argument we sketch as follows: (1) any ν -decomposition of σ can be improved upon (i.e., re-

ducing $H_v(AB|XYE)$ by considering only extremal v_e , by the concavity of conditional Shannon entropy; (2) any decomposition including positive weights on more than one PR box can be *strictly* improved upon by one with weights on a single PR box, by Theorem 2.1 of [24], which shows how to replace equal mixtures of two PR boxes with mixtures of a single PR box and local deterministic distributions; (3) this decomposition can be further strictly improved via Theorem 2.2 of [24] by removing any local deterministic distributions not saturating the CHSH–Bell inequality with those that do (the improvement being obtained by decreasing the weight on the sole remaining PR box). The resulting decomposition—that of (41)—is thus the unique optimiser of (44). It witnesses the bound of $1 + \dim \Pi_{\text{NS}} = 1 + 8 = 9$ on the set $\mathcal{E}$ (as shown in Theorem 4). In general, positive weight on all nine extremal boxes may be necessary, due to their affine independence which was noted in Section 4.1. One can confirm this visually from Table 1: weight on the (only) nonlocal distribution, the PR box, is necessary to violate the CHSH–Bell inequality and any distribution with non-zero probabilities for each possible outcome (a property possessed by, for example, the quantum distribution saturating Tsirelson’s bound) will require positive weight on all the local deterministic behaviours, as each LD behaviour corresponds to a distinct sole appearance of the number “1” in a column otherwise populated by zeroes in Table 1. This witnesses that further reduction of the $1 + \dim \Pi_{\text{NS}}$ bound on $|\mathcal{E}|$ in Theorem 4 is impossible and so this bound is optimal.

It is an important observation that the adversary needs to prepare only one non-classical state in her realisation of the optimal attack, since the preparation of a non-classical state is likely the most difficult aspect of the attack. We now explore possible generalisations of this feature to other trial models.

4.3. Characterising the Optimal Attack in Different Scenarios

We start by exploring the possibility of arriving at a similar analytic characterisation of the optimal adversarial attack when the adversary is limited to only quantum-realizable distributions. Suppose now that our trial model is the set Π_Q of quantum-achievable distributions for the (2,2,2) scenario. The adversary is still constrained to performing probabilistic attacks to simulate the trial statistics, while generating the least amount of randomness possible; however, she now tries to mimic the trial statistics using quantum-achievable distributions. The optimisation routine depicting this goal is:

$$\tilde{h}_{\min}(\sigma) = \min \left\{ H_{\omega}(AB|XYE) : \omega_e \in \Pi_Q, \sum_e \omega(e) \omega_e = \sigma \right\}, \quad (45)$$

where $\omega_e = \omega(ABXY|e)$. The set Π_Q is compact and convex, but, unlike Π_{NS} , is not a polytope and so there is a continuum of extremal points.

We conjecture that the minimum in (45) is achieved at a distribution that marginalises to the observed trial distribution through a convex combination of (only) one quantum extremal distribution violating the standard CHSH–Bell inequality and no more than eight local deterministic distributions that saturate the same inequality.

An attempt to prove this will require an understanding of the geometry of the quantum set and in particular its extremal points. We do not yet have a complete characterisation of the set of behaviours Π_Q (in the true $\mathbb{R}^8$ space), although a recent work has conjectured an analytic criterion for extremality in the CHSH scenario [27]. However, a characterisation does exist when we make the assumption of unbiased marginals: $\mu(A = 0|x) = \mu(A = 1|x) = 1/2$ for all $x \in \{0, 1\}$ and $\mu(B = 0|y) = \mu(B = 1|y) = 1/2$ for all $y \in \{0, 1\}$, in which case the set of behaviours is four dimensional. The unbiased marginal case has been completely characterised, a detailed exposition of which can be found in [25] (see Theorem 1 therein).

A key enabling step in the direction of characterising the optimal attack in the unbiased marginals case would be to see if the following two conditions hold simultaneously: first, a convex combination of any two extremal quantum behaviours can be expressed equivalently as a different convex combination of one extremal quantum behaviour (different

from the previous two) and classical noise (mixtures of the local deterministic behaviours), i.e., for extremal quantum behaviours $\vec{\mu}_1, \vec{\mu}_2$, the convex combination $\lambda\vec{\mu}_1 + (1-\lambda)\vec{\mu}_2$ can be re-expressed as the convex combination $\delta\vec{\mu}_3 + (1-\delta)\vec{\mu}_0$, where $\lambda, \delta \in (0,1)$, $\vec{\mu}_3$ is a third extremal quantum behaviour and $\vec{\mu}_0$ is a mixture of the local deterministic behaviours; and, second, $\lambda\mathbb{H}_{\mu_1}(AB|XY) + (1-\lambda)\mathbb{H}_{\mu_2}(AB|XY) \geq \delta\mathbb{H}_{\mu_3}(AB|XY)$, where the term $(1-\delta)\mathbb{H}_{\mu_0}(AB|XY)$ that might be expected to appear on the right vanishes due to the concavity of conditional Shannon entropy and the fact that it is zero for local deterministic behaviours, into which $\vec{\mu}_0$ can be decomposed.

A numerical inspection to check—by means of an exhaustive search—whether these two conditions hold simultaneously (in the uniform marginals case) introduces a lot of free variables. If we add more symmetry to the behaviours with uniform marginals and constrain ourselves to the two-dimensional slice as shown in Figure 4a, one can perform a numerical search to see whether the two conditions mentioned above hold simultaneously and we did observe it to hold in some initial numerical investigations comparing the $\vec{\theta}$ decompositions against the $\vec{v}$ decompositions as depicted in Figure 4b. The behaviours in the two-dimensional are represented by the formula:

$$\vec{\mu} = \frac{S}{4}\vec{\mu}_{PR,1} + \frac{S'}{4}\vec{\mu}_{PR,2} + \left(1 - \frac{S+S'}{4}\right)\vec{\mu}_0, \quad S, S' \in [-4, 4], \quad (46)$$

where $\vec{\mu}_0$ is the maximally random behaviour obtained as the equal mixtures of all 16 local deterministic behaviours. The disk $S^2 + (S')^2 \leq 8$ represents the set of quantum behaviours. Table 3 depicts a tabular representation of the behaviours expressible as (46). (As a side note, one way to add more symmetry to the behaviours with uniform marginals is as follows: a behaviour with uniform marginals can be completely specified by the correlators $(E_{00}, E_{01}, E_{10}, E_{11})$, where $-1 \leq E_{xy} \leq 1, \forall x, y$; see the line following (37) for the definition of E_{xy} . To obtain behaviours in the two-dimensional slice as shown in Figure 4a one can restrict attention to distributions of the form $\mu(ab|xy) = \frac{1}{4}(1 + (-1)^{a+b}C_{xy})$ where $C_{00} = -C_{11} = \frac{E_{00}-E_{11}}{2}$ and $C_{01} = C_{10} = \frac{E_{01}+E_{10}}{2}$.)

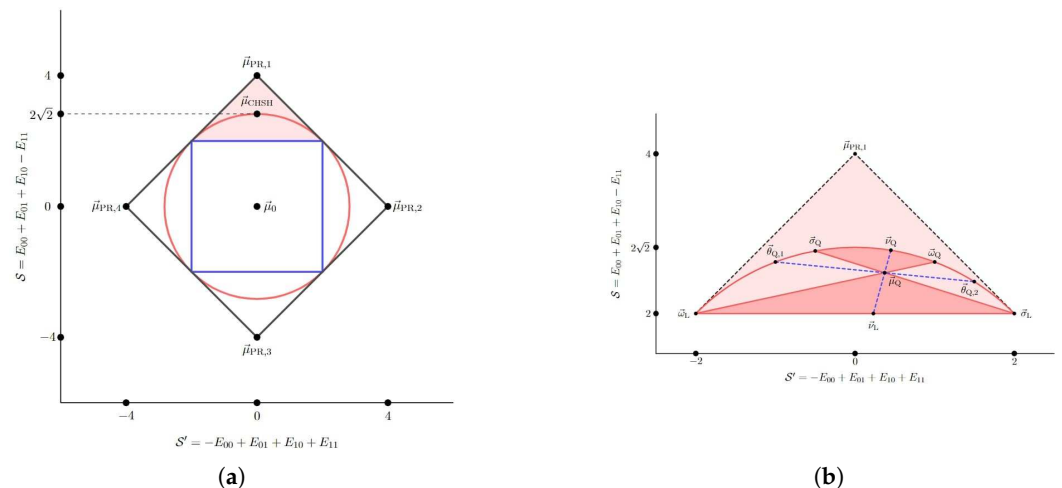


Figure 4. (a) A two-dimensional slice of the set of no-signalling behaviours (containing the quantum and the local set). The behaviours can be parameterised as the CHSH–Bell values S and S' obtained by two different versions of the CHSH–Bell expression in (37). Any behaviour on the slice can be represented as in (46). (b) The portion of the two-dimensional slice containing the no-signalling (including quantum-achievable) behaviours above the standard CHSH–Bell facet. For a fixed behaviour $\vec{\mu}_Q$ in the interior of the quantum region, the darker shaded region corresponds to possible ways of expressing $\vec{\mu}_Q$ as a convex combination of a behaviour on the quantum boundary and a behaviour on the local boundary (for example, $\vec{\mu}_Q = \lambda\vec{v}_Q + (1-\lambda)\vec{v}_L, \lambda \in (0,1)$). For the same behaviour $\vec{\mu}_Q$, the lighter shaded region represents possible ways of expressing it as a convex combination of two behaviours on the quantum boundary (for example, $\vec{\mu}_Q = \delta\vec{\theta}_{Q,1} + (1-\delta)\vec{\theta}_{Q,2}, \delta \in (0,1)$).

Going beyond the minimal Bell scenario, we considered the possibility of a similar characterisation of optimal *no-signalling* adversarial attack in higher (n, m, k) Bell scenarios. In the $(2, 2, 2)$ Bell scenario the analytical characterisation of the optimal adversarial attack crucially relied upon the geometric features of the no-signalling polytope, namely Theorems 2.1 and 2.2 in [24]: that equal mixtures of two PR behaviours are expressible as equal mixtures of four distinct LD behaviours and, consequently, a behaviour violating any of the eight versions (up to local relabelling of the outcomes and settings) of the CHSH–Bell inequality is expressible as a convex combination of the one PR behaviour achieving the nonlocal maximum and (up to) eight LD behaviours achieving the local maximum of the corresponding CHSH–Bell expression. These geometric features, however, do not extend to the no-signalling polytopes of higher (n, m, k) Bell scenarios. Membership of equal mixtures of extremal no-signalling nonlocal behaviours in the local polytope holds solely in the $(2, 2, 2)$ Bell scenario.

Table 3. Tabular representation of the no-signalling behaviours on the two-dimensional slice shown in Figure 4a. The behaviours have uniform marginals, i.e., the probability of observing an outcome conditioned on a measurement setting is $1/2$ for each party for all outcomes and settings. The behaviours are further constrained in having the third and fourth row completely determined by the first and second, which need not hold in general for uniform marginal distributions, and brings the dimensionality down from four to two. Any behaviour represented as above is parameterised as the values S and S' of the two versions of the CHSH–Bell expression $E_{00} + E_{01} + E_{10} - E_{11}$ and $-E_{00} + E_{01} + E_{10} + E_{11}$, respectively: $s_1 = (4 + S - S')/16$, $s_2 = (4 + S' - S)/16$, $s_3 = (4 + S + S')/16$, $s_4 = (4 - S - S')/16$, where for the no-signalling set $-4 \leq S' + S \leq 4$, $-4 \leq S' - S \leq 4$ and for the quantum set $S^2 + (S')^2 \leq 8$.

		<i>ab</i>			
		00	01	10	11
<i>xy</i>	00	s_1	s_2	s_2	s_1
	01	s_3	s_4	s_4	s_3
	10	s_3	s_4	s_4	s_3
	11	s_2	s_1	s_1	s_2

Below, we provide examples of equal mixtures of no-signalling nonlocal extremal behaviours in the $(2, 2, 3)$, $(2, 3, 2)$ and $(3, 2, 2)$ Bell scenarios that do not belong to the local polytope. One can use linear programming to check the nonlocality of such examples. Assessment of the locality of a behaviour is an instance of the *membership problem of the local polytope*. Since the local deterministic (LD) behaviours are the extremal points of the local polytope, we can formulate our problem as a *feasibility linear program*. Suppose $\{\vec{\mu}_{LD,1}, \vec{\mu}_{LD,2}, \dots, \vec{\mu}_{LD,\#LD}\}$ is the set of LD behaviours for some Bell scenario. The vector $\vec{\mu}_{LD,i} \in \mathbb{R}^d$ denotes the joint probability of outcomes conditioned on the input choices and d is the dimension of the ambient space in which the vector lies. The feasibility linear program has the variable $\vec{x} \in \mathbb{R}^{\#LD}$. The inequality constraints comprise $x_i \geq 0$, $i \in [\#LD]$ and the equality constraints are $\sum_{i=1}^{\#LD} x_i = 1$ and the following:

$$\frac{1}{2} \vec{NS}_{\text{extr}} + \frac{1}{2} \vec{NS}'_{\text{extr}} = \sum_{k=1}^{\#LD} x_k \vec{\mu}_{LD,k} \quad (47)$$

where $\vec{NS}_{\text{extr}}$ is a nonlocal no-signalling extremal behaviour. The details on formulating the dual of this linear program can be found in section E.2.1 of the Appendix of [6].

Before presenting the counter-examples we briefly review the (n, m, k) Bell scenario: This scenario consists of n spatially separated parties, where each party $i \in [n]$ has a choice of m different k -outcome measurements. For $\mathcal{X} \equiv \{0, 1, \dots, m-1\}$ and $\mathcal{A} \equiv \{0, 1, \dots, k-1\}$ the joint probability $\mu(a_1 a_2 \dots a_n | x_1 x_2 \dots x_n)$ of obtaining the outcomes $(a_1, a_2, \dots, a_n) \in \mathcal{A}^n$

conditioned on the inputs $(x_1, x_2, \dots, x_n) \in \mathcal{X}^n$ can be viewed as a probability vector $\vec{\mu} \in \mathbb{R}^d$, where $d = (|\mathcal{A}||\mathcal{X}|)^n$.

The extremal points of the no-signalling polytope comprise the local deterministic (LD) behaviours and the nonlocal extremal behaviours. The LD behaviours consist of all possible assignments $\Lambda_{\text{LD}} = \{\{\lambda_{1x_1}\}_{x_1 \in \mathcal{X}}; \{\lambda_{2x_2}\}_{x_2 \in \mathcal{X}}; \dots; \{\lambda_{nx_n}\}_{x_n \in \mathcal{X}}\}$, where $\lambda_{ix_i} \in \mathcal{A}$ for $i \in [n]$. The number of such assignments is $\#\text{LD} = (|\mathcal{A}|)^{n|\mathcal{X}|}$. Corresponding to each assignment $\lambda \in \Lambda_{\text{LD}}$ the LD probabilities are expressed as

$$\mu_{\text{LD},k}(a_1 a_2 \dots a_n | x_1 x_2 \dots x_n) = \llbracket a_1 = \lambda_{1x_1} \rrbracket \llbracket a_2 = \lambda_{2x_2} \rrbracket \dots \llbracket a_n = \lambda_{nx_n} \rrbracket \quad (48)$$

where $\llbracket \cdot \rrbracket$ is the function that evaluates to 1 if the condition within holds, 0 otherwise. A behaviour $\vec{\mu}_{\mathcal{L}}$ is local if it can be expressed as $\vec{\mu}_{\mathcal{L}} = \sum_{k=1}^{\#\text{LD}} q_k \vec{\mu}_{\text{LD},k}$, where $q_k \geq 0$ and $\sum_{k=1}^{\#\text{LD}} q_k = 1$.

(2,2,3) Bell scenario: This scenario is an instance of the more general $(2,2,k)$ scenario, also known in the literature as the CGLMP scenario [28], for $k = 3$. In this bipartite scenario the parties have two three-output choices of settings. The extremal behaviours for the no-signalling polytope for the CGLMP scenario have been fully described in [29]. The nonlocal no-signalling extremal behaviours for the $(2,2,3)$ scenario, up to relabelling of inputs and outcomes, are given by the following formula:

$$\text{NL}_{\text{ext}}(ab|xy) := \begin{cases} \frac{1}{3} & : b - a \equiv xy \pmod{3} \\ 0 & : \text{otherwise} \end{cases} \quad (49)$$

where $a, b \in \{0, 1, 2\}$ and $x, y \in \{0, 1\}$ are the outputs and inputs for the parties, respectively. We found that (47) does not necessarily hold for all equal mixtures of a pair of distinct nonlocal extremal behaviours. Among the several examples we found that violate (47), Table 4 shows one such example.

Table 4. Two nonlocal extremal behaviours for the CGLMP scenario with 3 outcomes whose equal mixtures are nonlocal. The inputs $x, y \in \{0, 1\}$ and the outcomes $a, b \in \{0, 1, 2\}$ with $x' = x \oplus 1$, $y' = y \oplus 1$ and $a' = a \oplus_3 1$, $a'' = a \oplus_3 2$, $b' = b \oplus_3 1$, $b'' = b \oplus_3 2$. The symbol $\oplus$ denotes addition modulo 2 and $\oplus_3$ denotes addition modulo 3. The missing entries correspond to 0. The top behaviour comes directly from (49) while the bottom behaviour is obtained through the relabelling $x \leftrightarrow x'$ and $y \leftrightarrow y'$. An equal mixture of these two boxes lies outside the local polytope.

	ab	ab'	ab''	$a'b$	$a'b'$	$a'b''$	$a''b$	$a''b'$	$a''b''$
xy	1/3				1/3				1/3
xy'	1/3				1/3				1/3
$x'y$	1/3				1/3				1/3
$x'y'$		1/3				1/3	1/3		
xy		1/3				1/3	1/3		
xy'	1/3				1/3				1/3
$x'y$	1/3				1/3				1/3
$x'y'$	1/3				1/3				1/3

(2,3,2) Bell scenario: More generally, the extremal behaviours of $(2,k,2)$ no-signalling polytope, with $k > 2$, have been completely characterised in [30,31], of which $(2,3,2)$ is an instance. Following Table II of [31], we can obtain Tables 5 and 6 which are two representative examples of nonlocal no-signalling extremal behaviours, equal mixtures of which lie outside the local polytope. In Table 5 all input choices, $x, y \in \{0, 1, 2\}$, for Alice and Bob have uniform probabilities of outcomes; in Table 6 all inputs for Alice and inputs

$y \in \{0, 1\}$ for Bob have uniform probabilities of outcomes, with the exception that Bob's outcome for $y = 2$ is deterministic.

$$\begin{bmatrix} p(00|xy) & p(01|xy) \\ p(10|xy) & p(11|xy) \end{bmatrix} \quad \text{with} \quad \begin{bmatrix} ? \\ ? \end{bmatrix} \equiv \begin{bmatrix} 1/2 & 0 \\ 0 & 1/2 \end{bmatrix} \text{ or } \begin{bmatrix} 0 & 1/2 \\ 1/2 & 0 \end{bmatrix}$$

There are 16 possible mixtures of the two behaviours in Tables 5 and 6 corresponding to each '?' in each table being a perfect correlation or a perfect anti-correlation, all of which represent mixtures of extremal nonlocal boxes [31] and all lie outside the local polytope. The nonlocality of the mixtures is confirmed by noting that the four cells in the upper left corner, corresponding to restricting the settings choices to $x, y \in \{0, 1\}$, is the PR box distribution which is of course nonlocal.

(3, 2, 2) *Bell scenario*: This is a tripartite scenario with each party having binary input choices and outcomes. The no-signalling polytope consists of 46 inequivalent classes of extremal behaviours, of which one is the class comprising 64 LD behaviours. A complete characterisation can be found in [32]. As an example violating (47) we can refer to the observation made in Section 2.3 of [32] that equal mixtures of two behaviours in Class 46 (see Table 1 of [32]) are a GHZ correlation which is expressed (entirely in terms of correlators $\langle A_x B_y C_z \rangle$) as $P_{\text{GHZ}}(abc|xyz) = \frac{1}{8}(a + abc \langle A_x B_y C_z \rangle)$. $\vec{P}_{\text{GHZ}}$ is a nonlocal behaviour which is obtained by measuring $\frac{1}{\sqrt{2}}(|000\rangle + |111\rangle)$ in suitable local bases [33].

Table 5. Nonlocal no-signalling extremal behaviour with all input choices $x, y \in \{0, 1, 2\}$ for Alice and Bob having uniform probabilities of outcomes.

		y					
		0		1		2	
x	0	$\frac{1}{2}$ 0	0 $\frac{1}{2}$	$\frac{1}{2}$ 0	0 $\frac{1}{2}$	$\frac{1}{2}$ 0	0 $\frac{1}{2}$
	1	$\frac{1}{2}$ 0	0 $\frac{1}{2}$	0 $\frac{1}{2}$	$\frac{1}{2}$ 0	?	
	2	$\frac{1}{2}$ 0	0 $\frac{1}{2}$	?		?	

Table 6. All inputs for Alice and inputs $y \in \{0, 1\}$ for Bob have uniform probabilities of outcomes, while Bob's outcome for $y = 2$ is deterministic.

		y					
		0		1		2	
x	0	1/2	0	1/2	0	1/2	0
		0	1/2	0	1/2	1/2	0
	1	1/2	0	0	1/2	1/2	0
		0	1/2	1/2	0	1/2	0
	2	1/2	0	?		1/2	0
		0	1/2			1/2	0

5. Conclusions

In this work, we revisited the probability estimation framework with the goal of presenting a complete and self-contained proof of its optimality in the asymptotic regime and obtaining a better characterisation of optimal adversarial attack strategies on the protocol. We obtained in Theorem 4 an improved and tight upper bound on the cardinality of the set of states needed in the optimal attack, and studied the implications of this result for specific scenarios in Section 4. We also considered the question of *robustness* for the PEF method, finding that asymptotic optimality of PEFs (in terms of randomness

generation rate) need not entail a trade-off with robustness to small deviations from expected experimental behaviour.

In proving the optimality of the framework, our results show that there remains nothing to be gained, asymptotically, for an adversary implementing memory attacks—an i.i.d. attack is asymptotically optimal. However, in real world applications this may not hold. The number of trials in a Bell experiment is finite, albeit large, and there are unavoidable correlations between the successive trials (referred to as memory effects). We leave to future work considerations of side-channel attacks in the non-asymptotic (finite trials) regime for the probability estimation framework.

Author Contributions: Conceptualisation, S.P. and P.B.; Formal analysis, S.P. and P.B.; Funding acquisition, P.B.; Investigation, S.P. and P.B.; Methodology, S.P. and P.B.; Software, S.P.; Supervision, P.B.; Writing—original draft, S.P. and P.B.; Writing—review and editing, S.P. and P.B. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by AFOSR Grant FA9550-20-1-0067, NSF Award 1839223 and Louisiana Board of Regents Award LEQSF (2019-22)-RD-A-27. The APC was funded by NSF Grant 2210399.

Institutional Review Board Statement: Not applicable.

Data Availability Statement: Not applicable.

Acknowledgments: We acknowledge helpful discussions with Jitendra Prakash and Mark Wilde.

Conflicts of Interest: The authors declare no conflict of interest.

Appendix A

Appendix A.1. Proofs for Theorems 1 and 2

Appendix A.1.1. Proof for Theorem 1

Theorem. Suppose $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ is a distribution of **CZE** such that $\mu_e(\mathbf{CZ}) \in \Theta$ for each $e \in \mathcal{E}$. Then, for fixed $\beta, \epsilon > 0$

$$\mathbb{P}_{\mu_e} \left(\mu_e(\mathbf{C}|\mathbf{Z}) \geq \left(\epsilon \prod_{i=1}^n F_i(C_i Z_i) \right)^{-1/\beta} \right) \leq \epsilon \quad (\text{A1})$$

holds for each $e \in \mathcal{E}$, where $F_i(C_i Z_i)$ is the probability estimation factor for the i 'th trial.

Proof. The sequence of random variables $\mathbf{C}, \mathbf{Z}$ represent the time-ordered sequence of n trial results. For the remainder of the proof we omit conditioning on $E = e$ since the result holds for each realisation. Hence, $\mu(\cdots)$, $\mathbb{P}_\mu(\cdots)$ and $\mathbb{E}_\mu[\cdots]$ must be understood to mean $\mu_e(\cdots)$, $\mathbb{P}_{\mu_e}(\cdots)$ and $\mathbb{E}_{\mu_e}[\cdots]$.

Observe that for any $i \in \{1, \dots, n-1\}$ we have

$$\mu(\mathbf{C}_{\leq i+1} | \mathbf{Z}_{\leq i+1}) = \mu(C_{i+1} | \mathbf{C}_{\leq i} \mathbf{Z}_{\leq i+1}) \mu(\mathbf{C}_{\leq i} | Z_{i+1} \mathbf{Z}_{\leq i}) = \mu(C_{i+1} | \mathbf{C}_{\leq i} \mathbf{Z}_{\leq i+1}) \mu(\mathbf{C}_{\leq i} | \mathbf{Z}_{\leq i}) \quad (\text{A2})$$

where the first equality is an elementary manipulation of conditional probabilities and the second equality follows from

$$\begin{aligned} \mu(\mathbf{C}_{\leq i} | Z_{i+1} \mathbf{Z}_{\leq i}) &= \mu(\mathbf{C}_{\leq i} Z_{i+1} \mathbf{Z}_{\leq i}) / \mu(Z_{i+1} \mathbf{Z}_{\leq i}) \\ &= \mu(\mathbf{C}_{\leq i} \mathbf{Z}_{\leq i}) \mu(Z_{i+1}) / \mu(Z_{i+1}) \mu(\mathbf{Z}_{\leq i}) \\ &= \mu(\mathbf{C}_{\leq i} | \mathbf{Z}_{\leq i}), \end{aligned}$$

with the second step above following from the second condition in (1), applied directly in the numerator and in the denominator via

$$\mu(z_{j+1}\mathbf{z}_{\leq j}) = \sum_{\mathbf{c}_{\leq j}} \mu(z_{j+1}\mathbf{c}_{\leq j}\mathbf{z}_{\leq j}) = \mu(z_{j+1}) \sum_{\mathbf{c}_{\leq j}} \mu(\mathbf{c}_{\leq j}\mathbf{z}_{\leq j}) = \mu(z_{j+1})\mu(\mathbf{z}_{\leq j}).$$

Now, consider the sequence $Q_i = \mu(\mathbf{C}_{\leq i}|\mathbf{Z}_{\leq i})^\beta \prod_{j=1}^i F_j$, for $i \geq 1$, where we note Q_i is a random variable that is determined by $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$. We begin by showing that conditioned on $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$ the expectation of Q_{i+1} is at most Q_i for all $i \in \{1, \dots, n-1\}$. Applying (A2), we can write

$$\begin{aligned} Q_{i+1} &= F_{i+1} \mu(C_{i+1}|Z_{i+1}\mathbf{C}_{\leq i}\mathbf{Z}_{\leq i})^\beta \mu(\mathbf{C}_{\leq i}|\mathbf{Z}_{\leq i})^\beta \prod_{j=1}^i F_j \\ &= F_{i+1} \mu(C_{i+1}|Z_{i+1}\mathbf{C}_{\leq i}\mathbf{Z}_{\leq i})^\beta Q_i, \\ \Rightarrow \mathbb{E}_\mu[Q_{i+1}|\mathbf{C}_{\leq i}\mathbf{Z}_{\leq i}] &= Q_i \mathbb{E}_\mu \left[F_{i+1} \mu(C_{i+1}|Z_{i+1}\mathbf{C}_{\leq i}\mathbf{Z}_{\leq i})^\beta \middle| \mathbf{C}_{\leq i}\mathbf{Z}_{\leq i} \right] \leq Q_i \end{aligned} \quad (\text{A3})$$

where the fact that Q_i is determined by $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$ allows us to pull it out of the conditional expectation and the inequality follows from the fact that $\mathbb{E}_\mu[F_{i+1} \mu(C_{i+1}|Z_{i+1}\mathbf{c}_{\leq i}\mathbf{z}_{\leq i})^\beta] \leq 1$ for all realisations $\mathbf{c}_{\leq i}, \mathbf{z}_{\leq i}$ of $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$, as ensured by Definition 1. We remark that Q_i is a supermartingale as indicated by the inequality in (A3): the term $F_i \mu(C_i|Z_i\mathbf{Z}_{\leq i-1}\mathbf{C}_{\leq i-1})^\beta$ is non-negative, is determined by $\mathbf{C}_{\leq i}, \mathbf{Z}_{\leq i}$ and satisfies $\mathbb{E}_\mu[F_i \mu(C_i|Z_i\mathbf{C}_{\leq i-1}\mathbf{Z}_{\leq i-1})^\beta | \mathbf{C}_{\leq i-1}\mathbf{Z}_{\leq i-1}] \leq 1$. Now, using the law of iterated expectation we obtain:

$$\mathbb{E}_\mu[Q_{i+1}] = \mathbb{E}_\mu[\mathbb{E}_\mu[Q_{i+1}|\mathbf{C}_{\leq i}\mathbf{Z}_{\leq i}]] \leq \mathbb{E}_\mu[Q_i] \quad (\text{A4})$$

Since Q_1 equals $\mu(C_1|Z_1)^\beta F_1$, it satisfies $\mathbb{E}_\mu[Q_1] \leq 1$ directly from Definition 1 and so repeated applications of (A4) yield $\mathbb{E}_\mu[Q_n] \leq \mathbb{E}_\mu[Q_{n-1}] \leq \dots \leq \mathbb{E}_\mu[Q_1] \leq 1$. Since $Q_n = \mu(\mathbf{C}|\mathbf{Z})^\beta \prod_{i=1}^n F_i$ is non-negative, we can use Markov's inequality and obtain the required result as shown below.

$$\begin{aligned} \mathbb{P}_\mu \left(\mu(\mathbf{C}|\mathbf{Z})^\beta \prod_{i=1}^n F_i \geq 1/\epsilon \right) &\leq \epsilon \mathbb{E}_\mu \left[\mu(\mathbf{C}|\mathbf{Z})^\beta \prod_{i=1}^n F_i \right] \leq \epsilon \\ \Rightarrow \mathbb{P}_\mu \left(\mu(\mathbf{C}|\mathbf{Z}) \geq \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \right) &\leq \epsilon. \end{aligned}$$

□

Appendix A.1.2. Proof for Theorem 2

Theorem. Let μ be a distribution $\mu: \mathcal{C}^n \times \mathcal{Z}^n \times \mathcal{E} \rightarrow [0, 1]$ of CZE such that, for each $e \in \mathcal{E}$, the following holds for every $\epsilon \in (0, 1)$:

$$\mathbb{P}_{\mu_e} \left(\mu_e(\mathbf{C}|\mathbf{Z}) \leq \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \right) \geq 1 - \epsilon, \quad (\text{A5})$$

where F_i is a PEF with power β for the i 'th trial. For a fixed choice of $\epsilon \in (0, 1)$ and $p \geq |\mathcal{C}|^{-n}$, define the event $S := \left\{ \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \leq p \right\}$. Then, if κ is a positive number for which $\mathbb{P}_\mu(S) \geq \kappa$, the following holds:

$$\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon/\kappa}(\mathbf{C}|\mathbf{Z}\mathbf{E}; S) \geq \log_2(\kappa) - \log_2(p) \quad (\text{A6})$$

Proof. The goal is to construct a distribution ω of CZE such that it is within ϵ/κ TV distance from $\mu(\mathbf{C}|\mathbf{Z}\mathbf{E}|S)$ and such that the average conditional maximum probability of $\mathbf{C}$ conditioned on (and averaged over) $\mathbf{Z}\mathbf{E}$ is bounded below by p/κ . We will construct ω to satisfy $\omega(\mathbf{C}ze|S) = 0$ for all values of $\mathbf{z}$ and e for which $\mu(\mathbf{z}e|S) = 0$. Hence, for the

rest of the construction, we will restrict attention to cases where $\mu(\mathbf{ze}|\mathbf{S}) > 0$. We will use expressions such as $\mathbb{P}_{\mu_e}(\mathbf{S})$ and $\mu_e(\mathbf{S})$ interchangeably.

We start by defining the event

$$\mathbf{R} := \left\{ \mu_e(\mathbf{C}|\mathbf{Z}) \leq \left(\epsilon \prod_{i=1}^n F_i \right)^{-1/\beta} \right\},$$

whose occurrence or non-occurrence is determined by the particular realisation of e , $\mathbf{c}$ and $\mathbf{z}$. The event $\mathbf{R}$ corresponds to the desired probability bound holding; (A5) ensures that this event occurs with high probability and we will construct our distribution ω to, in an intuitive sense, extend this desirable behaviour from $\mathbf{R} \cap \mathbf{S}$ to all of $\mathbf{S}$.

We begin the construction by defining, for each fixed e satisfying $\mu_e(\mathbf{S}) > 0$, a non-negative function $f: \mathcal{C}^n \times \mathcal{Z}^n \rightarrow \mathbb{R}^+$ as shown below.

$$f(\mathbf{cz}) = \begin{cases} \mu_e(\mathbf{cz})/\mathbb{P}_{\mu_e}(\mathbf{S}), & \mathbf{S} \cap \mathbf{R} \text{ holds} \\ 0, & \text{otherwise} \end{cases} \quad (\text{A7})$$

The weight w of f , defined as $w(f) = \sum_{\mathbf{c}, \mathbf{z}} f(\mathbf{cz})$, satisfies $w(f) \leq 1$ as shown below:

$$w(f) = \sum_{\mathbf{c}, \mathbf{z}} f(\mathbf{cz}) = \sum_{\mathbf{c}, \mathbf{z}} \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \cap \mathbf{R} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) \leq \sum_{\mathbf{c}, \mathbf{z}} \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) = \mathbb{P}_{\mu_e}(\mathbf{S}) / \mathbb{P}_{\mu_e}(\mathbf{S}) = 1,$$

where $\llbracket \cdot \rrbracket$ is equal to 1, if the condition or expression within holds, 0 otherwise. (Note that f is a *sub-probability distribution* on $\mathbf{cz}$: a set of non-negative numbers whose sum is less than or equal to 1. Defining a sub-probability distribution is a standard trick to construct a distribution by invoking certain lemmas.) Below we show that w satisfies $w(f) \geq 1 - \epsilon / \mathbb{P}_{\mu_e}(\mathbf{S})$.

$$\begin{aligned} w(f) &= 1 - 1 + \sum_{\mathbf{c}, \mathbf{z}} \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket \llbracket \mathbf{R} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) \\ &= 1 - \sum_{\mathbf{c}, \mathbf{z}} \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) + \sum_{\mathbf{c}, \mathbf{z}} \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket \llbracket \mathbf{R} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) \\ &= 1 - \sum_{\mathbf{c}, \mathbf{z}} (\mu_e(\mathbf{cz}) - \mu_e(\mathbf{cz}) \llbracket \mathbf{R} \rrbracket) \llbracket \mathbf{S} \rrbracket / \mathbb{P}_{\mu_e}(\mathbf{S}) \geq 1 - \sum_{\mathbf{c}, \mathbf{z}} (\mu_e(\mathbf{cz}) - \mu_e(\mathbf{cz}) \llbracket \mathbf{R} \rrbracket) / \mathbb{P}_{\mu_e}(\mathbf{S}) \\ &= 1 - (1 - \mathbb{P}_{\mu_e}(\mathbf{R})) / \mathbb{P}_{\mu_e}(\mathbf{S}) \geq 1 - \epsilon / \mathbb{P}_{\mu_e}(\mathbf{S}), \end{aligned} \quad (\text{A8})$$

where in (A8) we have used the fact that $\mathbb{P}_{\mu_e}(\mathbf{R}) \geq 1 - \epsilon$ holds for each $e \in \mathcal{E}$, as follows from (A5). Next, we define a non-negative function $\tilde{f}_{\mathbf{z}}: \mathcal{C}^n \rightarrow \mathbb{R}^+$ for each $\mathbf{z} \in \mathcal{Z}^n$ for which $\mu_e(\mathbf{z}|\mathbf{S}) > 0$:

$$\tilde{f}_{\mathbf{z}}(\mathbf{c}) = f(\mathbf{cz}) / \mu_e(\mathbf{z}|\mathbf{S}) \quad (\text{A9})$$

We show below that, for each such $\mathbf{z}$, $\tilde{f}_{\mathbf{z}}(\mathbf{c})$ is bounded by $\mu_e(\mathbf{c}|\mathbf{z}, \mathbf{S})$, $\forall \mathbf{c} \in \mathcal{C}^n$. We have:

$$\begin{aligned} \tilde{f}_{\mathbf{z}}(\mathbf{c}) &= \mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket \llbracket \mathbf{R} \rrbracket / (\mathbb{P}_{\mu_e}(\mathbf{S}) \mu_e(\mathbf{z}|\mathbf{S})) = \mu_e(\mathbf{cz}, \mathbf{S}) \llbracket \mathbf{R} \rrbracket / \mu_e(\mathbf{z}, \mathbf{S}) \\ &\leq \mu_e(\mathbf{cz}, \mathbf{S}) / \mu_e(\mathbf{z}, \mathbf{S}) = \mu_e(\mathbf{c}|\mathbf{z}, \mathbf{S}), \end{aligned} \quad (\text{A10})$$

where the equality $\mu_e(\mathbf{cz}) \llbracket \mathbf{S} \rrbracket = \mu_e(\mathbf{cz}, \mathbf{S})$ makes sense because whether or not $\mathbf{S}$ holds is determined by $\mathbf{cz}$. Since (A10) holds for all $\mathbf{c}$, we conclude $\tilde{f}_{\mathbf{z}}(\mathbf{C}) \leq \mu_e(\mathbf{C}|\mathbf{z}, \mathbf{S})$. This proves that $\tilde{f}_{\mathbf{z}}(\mathbf{C})$ is dominated by $\mu_e(\mathbf{C}|\mathbf{z}, \mathbf{S})$. From the definition of $\tilde{f}_{\mathbf{z}}$ we also have another upper bound for all $\mathbf{c}$:

$$\tilde{f}_{\mathbf{z}}(\mathbf{c}) = \mu_e(\mathbf{c}|\mathbf{z}) \mu_e(\mathbf{z}) \llbracket \mathbf{S} \cap \mathbf{R} \rrbracket / \mu_e(\mathbf{z}, \mathbf{S}) \leq p \mu_e(\mathbf{z}) / \mu_e(\mathbf{z}, \mathbf{S}).$$

Above, we have used the fact that the event $\mathbf{S} \cap \mathbf{R}$ implies $\mu_e(\mathbf{C}|\mathbf{Z}) \leq (\epsilon \prod_{i=1}^n F_i)^{-1/\beta} \leq p$. The bound $p \mu_e(\mathbf{z}) / \mu_e(\mathbf{z}, \mathbf{S}) \geq p \geq |\mathcal{C}|^{-n}$ also holds, since $\mu_e(\mathbf{z}) / \mu_e(\mathbf{z}, \mathbf{S}) \geq 1$. Hence, using the lemmas in Appendix D we can construct, for each $\mathbf{z}$ under consideration, a distribution

$\mu'_z(\mathbf{C})$ such that $\mu'_z(\mathbf{C}) \geq \tilde{f}_z(\mathbf{C})$, $\mu'_z(\mathbf{C}) \leq p\mu_e(\mathbf{z})/\mu_e(\mathbf{z}, S)$ and $\text{TV}(\mu'_z(\mathbf{C}), \mu_e(\mathbf{C}|\mathbf{z}, S)) \leq 1 - w(\tilde{f}_z)$, where $w(\tilde{f}_z) \leq 1$ is the weight of $\tilde{f}_z(\mathbf{C})$. Now we are ready to define the distribution $\omega(\mathbf{CZE})$ as

$$\omega(\mathbf{cze}) = \begin{cases} \mu'_z(\mathbf{c})\mu_e(\mathbf{z}|S)\mu(e|S) & \text{if } \mu(\mathbf{ze}|S) > 0 \\ 0 & \text{if } \mu(\mathbf{ze}|S) = 0 \end{cases}$$

We show that the total variation distance between ω and $\mu(\mathbf{CZE}|S)$ is bounded by ϵ/κ and that the average $\mathbf{ze}$ -conditional maximum probability of $\mathbf{C}$ is bounded by p/κ . First,

$$\begin{aligned} d_{\text{TV}}(\omega(\mathbf{CZE}), \mu(\mathbf{CZE}|S)) &= \frac{1}{2} \sum_{\mathbf{cze}} |\omega(\mathbf{cze}) - \mu(\mathbf{cze}|S)| \\ &= \frac{1}{2} \sum_{e:\mu(e|S)>0} \sum_{\mathbf{z}:\mu_e(\mathbf{z}|S)>0} \sum_{\mathbf{c}} |\mu'_z(\mathbf{c}) - \mu(\mathbf{c}|\mathbf{ze}, S)| \mu_e(\mathbf{z}|S)\mu(e|S) \end{aligned} \quad (\text{A11})$$

$$\leq \frac{1}{2} \sum_{e:\mu(e|S)>0} \sum_{\mathbf{z}:\mu_e(\mathbf{z}|S)>0} \sum_{\mathbf{c}} (\mu'_z(\mathbf{c}) - \tilde{f}_z(\mathbf{c}) + \mu(\mathbf{c}|\mathbf{ze}, S) - \tilde{f}_z(\mathbf{c})) \mu_e(\mathbf{z}|S)\mu(e|S) \quad (\text{A12})$$

$$= \frac{1}{2} \sum_{e:\mu(e|S)>0} \sum_{\mathbf{z}:\mu_e(\mathbf{z}|S)>0} \mu_e(\mathbf{z}|S)\mu(e|S) \left(1 - \sum_{\mathbf{c}} \tilde{f}_z(\mathbf{c}) + 1 - \sum_{\mathbf{c}} \tilde{f}_z(\mathbf{c}) \right) \quad (\text{A13})$$

$$= \frac{1}{2} \sum_{e:\mu(e|S)>0} 2 \left(1 - \sum_{\mathbf{cz}} f(\mathbf{cz}) \right) \mu(e|S) \quad (\text{A14})$$

$$= \sum_{e:\mu(e|S)>0} (1 - w(f)) \mu(e|S) \leq \sum_{e:\mu(e|S)>0} \frac{\epsilon}{\mathbb{P}_{\mu_e}(S)} \mu(e|S) = \sum_{e:\mu(e|S)>0} \frac{\epsilon \mu(e)}{\mu(S)} \leq \epsilon / \mathbb{P}_{\mu}(S) \leq \epsilon / \kappa \quad (\text{A15})$$

The equality in (A11) follows because $\omega(\mathbf{cze}) = \mu(\mathbf{cze}|S) = 0$ for values of $e, \mathbf{z}$ removed from the sums and $\mu'_z(\mathbf{c})$ is defined for the remaining values of $e, \mathbf{z}$. In (A12) we add and subtract with $\tilde{f}_z(\mathbf{c})$ inside the absolute value expression in the previous step and use the triangle inequality, following which we use the facts established above that both $\mu'_z(\mathbf{C})$ and $\mu(\mathbf{C}|\mathbf{ze}, S) = \mu_e(\mathbf{C}|\mathbf{z}, S)$ dominate $\tilde{f}_z(\mathbf{C})$; (A13) follows from the fact that $\mu'_z(\mathbf{c})\mu_e(\mathbf{z}|S)$ and $\mu(\mathbf{c}|\mathbf{ze}, S)$ sum to 1 over $\mathbf{c}$ (being distributions), and (A14) follows from $\tilde{f}_z(\mathbf{c}) = f(\mathbf{cz})/\mu_e(\mathbf{z}|S)$ and the fact that $f(\mathbf{cz}) = 0$ in cases where $\mu_e(\mathbf{z}|S) = 0$. Finally, the first inequality in (A15) follows from (A8) and the last inequality follows from $\mathbb{P}_{\mu}(S) \geq \kappa$. Next, we show the upper bound on the average conditional maximum probability.

$$\begin{aligned} \sum_{\mathbf{ze}} \left(\max_{\mathbf{c}} \omega(\mathbf{c}|\mathbf{ze}) \right) \omega(\mathbf{ze}) &= \sum_{\mathbf{ze}} \left(\max_{\mathbf{c}} \mu'_z(\mathbf{c}) \right) \mu_e(\mathbf{z}|S)\mu(e|S) \\ &\leq \sum_{\mathbf{ze}} \frac{\mu_e(\mathbf{z})p}{\mu_e(\mathbf{z}, S)} \mu_e(\mathbf{z}|S)\mu(e|S) = \sum_{\mathbf{ze}} \frac{\mu_e(\mathbf{z})p}{\mu_e(S)} \mu(e|S) \\ &= \sum_e \frac{p}{\mathbb{P}_{\mu_e}(S)} \mu(e|S) = \frac{p}{\mathbb{P}_{\mu}(S)} \leq \frac{p}{\kappa} \end{aligned} \quad (\text{A16})$$

$$\Rightarrow -\log_2 \left[\sum_{\mathbf{ze}} \left(\max_{\mathbf{c}} \omega(\mathbf{c}|\mathbf{ze}) \right) \omega(\mathbf{ze}) \right] \geq \log_2(\kappa) - \log_2(p) \quad (\text{A17})$$

Hence, we have obtained an upper bound on the average conditional maximum probability in (A16). Since by definition the ϵ/κ -smooth average conditional min-entropy involves a maximum (over the set $\mathcal{B}^{\epsilon/\kappa}(\mu)$) of the quantity on the left hand side of (A17), the final result follows. $\square$

Appendix B

Appendix B.1. Proofs Using Convex Geometry

Here we prove Theorems 4 and 5 using arguments from convex geometry.

Appendix B.1.1. Proof for Theorem 4

Theorem. Suppose Π is closed and equal to the convex hull of its extreme points. Then, there is a distribution $\omega(CZE) \in \Sigma_E$ with $|\mathcal{E}| = 1 + \dim \Pi$ such that $H_\mu(C|ZE) = h_{\min}(\rho(CZ))$.

Proof. We will be analysing $h_{\min}(\cdot)$ as a function with domain Π . It is useful to re-write $h_{\min}(\cdot)$ in the form

$$h_{\min}(\rho) = \inf_{\{\sigma_i\}_{i \in I}} \sum_i p_i H_{\sigma_i}(C|Z),$$

where the infimum is taken over all finite subsets $\{\sigma_i\}_{i \in I} \subseteq \Pi$ for which $\sum_{i \in I} p_i \sigma_i = \rho$ for some collection of non-negative p_i summing to 1. This is equivalent to the earlier definition if we set $\omega(CZ|e_i) = \sigma_i(CZ)$ and $\omega(e_i) = p_i$, yielding $H_\omega(C|ZE) = \sum_{i,j} H_{\omega(C|z_j, e_i)}(C) \omega(z_j, e_i) = \sum_i \left(\sum_j H_{\omega(C|z_j, e_i)}(C) \omega(z_j|e_i) \right) \omega(e_i) = \sum_i \left(\sum_j H_{\sigma_i(C|z_j)}(C) \sigma_i(z_j) \right) p_i = \sum_i p_i H_{\sigma_i}(C|Z)$. We first observe that the scope of the infimum can be reduced to consider only sets of σ_i belonging to Π_{extr} , the set of extreme points of Π . This follows from the fact that conditional Shannon entropy is concave. (The proof of Theorem 43 in [15] correctly notes that the concavity of conditional Shannon entropy can be obtained as a specialisation of the concavity of the quantum conditional entropy. It is worth noting, however, that the classical-only result can be obtained much more quickly and directly as shown in Appendix C.) Hence, any expression in the scope of the infimum defining $h_{\min}$ can always be decreased (or at least unchanged) by replacing each σ_i in the expression with a convex combination of extremal behaviours replicating σ_i .

Π is a subset of $\mathbb{R}^N$ where $N = |Z| \times |C|$ is the number of conditional probabilities appearing in the behaviour. In general, N is strictly larger than $\dim \Pi$: the constraint that certain elements of Π need to form valid probability distributions reduces the dimension and no-signalling equalities can reduce the dimension further. So, we seek to re-parameterise the elements of Π using only the number of coordinates necessary based on its dimension. The (affine) dimension of Π is by definition the dimension of the smallest affine space containing it—that is, the intersection of all affine subspaces of $\mathbb{R}^N$ containing Π , which is itself affine space. Let us call this smallest affine space $\mathcal{A}$. If $\dim \mathcal{A} = m$, then there is a set of m linearly independent vectors $\vec{v}_i$ and a displacement/base vector $\vec{b}$ such that any $\sigma \in \mathcal{A}$ has a unique representation as

$$\sigma = \vec{b} + \sum_{i=1}^m c_i \vec{v}_i. \quad (\text{A18})$$

For any $\sigma \in \Pi \subseteq \mathcal{A}$, then, we can uniquely represent σ as a vector of these coefficients, $(c_1, c_2, \dots, c_m)$.

Our approach now makes explicit the arguments only alluded to in the proof of Theorem 43 in [15] through general referral to existence and extension theorems in convex analysis, and takes full advantage of the fact that we are always working in a large ambient $\mathbb{R}^n$, allowing us to harness the strength of linear algebra. We would like to construct an affine-linear map $g : \mathbb{R}^N \rightarrow \mathbb{R}^m$ whose restriction to $\mathcal{A}$ maps the N -coordinate vector σ to its m -coordinate representation $(c_1, c_2, \dots, c_m)$. Our affine-linear map will be represented by a matrix M and a vector $\vec{k}$ such that $g(\sigma) = M\sigma + \vec{k} = (c_1, \dots, c_m)$. To construct M and $\vec{k}$, let A be the $N \times m$ matrix whose m columns are the vectors $\vec{v}_i$ appearing in (A18). Since the columns of A are linearly independent, $A^T A$ is invertible as its kernel consists only of the zero vector:

$$A^T A \vec{v} = \vec{0} \Rightarrow \vec{v}^T A^T A \vec{v} = 0 \Rightarrow (A \vec{v})^T A \vec{v} = 0 \Rightarrow \|A \vec{v}\| = 0 \Rightarrow A \vec{v} = \vec{0} \Rightarrow \vec{v} = \vec{0}$$

We can thus define $M = (A^T A)^{-1} A^T$ which will satisfy $MA = I$ (M is a *pseudo-inverse* of A), and so M maps the vectors $\vec{v}_i$ to the standard basis vectors in $\mathbb{R}^m$. Setting $\vec{k} = -M\vec{b}$ yields the desired $g(\cdot)$.

We point out a couple of properties of g that we will use in our arguments. First, it commutes with convex combinations: for a set of non-negative p_i satisfying $\sum_i p_i = 1$ and a collection of elements σ_i of $\mathcal{A}$,

$$\sum_i p_i g(\sigma_i) = g\left(\sum_i p_i \sigma_i\right), \quad (\text{A19})$$

which follows directly from expressing g as $M(\cdot) + \vec{k}$ and noticing that $\sum_i p_i \vec{k} = \vec{k}$. Second, M is injective when restricted to $\mathcal{A}$, so consequently g is a bijection between $\mathcal{A}$ and R^m , and in particular

$$g\left(\sum_i p_i \sigma_i\right) = g(\sigma) \Leftrightarrow \sum_i p_i \sigma_i = \sigma. \quad (\text{A20})$$

The following development is inspired by the arguments in the appendix of [34], though the assumptions and conclusions differ somewhat. Let us consider the following subset of $\mathbb{R}^{m+1}$,

$$\Xi_{\text{extr}} = \{(g(\sigma), H_\sigma(C|Z)) : \sigma \in \Pi_{\text{extr}}\},$$

where the first m coordinates of an element of Ξ_{extr} are the coordinates of $g(\sigma)$ and the $m + 1$ coordinate is $H_\sigma(C|Z)$. Define

$$\Xi = \text{conv}(\Xi_{\text{extr}}),$$

where ‘conv’ denotes the convex hull. Ξ_{extr} and Ξ are artificial constructions but by studying their geometry we can prove the existence of a convex combination achieving the infimum defining $h_{\min}(\rho)$.

We first confirm that Ξ_{extr} is indeed the set of extremal points Ξ (as suggested by our choice in names), i.e., we confirm that Ξ_{extr} contains only trivial convex combinations of its elements. To see this, note if $\sum_i p_i (g(\sigma_i), H_{\sigma_i}(C, Z)) = (g(\sigma), H_\sigma(C, Z))$ holds for some $\sigma_i, \sigma \in \Pi_{\text{extr}}$ and non-negative p_i satisfying $\sum_i p_i = 1$, then we must have $\sum_i p_i g(\sigma_i) = g(\sigma)$ and so $\sum_i p_i \sigma_i = \sigma$ by (A19) and (A20). This can only be a trivial convex combination (i.e., all σ_i with nonzero p_i coefficient must equal σ) as the σ_i and σ are assumed to be in Π_{extr} .

Second, we show that the point $(g(\rho), h_{\min}(\rho))$ is on the boundary of Ξ , i.e., that $(g(\rho), h_{\min}(\rho))$ is a limit point of Ξ and also a limit point of Ξ^C . To see that we can converge to this point from *within* the set, note that, for any set of $\sigma_i \in \Pi_{\text{extr}}$ satisfying $\sum_i p_i \sigma_i = \rho$, we have by definition $\sum_i p_i (g(\sigma_i), H_{\sigma_i}(C|Z)) \in \Xi$ which can be re-expressed as $(g(\rho), \sum_i p_i H_{\sigma_i}(C|Z)) \in \Xi$ by invoking (A19). By the nature of the infimum defining $h_{\min}(\rho)$, there must be a sequence of such elements of Ξ whose last component forms a non-increasing sequence converging to $h_{\min}(\rho)$; since the first m components are identically $g(\rho)$, this sequence converges to $(g(\rho), h_{\min}(\rho))$ as desired. Similarly, one can also converge to $(g(\rho), h_{\min}(\rho))$ from *outside* the set Ξ : $(g(\rho), h_{\min}(\rho) - \epsilon) \notin \Xi$ for all $\epsilon > 0$; this is because all elements of Ξ take the form

$$\sum_i p_i (g(\sigma_i), H_{\sigma_i}(C|Z)) = \left(\sum_i p_i g(\sigma_i), \sum_i p_i H_{\sigma_i}(C|Z) \right),$$

for some collection $\sigma_i \in \Pi_{\text{extr}}$ and if the first m coordinates are equal to $g(\rho)$, then by (A19) and (A20) we must have $\sum_i p_i \sigma_i = \rho$ and so the $m + 1$ coordinate is a term contributing to the infimum defining $h_{\min}(\rho)$; it cannot be less than $h_{\min}(\rho)$.

We now would like to demonstrate that $(g(\rho), h_{\min}(\rho))$ is contained in Ξ . As a first step, we show that

$$(g(\rho), h_{\min}(\rho)) \in \text{conv}(\overline{\Xi_{\text{extr}}}), \quad (\text{A21})$$

where $\text{conv}(\overline{\Xi_{\text{extr}}})$ denotes the convex hull of the closure of Ξ_{extr} . To see this, first note that Ξ_{extr} is bounded—for the $m + 1$ coordinate, Shannon entropy is non-negative with a maximum value set by the cardinality of the value space of C and, for the first m coordinates, these are contained in the image of the set Π_{extr} through the continuous map g —and since Π_{extr} is contained in the compact set $\mathcal{P} = [0, 1]^n$ ($\mathcal{P}$ contains all probability distributions), its

image must be contained in the compact (and thus bounded) set $g(\mathcal{P})$. As Ξ_{extr} is bounded, its closure, denoted $\overline{\Xi_{\text{extr}}}$, must be bounded as well and so is compact. It is a known fact that the convex hull of a compact set in $\mathbb{R}^n$ is compact, so $\text{conv}(\overline{\Xi_{\text{extr}}})$ is compact—and so, in particular, closed. Finally $\text{conv}(\overline{\Xi_{\text{extr}}})$ clearly contains $\Xi = \text{conv}(\Xi_{\text{extr}})$, the convex hull of a smaller set; as a closed set containing Ξ , it will contain the Ξ -boundary point $(g(\rho), h_{\min}(\rho))$.

Now we show that this implies containment in Ξ proper. Since the map

$$h(\rho) := (g(\rho), H_{\rho}(C|Z))$$

with image in $\mathbb{R}^{m+1}$ is continuous on the domain of n -dimensional probability distributions and Π_{extr} is bounded, we have $\overline{h(\Pi_{\text{extr}})} \subseteq h(\overline{\Pi_{\text{extr}}})$. (For any bounded subset S in $\mathbb{R}^n$ like Π_{extr} and continuous h , we have $\overline{h(S)} \subseteq h(\overline{S})$, the proof of which is as follows: Any $x \in \overline{h(S)}$ must be the limit of a sequence in $h(S)$; let $\{s_i\}_{i=1}^{\infty} \subseteq S$ satisfy $h(s_i) \rightarrow x$. Since $\overline{S}$ is compact, $\{s_i\}_{i=1}^{\infty} \subseteq S \subseteq \overline{S}$ has a convergent sub-sequence $\{s_j\}_{j=1}^{\infty}$ with limit in $\overline{S}$; let $s \in \overline{S}$ be this limit. By continuity, $h(s_j) \rightarrow h(s)$; considered as a sub-sequence of $\{h(s_i)\}_{i=1}^{\infty}$, we also have $h(s_j) \rightarrow x$ and so uniqueness of limits implies $x = h(s) \in h(\overline{S})$.)

Now, since by definition $\Xi_{\text{extr}} = h(\Pi_{\text{extr}})$, we write

$$\overline{\Xi_{\text{extr}}} \subseteq h(\overline{\Pi_{\text{extr}}}). \quad (\text{A22})$$

Next, using (A21), (A22), the definition of $h(\cdot)$ and finally (A19), we can write

$$\begin{aligned} (g(\rho), h_{\min}(\rho)) &= \sum_i p_i \vec{w}_i \text{ for some } \{\vec{w}_i\}_{i \in I} \subseteq \overline{\Xi_{\text{extr}}} \\ &= \sum_i p_i h(\tau_i) \text{ for some } \{\tau_i\}_{i \in I} \subseteq \overline{\Pi_{\text{extr}}} \\ &= \sum_i p_i (g(\tau_i), H_{\tau_i}(C|Z)) \text{ for some } \{\tau_i\}_{i \in I} \subseteq \overline{\Pi_{\text{extr}}} \\ &= \left(g\left(\sum_i p_i \tau_i\right), \sum_i p_i H_{\tau_i}(C|Z) \right) \text{ for some } \{\tau_i\}_{i \in I} \subseteq \overline{\Pi_{\text{extr}}}. \end{aligned} \quad (\text{A23})$$

Comparing the first expression in the above sequence to the last and applying (A20) implies that $\sum_i p_i \tau_i = \rho$. Now, since by assumption Π is closed, $\Pi_{\text{extr}} \subseteq \Pi$ implies $\overline{\Pi_{\text{extr}}} \subseteq \Pi$, so $\Pi = \text{conv}(\Pi_{\text{extr}})$ implies that elements of $\overline{\Pi_{\text{extr}}}$ can be expressed as convex combinations of elements of Π_{extr} . Thus, in the expression $\sum_i p_i \tau_i$, if there are any non-extremal τ_i elements they can be replaced with convex combinations of elements of Π_{extr} to yield a convex combination $\sum_j q_j \sigma_j$ equalling ρ where the concavity of conditional Shannon entropy implies that $\sum_j q_j H_{\sigma_j}(C|Z)$ is not larger than $\sum_i p_i H_{\tau_i}(C|Z)$. However, by (A23), $\sum_i p_i H_{\tau_i}(C|Z) = h_{\min}(\rho)$ and, since $\sum_j q_j H_{\sigma_j}(C|Z)$ cannot be smaller than $h_{\min}(\rho)$, it must equal $h_{\min}(\rho)$. As $\sum_j q_j \sigma_j = \rho$ and $\sum_j q_j H_{\sigma_j}(C|Z) = h_{\min}(\rho)$, one more application of (A19) yields

$$\begin{aligned} (g(\rho), h_{\min}(\rho)) &= \left(g\left(\sum_j q_j \sigma_j\right), \sum_j q_j H_{\sigma_j}(C|Z) \right) \text{ for some } \{\sigma_j\}_{j \in J} \subseteq \Pi_{\text{extr}} \\ &= \sum_j q_j (g(\sigma_j), H_{\sigma_j}(C|Z)) \text{ for some } \{\sigma_j\}_{j \in J} \subseteq \Pi_{\text{extr}}, \end{aligned}$$

which is in Ξ .

The argument thus far demonstrates the existence of a convex combination of Π_{extr} elements explicitly achieving the infimum in the definition of $h_{\min}(\rho)$. We continue with our argument to further demonstrate that the number of required Π_{extr} elements in such an optimal decomposition is not greater than $m + 1$.

We first note that, since $(g(\rho), h_{\min}(\rho))$ is on the boundary of the convex set Ξ , the supporting hyperplane theorem says there is a supporting hyperplane H_ρ with $(g(\rho), h_{\min}(\rho)) \in H_\rho$ and Ξ entirely on one side of H_ρ . Now, notice that if we decompose $(g(\rho), h_{\min}(\rho))$ as a convex combination of Ξ_{extr} elements, these elements must all lie in the hyperplane H_ρ : this is because any elements strictly on one side of H_ρ would have to be counterbalanced by elements strictly on the other side of H_ρ —but, since one side of H_ρ is disjoint from Ξ , this is not possible. Applying the same observation to any other element of $H_\rho \cap \Xi$, it follows that $H_\rho \cap \Xi$ is contained in the convex hull of $H_\rho \cap \Xi_{\text{extr}}$. As the reverse inclusion follows from the convexity of H_ρ and the fact that $\Xi = \text{conv}(\Xi_{\text{extr}})$, we can write $\text{conv}(H_\rho \cap \Xi_{\text{extr}}) = H_\rho \cap \Xi$. Now, since $H_\rho \cap \Xi$ is at most m dimensional (H_ρ , as a hyperplane, has one fewer dimension than the ambient $(m+1)$ -dimensional space), we can invoke Carathéodory's theorem to see that at most $m+1$ points of $H_\rho \cap \Xi_{\text{extr}}$ are required to replicate $(g(\rho), h_{\min}(\rho))$ as a convex combination. Thus, we have

$$(g(\rho), h_{\min}(\rho)) = \sum_i p_i \bar{w}_i \text{ for some } \{\bar{w}_i\}_{i \in I} \subseteq \Xi_{\text{extr}}, |I| \leq m+1$$

and so, recalling the definition of Ξ_{extr} and invoking (A19) one last time, we can write that, for some integer m^* satisfying $1 \leq m^* \leq m+1$,

$$\begin{aligned} (g(\rho), h_{\min}(\rho)) &= \sum_{i=1}^{m^*} p_i (g(\sigma_i), H_{\sigma_i}(C|Z)) \text{ for some } \{\sigma_i\}_{i=1}^{m^*} \subseteq \Pi_{\text{extr}} \\ &= \left(g\left(\sum_{i=1}^{m^*} p_i \sigma_i\right), \sum_{i=1}^{m^*} p_i H_{\sigma_i}(C|Z) \right) \text{ for some } \{\sigma_i\}_{i=1}^{m^*} \subseteq \Pi_{\text{extr}}. \end{aligned}$$

By (A20), $\sum_{i=1}^{m^*} p_i \sigma_i = \rho$ and so $\{\sigma_i\}_{i=1}^{m^*}$ induces the desired distribution $\omega(\text{CZE})$ by setting $\omega(\text{CZ}|e_i) = \sigma_i(\text{CZ})$ and $\omega(e_i) = p_i$. $\square$

Appendix B.1.2. Proof for Theorem 5

Theorem. Suppose Π satisfies the conditions of Theorem 4 and ρ is in the interior of Π . Then, there exists an entropy estimator whose entropy estimate at ρ is equal to $h_{\min}(\rho)$.

Proof. We continue from where we left off in the proof of Theorem 4 and show that the supporting hyperplane H_ρ discussed in that proof can be used to construct an affine function that is the desired entropy estimator. Recall that the dimension of Π , which is embedded in a higher dimensional vector space $\mathbb{R}^N$, is defined as the affine dimension of $\mathcal{A}$, the smallest affine subspace containing Π . Given this context, the assumption that ρ is in the interior of Π means that there exists an open ϵ -ball U in $\mathbb{R}^N$ such that $U \cap \mathcal{A}$, which is open in the subspace topology, is contained in Π . (If this assumption is removed, a weaker form of the theorem demonstrating the existence of entropy estimators with estimate ϵ -close to $h_{\min}(\rho)$ can be proved with a similar argument to that of the current proof by invoking Exercise 3.28 of [35].)

First, we note that $g(\rho)$ is in the interior of $g(\Pi)$. To see this, consider the restriction $g|_{\mathcal{A}}$ of g to $\mathcal{A}$, which is a bijection with affine-linear inverse map $(g|_{\mathcal{A}})^{-1} : \mathbb{R}^m \rightarrow \mathcal{A}$ given by $A(\cdot) + \vec{b}$ (recalling the construction following (A18) in the proof of Theorem 4). This ensures that the set $g|_{\mathcal{A}}(U \cap \mathcal{A})$ must be open, as it is equal to the inverse image of $U \cap \mathcal{A}$ under the map $(g|_{\mathcal{A}})^{-1}$ which is equal to the inverse image of the open set U under the (continuous) map $A(\cdot) + \vec{b} : \mathbb{R}^m \rightarrow \mathbb{R}^N$. Hence, $g(\rho)$ is contained in the open set $g(U \cap \mathcal{A})$ which is a subset of $g(\Pi)$ as $U \cap \mathcal{A} \subseteq \Pi$.

Now, we take a closer look at H_ρ , the supporting hyperplane touching Ξ at $(g(\rho), h_{\min}(\rho))$. As a hyperplane, H_ρ will be equal to the set of $\vec{x}$ satisfying an equation of the form $\vec{a} \cdot \vec{x} = b$ for some fixed $\vec{a} \in \mathbb{R}^{m+1}$ and $b \in \mathbb{R}$, where $\cdot$ denotes the dot product, and the condition

$$\zeta \in \Xi \Rightarrow \vec{a} \cdot \zeta \geq b \quad (\text{A24})$$

expresses algebraically the notion that Ξ is on one side of H_ρ . We argue that the fact that $g(\rho)$ is in the interior of $g(\Pi)$ implies the $m+1$ component of $\vec{a}$, denoted $\vec{a}_{m+1}$, must be nonzero. Assume $\vec{a}_{m+1} = 0$ for a proof by contradiction: since $(g(\rho), h_{\min}(\rho))$ is the point of contact of the supporting hyperplane H_ρ , we have $\vec{a} \cdot (g(\rho), h_{\min}(\rho)) = b$, which implies $\vec{a}_{[m]} \cdot g(\rho) = b$ where $\vec{a}_{[m]} \in \mathbb{R}^m$ denotes the vector consisting of the first m coordinates of $\vec{a}$. Since the previous paragraph demonstrated there is an open subset of $g(\Pi)$ containing $g(\rho)$, this means $g(\rho) - c\vec{a}_{[m]}$ for a sufficiently small positive c is equal to $g(\phi)$ for some ϕ in Π . By construction, ϕ will satisfy $\vec{a}_{[m]} \cdot g(\phi) < b$ but since $\vec{a}_{m+1} = 0$ this requires $\vec{a} \cdot (g(\phi), h_{\min}(\phi)) < b$ as well. This would imply $(g(\phi), h_{\min}(\phi)) \notin \Xi$; however, this is a contradiction as the arguments of Theorem 4 show that, for any $\phi \in \Pi$, $(g(\phi), h_{\min}(\phi))$ belongs to Ξ (the arguments of Theorem 4 demonstrated this for ρ but they apply to any element of Π).

Having demonstrated $\vec{a}_{m+1} \neq 0$, we can define a function $f_\rho : \mathbb{R}^m \rightarrow \mathbb{R}$ as follows:

$$f_\rho(\vec{x}) = \frac{b - \vec{a}_{[m]} \cdot \vec{x}}{\vec{a}_{m+1}}. \quad (\text{A25})$$

Composing this function with g , we find that $f_\rho \circ g(\rho) = h_{\min}(\rho)$. Furthermore, for general $\phi \in \Pi$ the fact that $(g(\phi), h_{\min}(\phi)) \in \Xi$ ensures $f_\rho \circ g(\phi) \leq h_{\min}(\phi)$, and $h_{\min}(\phi) \leq H_\phi(C|Z)$ by concavity of conditional Shannon entropy, so the map $f_\rho \circ g$, applied to any $\phi \in \Pi$, satisfies

$$f_\rho \circ g(\phi) \leq H_\phi(C|Z) = E_\phi(-\log_2[\phi(C|Z)]).$$

We now use $f_\rho \circ g$ to construct the desired entropy estimator as follows. We have

$$f_\rho \circ g(\phi) = \frac{b - \vec{a}_{[m]} \cdot (M\phi + \vec{k})}{\vec{a}_{m+1}} = \vec{n} \cdot \phi + d$$

where $d = (b - \vec{a}_{[m]} \cdot \vec{k})/\vec{a}_{m+1}$ is a constant and $\vec{n} = -(1/\vec{a}_{m+1})M^T\vec{a}_{[m]}$ is an N -dimensional vector; that is, it has one component for each possible distinct outcome pair c, z for the random variable pair C, Z . Now we can define $K(c, z) := \vec{n}_{cz} + d$ to obtain a function of C, Z satisfying

$$\mathbb{E}_\phi[K(CZ)] = \vec{n} \cdot \phi + d = f_\rho \circ g(\phi) \leq E_\phi(-\log_2[\phi(C|Z)]),$$

and thus K is an entropy estimator satisfying the conditions of the theorem. $\square$

Appendix C

Concavity of Conditional Shannon Entropy

It is known that conditional Shannon entropy is concave. For completeness, we provide a brief proof of how this follows from the concavity of (unconditional) Shannon entropy. Let ν be a convex combination of ν_1 and ν_2 , so that for all $(c, z) \in \mathcal{C} \times \mathcal{Z}$ we have $\nu(c, z) = \lambda\nu_1(c, z) + (1 - \lambda)\nu_2(c, z)$ for some $\lambda \in [0, 1]$. Then, it follows that $\nu(C|z)$ is a convex mixture of $\nu_1(C|z)$ and $\nu_2(C|z)$ for each fixed z for which $\nu(z) > 0$:

$$\nu(C|z) = \lambda \frac{\nu_1(z)}{\nu(z)} \nu_1(C|z) + (1 - \lambda) \frac{\nu_2(z)}{\nu(z)} \nu_2(C|z), \quad (\text{A26})$$

where it is straightforward to check that the coefficients of $\nu_1(C|z)$ and $\nu_2(C|z)$ are non-negative numbers summing to one. Then, using (A26) and the concavity of (unconditional) Shannon entropy we have the following.

$$\begin{aligned} H_\nu(C|Z) &= \sum_z H_\nu(C|z)\nu(z) \geq \sum_z \left(\lambda \frac{\nu_1(z)}{\nu(z)} H_{\nu_1}(C|z) + (1 - \lambda) \frac{\nu_2(z)}{\nu(z)} H_{\nu_2}(C|z) \right) \nu(z) \\ &= \sum_z (\lambda H_{\nu_1}(C|z)\nu_1(z) + (1 - \lambda) H_{\nu_2}(C|z)\nu_2(z)) = \lambda H_{\nu_1}(C|Z) + (1 - \lambda) H_{\nu_2}(C|Z) \end{aligned} \quad (\text{A27})$$

Appendix D

Useful Lemmas

The following lemmas are used for arguments in Appendices A.1.1 and E.

Lemma A1. *If the distributions $\mu(X)$ and $\mu'(X)$ dominate the non-negative function $f: \mathcal{X} \rightarrow \mathbb{R}^+$ with weight $w(f) = \sum_{x \in \mathcal{X}} f(x) = 1 - \epsilon$ for $\epsilon \in [0, 1]$, i.e., $\mu(x) \geq f(x)$, $\mu'(x) \geq f(x)$, for all $x \in \mathcal{X}$, then $d_{TV}(\mu, \mu') \leq \epsilon$.*

Proof. Using the definition of TV distance we have the required result as shown below.

$$\begin{aligned} d_{TV}(\mu, \mu') &= \frac{1}{2} \sum_{x \in \mathcal{X}} |\mu(x) - f(x) + f(x) - \mu'(x)| \\ &\leq \frac{1}{2} \sum_{x \in \mathcal{X}} (|\mu(x) - f(x)| + |\mu'(x) - f(x)|) = 1 - w(f) = \epsilon. \end{aligned}$$

□

Lemma A2. *Suppose the function $f: \mathcal{X} \rightarrow \mathbb{R}^+$ has weight $w(f) = \sum_{x \in \mathcal{X}} f(x) = 1 - \epsilon$ where $\epsilon \in [0, 1]$ and satisfies $f(x) \leq p$, $\forall x \in \mathcal{X}$ for some fixed $p \geq 1/|\mathcal{X}|$. Then, there exists a distribution $\mu'(X)$ such that $f(x) \leq \mu'(x) \leq p$ holds for all $x \in \mathcal{X}$.*

Proof. If $\epsilon = 0$, it suffices to take $\mu'(x) = f(x)$. If $\epsilon > 0$, we construct a distribution satisfying the two properties as follows. Define a function μ_λ with domain $\mathcal{X}$ as

$$\mu_\lambda(x) = (1 - \lambda)f(x) + \lambda p \quad (\text{A28})$$

Then, for any fixed $\lambda \in [0, 1]$, $\mu_\lambda(x)$ is a convex combination of non-negative numbers and thus non-negative for any choice of x . We show that there exists a $\lambda \in [0, 1]$ for which $\sum_x \mu_\lambda(x) = 1$, making μ_λ a distribution. It is easy to verify that for $\lambda' = \epsilon/(p|\mathcal{X}| + \epsilon - 1)$ the above function adds up to unity when summed over $x \in \mathcal{X}$. We just need to ensure that $\epsilon/(p|\mathcal{X}| + \epsilon - 1) \in [0, 1]$ holds. To see this, note that $p|\mathcal{X}| \geq 1$ so we have $p|\mathcal{X}| + \epsilon - 1 \geq \epsilon$ and since $\epsilon > 0$ the quotient must indeed lie in $[0, 1]$. Finally, $\mu_{\lambda'}(X)$ satisfies the bounds in the lemma: since $f(x) \leq p$ for all $x \in \mathcal{X}$, for any $\lambda \in [0, 1]$ we have

$$p \geq p + (1 - \lambda)(f(x) - p) = (1 - \lambda)f(x) + \lambda p = f(x) + \lambda(p - f(x)) \geq f(x), \forall x \in \mathcal{X} \quad (\text{A29})$$

and the middle term above is $\mu_{\lambda'}(X)$ for $\lambda = \lambda'$. □

Appendix E

Inequalities Relating Smooth Average Conditional Min-Entropy and Smooth Worst-Case Conditional Min-Entropy

Here we state and prove a known inequality that relates two notions of smooth conditional min-entropy. We present this result without structuring random variables as stochastic sequences, i.e., instead of considering distributions of $\mathbf{C}, \mathbf{Z}, E$ we consider distributions of X, Y . The result and its proof can be adapted to the more general case involving sequence of random variables.

For a distribution $\mu: \mathcal{X} \times \mathcal{Y} \rightarrow [0, 1]$ of X, Y and the set $\mathcal{B}^\epsilon(\mu)$ of distributions of X, Y defined as $\mathcal{B}^\epsilon(\mu) := \{\sigma: \mathcal{X} \times \mathcal{Y} \rightarrow [0, 1] \mid d_{TV}(\mu, \sigma) \leq \epsilon\}$, the ϵ -smooth average conditional min-entropy is:

$$\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y) := \max_{\sigma \in \mathcal{B}^\epsilon(\mu)} \left[-\log_2 \left(\sum_{y \in \mathcal{Y}} \max_{x \in \mathcal{X}} \sigma(x|y) \sigma(y) \right) \right].$$

A stricter definition of smooth conditional min-entropy than the one stated above is the ϵ -smooth “worst-case” conditional min-entropy, introduced in [19]. It reads as follows:

$$\mathbb{H}_{\infty, \mu}^{\text{wst}, \epsilon}(X|Y) = \max_{\sigma \in \mathcal{B}^{\epsilon}(\mu)} \left[-\log_2 \left(\max_{x \in \mathcal{X}, y \in \mathcal{Y}} \sigma(x|y) \right) \right]. \quad (\text{A30})$$

For purposes of randomness extraction or scenarios involving predictability of an adversary, the smooth average conditional min-entropy suffices. One can show that the notions of *average-case* and *worst-case* are equivalent up to an additive factor [36]. This is formalised in Proposition A1.

Proposition A1. For a distribution $\mu: \mathcal{X} \times \mathcal{Y} \rightarrow [0, 1]$ of X, Y and $1 \geq \epsilon \geq 0, 1 > \epsilon' > 0$, the smooth worst case conditional min-entropy and smooth average case conditional min-entropy are related by the following inequalities:

$$\mathbb{H}_{\infty, \mu}^{\text{wst}, \epsilon}(X|Y) \leq \mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y) \leq \mathbb{H}_{\infty, \mu}^{\text{wst}, \epsilon + \epsilon'}(X|Y) + \log_2(1/\epsilon') \quad (\text{A31})$$

Proof. The first inequality holds immediately, since for every $\sigma(X, Y) \in \mathcal{B}^{\epsilon}(\mu)$ we have

$$\max_{x \in \mathcal{X}, y \in \mathcal{Y}} \sigma(x|y) \geq \sum_{y \in \mathcal{Y}} \left(\max_{x \in \mathcal{X}} \sigma(x|y) \right) \sigma(y). \quad (\text{A32})$$

Taking $-\log_2$ of both sides we obtain $\mathbb{H}_{\infty, \sigma}^{\text{wst}}(X|Y) \leq \mathbb{H}_{\infty, \sigma}^{\text{avg}}(X|Y)$, where $\mathbb{H}_{\infty, \sigma}^{\text{wst}}(X|Y)$ is the bracketed quantity in (A30) and, since this inequality holds for every $\sigma \in \mathcal{B}^{\epsilon}(\mu)$, we have $\mathbb{H}_{\infty, \mu}^{\text{wst}, \epsilon}(X|Y) \leq \mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y)$. For the second inequality of (A31), we want to show that $\mathbb{H}_{\infty, \mu}^{\text{wst}, \epsilon + \epsilon'}(X|Y) \geq \mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y) - \log(1/\epsilon')$ holds. Suppose the distribution $\nu(X, Y) \in \mathcal{B}^{\epsilon}(\mu)$ witnesses $\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y)$, i.e., $\mathbb{H}_{\infty, \mu}^{\text{avg}, \epsilon}(X|Y) = \mathbb{H}_{\infty, \nu}^{\text{avg}}(X|Y)$. The existence of such a witness follows from the compactness of $\mathcal{B}^{\epsilon}(\mu)$ and the continuity of $\mathbb{H}_{\infty, \mu}^{\text{avg}}(X|Y)$. It suffices to construct a distribution $\sigma(X, Y) \in \mathcal{B}^{\epsilon'}(\nu)$ such that $\max_{x, y} \sigma(x|y) \leq p/\epsilon'$ holds, where $p = \sum_y \max_x \nu(x|y) \nu(y) = \mathbb{E}_{\nu(Y)}[\max_{x \in \mathcal{X}} \nu(x|Y)]$. We begin by defining the sub-probability distribution $\tilde{\nu}(X, Y)$ as shown below.

$$\tilde{\nu}(x, y) = \nu(x, y) \mathbb{I}[\max_{x \in \mathcal{X}} \nu(x|y) \leq p/\epsilon'], \quad (\text{A33})$$

where the notation $\mathbb{I}[\cdot]$ represents the function that evaluates to 1 if the enclosed condition holds and zero if it does not. Basically, the definition of $\tilde{\nu}$ in (A33) involves discarding ν corresponding to those $y \in \mathcal{Y}$ for which $\max_x \nu(x|y) > p/\epsilon'$ holds. An application of Markov’s inequality then shows that the weight $w(\tilde{\nu})$ of $\tilde{\nu}$ is at least $1 - \epsilon'$:

$$\begin{aligned} w(\tilde{\nu}) &= \sum_{x, y} \tilde{\nu}(x, y) = \sum_{y \in \mathcal{Y}: \max_x \nu(x|y) \leq p/\epsilon'} \sum_x \nu(x, y) \\ &= \mathbb{P}_{\nu(Y)} \left(\max_{x \in \mathcal{X}} \nu(x|Y) \leq p/\epsilon' \right) \geq 1 - \epsilon' \end{aligned} \quad (\text{A34})$$

Since $p = \mathbb{E}_{\nu(Y)}[\max_x \nu(x|Y)]$, (A34) follows. One way to now construct a distribution $\sigma \in \mathcal{B}^{\epsilon'}(\mu)$ satisfying $\max_{x, y} \sigma(x|y) \leq p/\epsilon'$ is to scale $\tilde{\nu}$, i.e., we define $\sigma(X, Y)$ as $\sigma(x, y) = \tilde{\nu}(x, y)/w(\tilde{\nu})$. Note that, since $w(\tilde{\nu}) \geq 1 - \epsilon'$ and $\epsilon' \in (0, 1)$, $w(\tilde{\nu})$ is positive; also, $\sigma \geq \tilde{\nu}$ holds since $w(\tilde{\nu}) \leq 1$. Together with the fact that $\nu \geq \tilde{\nu}$, we can use Lemma A1 to show that $d_{\text{TV}}(\sigma, \nu) \leq 1 - w(\tilde{\nu}) \leq \epsilon'$. By definition of σ we have $\sigma(x, y) \leq p\sigma(y)/\epsilon'$ for all choices of x, y , where $\sigma(y) = \frac{\nu(y)}{w(\tilde{\nu})} \mathbb{I}[\max_x \nu(x|y) \leq p/\epsilon']$ for each y . With the convention that $\sigma(x|y)$ is assigned the value 0 when $\sigma(y) = 0$, we then have $\sigma(x|y) \leq p/\epsilon'$ for all choices of x, y . Membership of σ in the set $\mathcal{B}^{\epsilon + \epsilon'}(\mu)$ follows from the triangle inequality $d_{\text{TV}}(\mu, \sigma) \leq d_{\text{TV}}(\mu, \nu) + d_{\text{TV}}(\nu, \sigma) \leq \epsilon + \epsilon'$. And so we have constructed a distribution in $\mathcal{B}^{\epsilon + \epsilon'}(\mu)$ such that $\max_{x, y} \sigma(x|y) \leq p/\epsilon'$. Taking $-\log_2$ on both sides, we

obtain $\mathbb{H}_{\infty,\sigma}^{\text{wst}}(X|Y) \geq \mathbb{H}_{\infty,\nu}^{\text{avg}}(X|Y) - \log_2(1/\epsilon')$. As mentioned earlier, $\nu \in \mathcal{B}^\epsilon(\mu)$ witnesses $\mathbb{H}_{\infty,\mu}^{\text{avg},\epsilon}(X|Y)$; hence, we have shown:

$$\mathbb{H}_{\infty,\sigma}^{\text{wst}}(X|Y) \geq \mathbb{H}_{\infty,\mu}^{\text{avg},\epsilon}(X|Y) - \log_2(1/\epsilon') \quad (\text{A35})$$

Since, by definition, the smooth *worst-case* conditional min-entropy involves a maximum of the left hand side of (A35) over the set $\mathcal{B}^{\epsilon+\epsilon'}(\mu)$, this shows that $\mathbb{H}_{\infty,\mu}^{\text{wst},\epsilon+\epsilon'}(X|Y) \geq \mathbb{H}_{\infty,\mu}^{\text{avg},\epsilon}(X|Y) - \log_2(1/\epsilon')$ holds, from which the second inequality in (A31) follows. $\square$

In the asymptotic limit of a large number n of trials, constant factors vanish in measuring per-trial min entropy and, since ϵ' can be made arbitrarily small, (A31) enables us to consider either definition when considering asymptotic performance.

Appendix F

Proof of Proposition 1

Proposition. For the set of behaviours Π_{NS} , the PEF optimisation in (10) is independent of the power β for $\beta \geq \log_2(4/3)$.

Proof. For a fixed value of n and ϵ the optimisation problem in (10) is equivalent to the following:

$$\begin{aligned} &\text{Maximise: } \mathbb{E}_\rho[\log_2(F(ABXY))] \\ &\text{Subject to: } \mathbb{E}_{\mu_{\text{PR}}^i}[F(ABXY)\mu_{\text{PR}}^i(AB|XY)^\beta] \leq 1, \text{ for all } i \in \{0,1\}^3, \\ &\quad \mathbb{E}_{\mu_{\text{LD}}^j}[F(ABXY)\mu_{\text{LD}}^j(AB|XY)^\beta] \leq 1, \text{ for all } j \in \{0,1\}^4, \\ &\quad F(abxy) \geq 0, \forall a,b,x,y \in \{0,1\}, \end{aligned} \quad (\text{A36})$$

where the constraints range over the extremal points of Π_{NS} as given in (35) and (36). We show that, for $\beta \geq \log_2(4/3)$, the above constraints are equivalent to

$$\begin{aligned} &\mathbb{E}_{\mu_{\text{LD}}^j}[F(ABXY)\mu_{\text{LD}}^j(AB|XY)] \leq 1, \text{ for all } j \in \{0,1\}^4, \\ &\quad F(abxy) \geq 0, \forall a,b,x,y \in \{0,1\}, \end{aligned} \quad (\text{A37})$$

noticing that β does not appear in (A37).

It is immediate to see that the constraints of (A36) imply (A37): since $\mu(AB|XY)$ is always zero or one for local deterministic distributions, in this case we have $\mu(AB|XY)^\beta = \mu(AB|XY)$ and thus for each choice of j we have $\mathbb{E}_{\mu_{\text{LD}}^j}[F(ABXY)\mu_{\text{LD}}^j(AB|XY)^\beta] \leq 1$ implying the non- β counterpart $\mathbb{E}_{\mu_{\text{LD}}^j}[F(ABXY)\mu_{\text{LD}}^j(AB|XY)] \leq 1$ in (A37). Now, we demonstrate the reverse implication. First, the argument just given also works in the opposite direction to show that the non- β constraints of (A37) imply the corresponding constraints (with β) in (A36). We thus need only to show that the $\mathbb{E}_{\mu_{\text{PR}}^i}[\dots] \leq 1$ in (A36) are implied as well. We give a specific argument for the PR box given in Table 1; symmetric arguments apply for the other PR boxes. Since any distribution $\mu(ABXY)$ is the behaviour $\mu(AB|XY)$ times a fixed settings distribution $\sigma_s(XY)$, we can express the product $F(abxy)\sigma_s(xy)$ as $F'(abxy)$ for all choices of (a,b,x,y) when the expectation functional $\mathbb{E}[\cdot]$ is written out in full. The constraints (A37) then imply, by summing over the eight of them corresponding to the eight local deterministic distributions appearing in Table 1 (a set we denote LD_1), that

$$\sum_{a,b,x,y} F'(abxy) \sum_{\mu_{\text{LD}} \in \text{LD}_1} \mu_{\text{LD}}(ab|xy)^2 \leq 8. \quad (\text{A38})$$

Noticing that the inner sum above is always 3 or 1 (this corresponds to the number of 1s appearing in each column of Table A1, with the result given in Table A2), we can now rewrite (A38) as $3M + N \leq 8$, where $M = F'(0000) + F'(0001) + F'(0010) + F'(0111) + F'(1011) + F'(1100) + F'(1101) + F'(1110)$ and $N = F'(0011) + F'(1111) + F'(1001) + F'(0101) + F'(0110) + F'(1010) + F'(0100) + F'(1000)$.

Table A1. Tabular representation for $\mu_{\text{PR},1}(AB|XY)^{1+\beta}$.

		<i>ab</i>			
		00	01	10	11
<i>xy</i>	00	$\frac{1}{2^{1+\beta}}$	0	0	$\frac{1}{2^{1+\beta}}$
	01	$\frac{1}{2^{1+\beta}}$	0	0	$\frac{1}{2^{1+\beta}}$
	10	$\frac{1}{2^{1+\beta}}$	0	0	$\frac{1}{2^{1+\beta}}$
	11	0	$\frac{1}{2^{1+\beta}}$	$\frac{1}{2^{1+\beta}}$	0

Table A2. Tabular representation of the values of the sum $\sum_{\mu_{\text{LD}} \in \text{LD}_1} \mu_{\text{LD}}(ab|xy)^2$.

		<i>ab</i>			
		00	01	10	11
<i>xy</i>	00	3	1	1	3
	01	3	1	1	3
	10	3	1	1	3
	11	1	3	3	1

Since M, N are both non-negative, we can drop N to find that $3M + N \leq 8$ implies $M \leq 8/3 = 2^{1+\log_2(4/3)}$ which in turn implies $M \leq 2^{1+\beta}$ whenever $\beta \geq \log_2(4/3)$. Since $\mathbb{E}_{\mu_{\text{PR},1}}[F(ABXY)\mu_{\text{PR},1}(AB|XY)^\beta]$ is equal to $M(1/2)^{1+\beta}$ (see Table A1) the constraint $\mathbb{E}_{\mu_{\text{PR},1}}[\dots] \leq 1$ follows. $\square$

We remark that this inequality condition $\beta \geq \log_2(4/3)$ is tight in the following sense: there exists a non-negative function $F(abxy)$ violating the PR box constraint of $\mu_{\text{PR},1}$ appearing in (A36) for any $\beta < \log_2(4/3)$, while satisfying (A37) for any positive β —and consequently satisfying all the constraints of (A36) for $\beta \geq \log_2(4/3)$ per the argument in the above proof. Thus, the feasible set of (A36) always excludes this particular choice of F for $\beta < \log_2(4/3)$ and includes it for $\beta \geq \log_2(4/3)$. This function is $F(abxy) = (1/3)[a \oplus b = xy]\sigma_s(xy)^{-1}$; fixing $\beta = \log_2(4/3) - \epsilon$ for some choice of ϵ in the interval $(0, \log_2(4/3))$, we can check that all the LD boxes satisfy the inequality $\sum_{a,b,x,y} F'(abxy)\mu_{\text{LD}}(ab|xy)^{1+\beta} \leq 1$; the value of the expression is always either $1/3$ or 1 . However, for the PR box $\mu_{\text{PR},1}$ in Table 2 we obtain $\sum_{a,b,x,y} F'(abxy)\mu_{\text{PR},1}(ab|xy)^{1+\beta} = 2^\epsilon > 1$, which is a violation.

References

1. Dodis, Y.; Ong, S.J.; Prabhakaran, M.; Sahai, A. On the (Im)Possibility of Cryptography with Imperfect Randomness. In Proceedings of the 45th Annual IEEE Symposium on Foundations of Computer Science, Rome, Italy, 17–19 October 2004; pp. 196–205. [\[CrossRef\]](#)
2. Austrin, P.; Chung, K.M.; Mahmoody, M.; Pass, R.; Seth, K. On the Impossibility of Cryptography with Tamperable Randomness. In Proceedings of the Advances in Cryptology—CRYPTO, Santa Barbara, CA, USA, 17–21 August 2014; Garay, J.A., Gennaro, R., Eds.; Publisher: Springer Berlin Heidelberg; Berlin/Heidelberg, Germany, 2014; pp. 462–479.
3. Dodis, Y.; Yao, Y. Privacy with Imperfect Randomness. In Advances in Cryptology—CRYPTO 2015; Gennaro, R., Robshaw, M., Eds.; Springer: Berlin/Heidelberg, Germany, 2015; pp. 463–482.
4. Orsini, C.; Dankulov, M.M.; Colomer-de Simón, P.; Jamakovic, A.; Mahadevan, P.; Vahdat, A.; Bassler, K.E.; Toroczkai, Z.; Boguñá, M.; Caldarelli, G.; et al. Quantifying randomness in real networks. *Nat. Commun.* **2015**, *6*, 8627. [\[CrossRef\]](#) [\[PubMed\]](#)

5. Motwani, R.; Raghavan, P. *Randomized Algorithms*; Cambridge University Press: Cambridge, NY, USA, 1995. [\[CrossRef\]](#)
6. Scarani, V. *Bell Nonlocality*; Oxford University Press: New York, NY, USA, 2019. [\[CrossRef\]](#)
7. Giustina, M.; Versteegh, M.A.M.; Wengerowsky, S.; Handsteiner, J.; Hochrainer, A.; Phelan, K.; Steinlechner, F.; Kofler, J.; Larsson, J.A.; Abellán, C.; et al. Significant-Loophole-Free Test of Bell's Theorem with Entangled Photons. *Phys. Rev. Lett.* **2015**, *115*, 250401. [\[CrossRef\]](#)
8. Shalm, L.K.; Meyer-Scott, E.; Christensen, B.G.; Bierhorst, P.; Wayne, M.A.; Stevens, M.J.; Gerrits, T.; Glancy, S.; Hamel, D.R.; Allman, M.S.; et al. Strong Loophole-Free Test of Local Realism. *Phys. Rev. Lett.* **2015**, *115*, 250402. [\[CrossRef\]](#)
9. Hensen, B.; Bernien, H.; Dréau, A.E.; Reiserer, A.; Kalb, N.; Blok, M.S.; Ruitenberg, J.; Vermeulen, R.F.L.; Schouten, R.N.; Abellán, C.; et al. Loophole-free Bell inequality violation using electron spins separated by 1.3 kilometres. *Nature* **2015**, *526*, 682–686. [\[CrossRef\]](#) [\[PubMed\]](#)
10. Rosenfeld, W.; Burchardt, D.; Garthoff, R.; Redeker, K.; Ortengel, N.; Rau, M.; Weinfurter, H. Event-Ready Bell Test Using Entangled Atoms Simultaneously Closing Detection and Locality Loopholes. *Phys. Rev. Lett.* **2017**, *119*, 010402. [\[CrossRef\]](#)
11. Bierhorst, P.; Knill, E.; Glancy, S.; Zhang, Y.; Mink, A.; Jordan, S.; Rommal, A.; Liu, Y.K.; Christensen, B.; Nam, S.W.; et al. Experimentally generated randomness certified by the impossibility of superluminal signals. *Nature* **2018**, *556*, 223–226. [\[CrossRef\]](#) [\[PubMed\]](#)
12. Shalm, L.K.; Zhang, Y.; Bienfang, J.C.; Schlager, C.; Stevens, M.J.; Mazurek, M.D.; Abellán, C.; Amaya, W.; Mitchell, M.W.; Alheji, M.A.; et al. Device-independent randomness expansion with entangled photons. *Nat. Phys.* **2021**, *17*, 452–456. [\[CrossRef\]](#)
13. Li, M.H.; Zhang, X.; Liu, W.Z.; Zhao, S.R.; Bai, B.; Liu, Y.; Zhao, Q.; Peng, Y.; Zhang, J.; Zhang, Y.; et al. Experimental Realization of Device-Independent Quantum Randomness Expansion. *Phys. Rev. Lett.* **2021**, *126*, 050503. [\[CrossRef\]](#)
14. Zhang, Y.; Knill, E.; Bierhorst, P. Certifying quantum randomness by probability estimation. *Phys. Rev. A* **2018**, *98*, 040304. [\[CrossRef\]](#)
15. Knill, E.; Zhang, Y.; Bierhorst, P. Generation of quantum randomness by probability estimation with classical side information. *Phys. Rev. Res.* **2020**, *2*, 033465. [\[CrossRef\]](#)
16. Bierhorst, P.; Zhang, Y. Tsirelson Polytopes and randomness generation. *New J. Phys.* **2020**, *22*, 083036. [\[CrossRef\]](#)
17. Zhang, Y.; Fu, H.; Knill, E. Efficient randomness certification by quantum probability estimation. *Phys. Rev. Res.* **2020**, *2*, 013016. [\[CrossRef\]](#)
18. Zhang, Y.; Shalm, L.K.; Bienfang, J.C.; Stevens, M.J.; Mazurek, M.D.; Nam, S.W.; Abellán, C.; Amaya, W.; Mitchell, M.W.; Fu, H.; et al. Experimental Low-Latency Device-Independent Quantum Randomness. *Phys. Rev. Lett.* **2020**, *124*, 010505. [\[CrossRef\]](#) [\[PubMed\]](#)
19. Renner, R.; Wolf, S. Simple and Tight Bounds for Information Reconciliation and Privacy Amplification. In Proceedings of the Advances in Cryptology—ASIACRYPT, Chennai, India, 4–8 December 2005; Springer: Berlin/Heidelberg, Germany, 2005; pp. 199–216.
20. Frank, R.L.; Lieb, E.H. Monotonicity of a relative Rényi entropy. *J. Math. Phys.* **2013**, *54*, 122201. [\[CrossRef\]](#)
21. Wilde, M.M. *Quantum Information Theory*; Cambridge University Press: Cambridge, NY, USA, 2013. [\[CrossRef\]](#)
22. Popescu, S.; Rohrlich, D. Quantum nonlocality as an axiom. *Found. Phys.* **1994**, *24*, 379–385. [\[CrossRef\]](#)
23. Fine, A. Hidden Variables, Joint Probability, and the Bell Inequalities. *Phys. Rev. Lett.* **1982**, *48*, 291–295. [\[CrossRef\]](#)
24. Bierhorst, P. Geometric decompositions of Bell polytopes with practical applications. *J. Phys. A Math. Theor.* **2016**, *49*, 215301. [\[CrossRef\]](#)
25. Le, T.P.; Meroni, C.; Sturmfels, B.; Werner, R.F.; Ziegler, T. Quantum Correlations in the Minimal Scenario. *Quantum* **2023**, *7*, 947. [\[CrossRef\]](#)
26. Brito, S.G.A.; Amaral, B.; Chaves, R. Quantifying Bell nonlocality with the trace distance. *Phys. Rev. A* **2018**, *97*, 022111. [\[CrossRef\]](#)
27. Mikos-Nuszkiewicz, A.; Kaniewski, J. Extremal points of the quantum set in the CHSH scenario: Conjectured analytical solution. *arXiv* **2023**, arXiv:2302.10658.
28. Collins, D.; Gisin, N.; Linden, N.; Massar, S.; Popescu, S. Bell Inequalities for Arbitrarily High-Dimensional Systems. *Phys. Rev. Lett.* **2002**, *88*, 040404. [\[CrossRef\]](#)
29. Barrett, J.; Linden, N.; Massar, S.; Pironio, S.; Popescu, S.; Roberts, D. Nonlocal correlations as an information-theoretic resource. *Phys. Rev. A* **2005**, *71*, 022101. [\[CrossRef\]](#)
30. Barrett, J.; Pironio, S. Popescu-Rohrlich Correlations as a Unit of Nonlocality. *Phys. Rev. Lett.* **2005**, *95*, 140401. [\[CrossRef\]](#)
31. Jones, N.S.; Masanes, L. Interconversion of nonlocal correlations. *Phys. Rev. A* **2005**, *72*, 052312. [\[CrossRef\]](#)
32. Pironio, S.; Bancal, J.D.; Scarani, V. Extremal correlations of the tripartite no-signaling polytope. *J. Phys. A Math. Theor.* **2011**, *44*, 065303. [\[CrossRef\]](#)
33. Greenberger, D.M.; Horne, M.A.; Zeilinger, A. Going Beyond Bell's Theorem. *arXiv* **2007**, arXiv:0712.0921.
34. Uhlmann, A. Entropy and Optimal Decompositions of States Relative to a Maximal Commutative Subalgebra. *Open Syst. Inf. Dyn.* **1998**, *5*, 209–228. [\[CrossRef\]](#)

35. Boyd, S.; Vandenberghe, L. *Convex Optimization*; Cambridge University Press: Cambridge, NY, USA, 2004.
36. Dodis, Y.; Reyzin, L.; Smith, A. Fuzzy Extractors: How to Generate Strong Keys from Biometrics and Other Noisy Data. In Proceedings of the Advances in Cryptology-EUROCRYPT 2004: International Conference on the Theory and Applications of Cryptographic Techniques, Interlaken, Switzerland, 2–6 May 2004; pp. 523–540. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.