

Parallelization of Adaptive Quantum Channel Discrimination in the Non-Asymptotic Regime

Bjarne Bergh^{ID}, Nilanjana Datta, Robert Salzmänn^{ID}, and Mark M. Wilde^{ID}, *Fellow, IEEE*

Abstract—We investigate the performance of parallel and adaptive quantum channel discrimination strategies for a finite number of channel uses. It has recently been shown that, in the asymmetric setting with asymptotically vanishing type I error probability, adaptive strategies are asymptotically not more powerful than parallel ones. We extend this result to the non-asymptotic regime with finitely many channel uses, by explicitly constructing a parallel strategy for any given adaptive strategy, and bounding the difference in their performances, measured in terms of the decay rate of the type II error probability per channel use. We further show that all parallel strategies can be optimized over in time polynomial in the number of channel uses, and hence our result can also be used to obtain a poly-time-computable asymptotically tight upper bound on the performance of general adaptive strategies.

Index Terms—Quantum information theory, Shannon theory, error exponents, channel discrimination, adaptive strategies, parallel strategies.

I. INTRODUCTION

BINARY quantum state discrimination is one of the oldest and most studied tasks in quantum information theory [1], [2]. The task involves determining the state of a quantum system, given the side information that it is in one of two possible states. This is done by performing suitable measurements on the state. This task has been studied in both the n -shot regime, in which a finite number (n) of identical copies of the unknown state are available, and in the asymptotic limit in which one assumes that an infinite number of copies are available (i.e., $n \rightarrow \infty$). The task of binary state discrimination,

which can be viewed as that of binary hypothesis testing for quantum states, is a fundamental primitive of quantum information theory since many other information processing tasks can be reduced to it. Known results for this problem now include optimal decision strategies and the behavior of the error probabilities of misidentification in both these finite and asymptotic regimes [1], [3], [4], [5], [6], [7], [8], [9], [10], [11], [12]. There are generally two possible errors that can be incurred: the system is inferred to be in the second state when it is actually in the first (*type I error*) or vice versa (*type II error*). There are two different settings in which the state discrimination task is usually studied: in the *symmetric setting*, the probabilities of these two errors are treated on an equal footing, while in the *asymmetric setting*, one minimizes the probability of type II error under the constraint that the type I error is below a given threshold.

Similar in nature, but a lot less understood, is the task of binary quantum channel discrimination: given an unknown quantum channel as a black box and the side information that it is one of two possible channels, the task is to determine the channel's identity [13], [14], [15], [16]. The additional complexity here comes from the fact that, on top of finding the best measurement to perform on the output of the channel, we also have to figure out which quantum states to send as input to the channel. Say we are given access to n copies of the channel (i.e., we are given n identical black boxes, each of which can be used once); then there are different strategies (sometimes also called protocols) in which we could set up our decision experiment – the so-called *parallel* and *adaptive* strategies. In a parallel strategy one prepares a joint state, usually entangled between the input systems of all the n copies of the channel and an additional reference (or memory) system. This state is then fed as input to all the n channels at once (with the state of the reference system being left undisturbed). Finally, a binary positive operator-valued measure (POVM) is performed on the joint state at the output of the channels and the reference system in order to arrive at a decision for the channel's identity. In an adaptive strategy, on the other hand, one prepares a state of the input system of a single copy of the channel (again usually entangled with a reference system) which is fed into the first copy of the channel, with the state of the reference system being left undisturbed. The input to the next use of the channel is then chosen depending on the output of the first channel and the state of our reference system. This is done, most generally, by subjecting the latter to an arbitrary quantum operation (or

Manuscript received 25 July 2022; revised 3 November 2023; accepted 29 November 2023. Date of publication 31 January 2024; date of current version 19 March 2024. The work of Bjarne Bergh was supported by the U.K. Engineering and Physical Sciences Research Council (EPSRC) under Grant EP/V52024X/1. The work of Robert Salzmänn was supported by the Cambridge Commonwealth, European and International Trust. The work of Mark M. Wilde was supported by the National Science Foundation under Grant 2315398. An earlier version of this paper was presented at the Beyond IID in Information Theory 10 Conference, 2022. (*Corresponding author: Bjarne Bergh.*)

Bjarne Bergh, Nilanjana Datta, and Robert Salzmänn are with the Department of Applied Mathematics and Theoretical Physics, University of Cambridge, CB3 0WA Cambridge, U.K. (e-mail: bb536@cam.ac.uk; n.datta@statslab.cam.ac.uk; rals.salzmänn@web.de).

Mark M. Wilde is with the Department of Physics and Astronomy, Center for Computation and Technology, Hearne Institute for Theoretical Physics, Louisiana State University, Baton Rouge, LA 70803 USA, and also with the School of Electrical and Computer Engineering, Cornell University, Ithaca, NY 14850 USA (e-mail: wilde@cornell.edu).

Communicated by M. Berta, Associate Editor for Quantum.

Color versions of one or more figures in this article are available at <https://doi.org/10.1109/TIT.2024.3355929>.

Digital Object Identifier 10.1109/TIT.2024.3355929

channel), which we call a preparation operation. This step is repeated for each successive use of the channel until all the n black boxes have been used. Then a binary POVM is performed on the joint state of the output of the last use of the channel and the reference system. See Figure 1 for a depiction of an adaptive strategy. Adaptive strategies are also sometimes called sequential, which however should not be confused with the setting of sequential hypothesis testing [17], [18], [19], where samples (i.e., states or channels) can be requested one by one.

One particularly interesting question is whether and to what degree adaptive strategies give an advantage over parallel ones. Note that any parallel strategy can be written as an adaptive strategy by taking all but one channel input as part of the reference system, and then choosing each preparation operation such that it extracts the next part of the joint input state for the next channel use and replaces it by the output of the previous channel use. However, the converse is not true, and so adaptive strategies are more general. Parallel strategies are conceptually a lot simpler than adaptive ones – aside from the measurement, everything is specified just by the joint input state – in contrast to adaptive strategies, in which after each channel use we can perform an arbitrary quantum operation to prepare the input to the next use of the channel. It is thus interesting to determine to what degree parallel strategies can still be optimal. It is known that in certain cases adaptive strategies can give an advantage over parallel ones. In [16] the authors constructed an example in which an adaptive strategy with only two channel uses could be used to discriminate the channels with certainty, which is not possible with a parallel strategy, even if arbitrarily many channel uses are allowed.

Interestingly, *asymptotically*, there are multiple known cases in which adaptive strategies give no advantage over parallel ones, i.e., the optimal exponential decay rate of the error probability per channel use is the same in the asymptotic limit. For example, this is the case both in the symmetric and asymmetric settings when the channels are classical [15] or classical-quantum [20], [21]. For arbitrary quantum channels, the recently shown chain rule for the quantum relative entropy [22] and the characterization of asymmetric channel discrimination in terms of amortized relative entropy [20], [23], also imply that in the asymmetric setting, in the asymptotic limit where we also require the type I error to vanish, adaptive strategies give no advantage over parallel ones (i.e., the optimal asymptotic exponential decay rate of the type II error per channel use is the same for parallel and adaptive strategies; see Section II below for more details). This is in contrast to the symmetric setting in which the example of [21] shows that there always is an advantage of adaptive strategies, also in the asymptotic limit.

In this paper, our aim is to move from these purely asymptotic results to statements comparing adaptive and parallel strategies for finite n . In our main result (Corollary 1 and Theorem 1), we relate the type I and type II errors of an arbitrary adaptive strategy with those of a suitably chosen parallel strategy. Specifically, given an adaptive strategy with n channel uses, we construct a parallel strategy with m channel uses (where m can be chosen at will), such that for arbitrary

fixed type I errors of the two strategies, we find an explicit bound on the difference between their type II error decay rates. This difference goes to zero in a suitably chosen asymptotic limit $m, n \rightarrow \infty$ if also the type I error vanishes, hence also implying the known asymptotic equivalence in the asymmetric case. Our result answers the following interesting question which is of practical relevance: *Given an adaptive strategy involving n uses of the channel, if one instead employs a parallel strategy with m channel uses, how much worse are the errors going to be?*

Note that the asymptotic results obtained in [20] and [23] do not purely come from finite-length considerations, and hence results analogous to ours are not directly obtainable from these references.

Our main result becomes particularly interesting considering that we additionally show that one can optimize over all parallel strategies in time polynomial in the number of channel uses (Lemma 3), whereas optimizing over all adaptive strategies seems to require exponential time (see Remark 6 below). Hence we can also use our theorem to give a bound related to the following practically relevant question: *Given that one computed the optimal parallel strategy involving m uses of the channel, how much better could the errors of the best adaptive strategy with n channel uses possibly be?*

Note that the upper bound on the finite- n performance of adaptive strategies that can be obtained like this by optimizing over all parallel strategies, is to our knowledge the first upper bound on this quantity that is computable in time polynomial in the number of channel uses and also asymptotically tight.

On the way to proving our main theorem, we also prove the following two technical results: first, a slightly refined bound on the difference between the Petz–Rényi relative entropy D_α and the Umegaki relative entropy D (Lemma 5), and secondly, a bound on the convergence speed in the asymptotic equipartition property of the smoothed max-relative entropy (Lemma 6), which turns out to be better suited to our application than previous results (see Remark 8 and Remark 9 for a detailed explanation of the differences).

We start by introducing all of the required notation and previous results in Section II. Then we state and prove our main theorem (Theorem 1) in Section III. Section IV illustrates how adaptive and parallel strategies can differ in the finite regime with an example, and it also demonstrates an application of our result.

II. PRELIMINARIES, NOTATION, AND PREVIOUS RESULTS

A. Notation

We write $\mathcal{H}$ for a complex finite-dimensional Hilbert space and $\mathcal{B}(\mathcal{H})$ for the set of linear operators acting on $\mathcal{H}$. We write $\mathcal{P}(\mathcal{H})$ for the set of positive semi-definite operators acting on $\mathcal{H}$. For $A, B \in \mathcal{P}(\mathcal{H})$ we further write $A \ll B$ if $\text{supp}(A) \subseteq \text{supp}(B)$ and $A \not\ll B$ if $\text{supp}(A) \not\subseteq \text{supp}(B)$. Let $\mathcal{D}(\mathcal{H})$ denote the set of density matrices, i.e., the set of positive semi-definite operators with trace one. A quantum channel (in this paper usually denoted as $\mathcal{E}$ or $\mathcal{F}$) is a completely positive, trace-preserving map between density operators. We will label different quantum systems by capital Roman letters (A, B ,

C , etc.) and often use these letters interchangeably with the corresponding Hilbert space or set of operators or density matrices (i.e., we write $\rho \in \mathcal{D}(A)$ instead of $\rho \in \mathcal{D}(\mathcal{H}_A)$ and $\mathcal{E} : A \rightarrow B$ instead of $\mathcal{E} : \mathcal{D}(\mathcal{H}_A) \rightarrow \mathcal{D}(\mathcal{H}_B)$). We write $\log$ for the logarithm to the base two.

B. Quantum Information Measures

For $\rho \in \mathcal{D}(\mathcal{H})$ and $\sigma \in \mathcal{P}(\mathcal{H})$ the (Umegaki) quantum relative entropy is defined as [24]

$$D(\rho\|\sigma) := \text{Tr}(\rho(\log \rho - \log \sigma)), \quad (1)$$

if $\rho \ll \sigma$ and $D(\rho\|\sigma) := \infty$ if $\rho \not\ll \sigma$. One of its most important properties is the data-processing inequality [25], which states that for every quantum channel $\mathcal{E}$:

$$D(\rho\|\sigma) \geq D(\mathcal{E}(\rho)\|\mathcal{E}(\sigma)). \quad (2)$$

A self-contained proof can be found, e.g., in [26]. More generally, we call a function of ρ and σ a divergence if it satisfies the data-processing inequality.

Below we also make use of the binary entropy, which is the Shannon entropy of a classical random variable that takes two possible values. It is uniquely specified by the probability $p \in [0, 1]$ of one of the values and is given by

$$h(p) := -p \log p - (1-p) \log(1-p). \quad (3)$$

It satisfies

$$0 \leq h(p) \leq 1 \quad \forall p \in [0, 1]. \quad (4)$$

1) *Fidelity and Sine Distance*: For two quantum states $\rho, \sigma \in \mathcal{D}(\mathcal{H})$, we define the fidelity as [27]

$$F(\rho, \sigma) := \text{Tr} \left(\sqrt{\sqrt{\sigma} \rho \sqrt{\sigma}} \right). \quad (5)$$

Note that this is sometimes also called the square root fidelity due to different conventions on whether to include a square or not. We define the sine distance as [28], [29], [30], [31]

$$P(\rho, \sigma) := \sqrt{1 - F(\rho, \sigma)^2}, \quad (6)$$

which satisfies the properties of a metric (especially the triangle inequality) and also the data-processing inequality (see, e.g., [26]). This quantity is also known as the purified distance.

2) *(Smoothed) Max-Divergence*: For $\rho \in \mathcal{D}(\mathcal{H})$ and $\sigma \in \mathcal{P}(\mathcal{H})$, define the quantum max-divergence (or the max-relative entropy) as [32]

$$D_{\max}(\rho\|\sigma) := \log \inf \{ \lambda \in \mathbb{R} \mid \rho \leq \lambda \sigma \}. \quad (7)$$

The quantum max-divergence also satisfies the data-processing inequality [32]. Let

$$B_{\varepsilon}^{\circ}(\rho) := \{ \tilde{\rho} \in \mathcal{D}(\mathcal{H}) \mid P(\rho, \tilde{\rho}) \leq \varepsilon \} \quad (8)$$

be the ε -ball of *normalized* states around ρ in sine distance. Then, we define the smoothed max-divergence as [32]

$$D_{\max}^{\varepsilon}(\rho\|\sigma) := D_{\max}^{\varepsilon, \circ}(\rho\|\sigma) = \inf_{\tilde{\rho} \in B_{\varepsilon}^{\circ}(\rho)} D_{\max}(\tilde{\rho}\|\sigma). \quad (9)$$

Note that this is sometimes defined differently in the literature, where one allows the infimum to also include sub-normalized states.

3) *Rényi Divergences*: Let $\rho \in \mathcal{D}(\mathcal{H})$ and $\sigma \in \mathcal{P}(\mathcal{H})$ be two operators such that $\rho \ll \sigma$. The definitions below can, in some cases, also be extended with finite values to the case where $\rho \not\ll \sigma$ (see the references below or also [26]); however, this is not going to be relevant for our work. For every $\alpha \in [0, 1) \cup (1, 2]$ define the *Petz–Rényi divergence* as [33]:

$$D_{\alpha}(\rho\|\sigma) := \frac{1}{\alpha - 1} \log \text{Tr}(\rho^{\alpha} \sigma^{1-\alpha}). \quad (10)$$

Similarly, for every $\alpha \in (0, 1) \cup (1, 2]$ define the *geometric Rényi divergence* as [34]:

$$\hat{D}_{\alpha}(\rho\|\sigma) := \frac{1}{\alpha - 1} \log \text{Tr} \left(\sigma^{\frac{1}{2}} (\sigma^{-\frac{1}{2}} \rho \sigma^{-\frac{1}{2}})^{\alpha} \sigma^{\frac{1}{2}} \right). \quad (11)$$

It was shown in [34] that the geometric Rényi divergence is the largest Rényi divergence for every $\alpha \in (0, 1) \cup (1, 2]$, and so

$$D_{\alpha}(\rho\|\sigma) \leq \hat{D}_{\alpha}(\rho\|\sigma). \quad (12)$$

4) *Channel Divergences*: We say that a function $\mathbf{D}$ of $\rho \in \mathcal{D}(\mathcal{H})$ and $\sigma \in \mathcal{P}(\mathcal{H})$ is a divergence if it satisfies the data-processing inequality (recall (2) here). For every given divergence $\mathbf{D}$ for states, one can define an associated channel divergence [35] by performing a (stabilized) maximization over all input states, i.e., with $\mathcal{E}, \mathcal{F} : A \rightarrow B$ being quantum channels

$$\mathbf{D}(\mathcal{E}\|\mathcal{F}) := \sup_{\rho_{RA} \in \mathcal{D}(R \otimes A)} \mathbf{D}((\text{id}_R \otimes \mathcal{E})(\rho) \| (\text{id}_R \otimes \mathcal{F})(\rho)). \quad (13)$$

Since $\mathbf{D}$ satisfies the data-processing inequality by definition, the supremum can be restricted to pure states such that the reference system R is isomorphic to the channel input system A .

Specifically relevant later will be the max channel divergence, which has the following simple representation (see [36, Definition 19] and [20, Lemma 12]):

$$\begin{aligned} D_{\max}(\mathcal{E}\|\mathcal{F}) &:= \sup_{\psi_{RA}} D_{\max}((\text{id}_R \otimes \mathcal{E})(\psi) \| (\text{id}_R \otimes \mathcal{F})(\psi)) \\ &= D_{\max}((\text{id}_R \otimes \mathcal{E})(\Phi) \| (\text{id}_R \otimes \mathcal{F})(\Phi)), \end{aligned} \quad (14)$$

where $R \simeq A$ and $\Phi = \Phi_{RA}$ is a maximally entangled state. Hence, the max channel divergence is easily computable as a semidefinite program (SDP).

C. Quantum State Discrimination

In binary quantum state discrimination, a quantum system with Hilbert space $\mathcal{H}$ is prepared in one of two states ρ or σ , and the objective is to perform a binary POVM $\{\Pi, \mathbb{1} - \Pi\}$ (where $\mathbb{1}$ denotes the identity operator acting on $\mathcal{H}$ and $\Pi \in \mathcal{B}(\mathcal{H})$ satisfying $0 \leq \Pi \leq \mathbb{1}$) in order to guess which state was prepared (see [12] for a review). Performing this measurement comes with the possibility of making an error of misidentification. Generally one defines the type I and type II error probabilities, respectively, as

$$\alpha(\Pi, \rho, \sigma) := \text{Tr}((\mathbb{1} - \Pi)\rho), \quad (16)$$

$$\beta(\Pi, \rho, \sigma) := \text{Tr}(\Pi\sigma). \quad (17)$$

In the rest of the paper, we refer to the above simply as the type I and type II errors. It is easy to see from the definition that different choices of Π lead to different trade-offs between the type I and type II errors.

In the asymmetric setting, we are interested in finding the optimal type II error while keeping the type I error below a threshold (usually denoted as ε). This optimal type II error is then defined for all $\varepsilon \in [0, 1]$ as

$$\beta_\varepsilon(\rho\|\sigma) := \min_{\substack{0 \leq \Pi \leq \mathbb{1} \\ \text{Tr}(\Pi\rho) \geq 1-\varepsilon}} \text{Tr}(\Pi\sigma). \quad (18)$$

The negative logarithm of this optimal type II error is called the hypothesis testing relative entropy [37], [38], [39]:

$$D_H^\varepsilon(\rho\|\sigma) := -\log \beta_\varepsilon(\rho\|\sigma). \quad (19)$$

It is also known as the smooth min-relative entropy [40], [41]. The hypothesis testing relative entropy satisfies the data-processing inequality [39] and can be related to other entropic quantities. We will make use of the following known relations:

Lemma 1 (Upper Bound on D_H^ε): Let $\rho, \sigma \in \mathcal{D}(\mathcal{H})$ be quantum states. Then for all $\varepsilon \in [0, 1]$

$$D_H^\varepsilon(\rho\|\sigma) \leq \frac{1}{1-\varepsilon} (D(\rho\|\sigma) + h(\varepsilon)), \quad (20)$$

where $h(\varepsilon)$ is the binary entropy function.

This first Lemma is a very well-known consequence of the data-processing inequality of the relative entropy and has been used in converse proofs all throughout information theory. A statement and proof using our notation can be found in [39], although the essence of the statement can already be found much earlier, for example in [4, Theorem 2.2] and also [42, Eq. (3.30)]. We will also make use of the following Lemma:

Lemma 2 (Relation to Smoothed Max-Divergence [43, Theorem 4]): Let $\rho \in \mathcal{D}(\mathcal{H})$ and $\sigma \in \mathcal{P}(\mathcal{H})$. Then, for $\varepsilon \in (0, 1)$ and $\delta \in (0, 1 - \varepsilon^2)$:

$$D_H^{1-\varepsilon^2-\delta}(\rho\|\sigma) - \log\left(\frac{4(1-\varepsilon^2)}{\delta^2}\right) \leq D_{\max}^\varepsilon(\rho\|\sigma) \leq D_H^{1-\varepsilon^2}(\rho\|\sigma) - \log(1-\varepsilon^2). \quad (21)$$

Remark 1: Note that [43] defines the smoothed max-divergence using sub-normalized smoothing; however, the proof of this lemma goes through also when restricting to normalized states. Note also that the arXiv version (v1) of [43] has a typo in the admissible range of δ , which has been corrected in the published version.

D. Quantum Channel Discrimination

Let $\mathcal{E}$ and $\mathcal{F}$ be two quantum channels, each taking system A to system B . In the discrimination setting we will generally use inputs to the channels that are entangled with a reference system. To simplify notation we will usually *not* make the identity channel on the reference system explicit. Hence, if $\rho_{RA} \in \mathcal{D}(R \otimes A)$ is a state, we will write

$$\mathcal{E}(\rho) := \mathcal{E}_{A \rightarrow B}(\rho_{RA}) = (\text{id}_R \otimes \mathcal{E}_{A \rightarrow B})(\rho_{RA}), \quad (22)$$

and similarly also for $\mathcal{F}$.

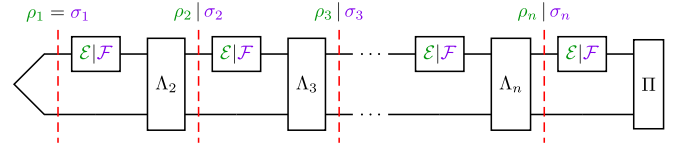


Fig. 1. Illustration of a general adaptive protocol with n uses of the black-box channel. The top row makes use of the given black-box $\mathcal{E}|\mathcal{F}$, which is either $\mathcal{E}$ or $\mathcal{F}$, while the bottom row depicts the memory system R . At various stages in the protocol, the green states ρ occur if the channel is $\mathcal{E}$ and the purple states σ occur if the channel is $\mathcal{F}$.

The most general channel discrimination protocol will choose input states based on the outputs of previous channel uses. This is called an adaptive protocol. A general adaptive channel discrimination protocol with n uses of the black-box channel ($\mathcal{E}$ or $\mathcal{F}$), can be fully specified by an initial state $\rho_1 = \sigma_1 \in \mathcal{D}(R \otimes A)$, a set of $n-1$ CPTP maps $\Lambda_i : R \otimes B \rightarrow R \otimes A$, that transform the state before it is fed into the next black-box channel, and a final binary POVM $\{\Pi, \mathbb{1} - \Pi\}$ on $R \otimes B$. We will assume the size of reference system R to be fixed and identical throughout the protocol (this is without loss of generality). The protocol consists of alternating applications of the black-box channel and the preparation CPTP maps Λ_i (see Figure 1). We define for $i \in \{2, \dots, n\}$:

$$\rho_i := \Lambda_i(\mathcal{E}(\rho_{i-1})), \quad \sigma_i := \Lambda_i(\mathcal{F}(\sigma_{i-1})), \quad (23)$$

and so the final state before the action of the POVM will be $\mathcal{E}(\rho_n)$ if the channel is $\mathcal{E}$ and $\mathcal{F}(\sigma_n)$ if the channel is $\mathcal{F}$. The distinguishability of the channels then comes down to the distinguishability of the two states $\mathcal{E}(\rho_n)$ and $\mathcal{F}(\sigma_n)$. We will be focussing again on the asymmetric setting, with the main object of study being the hypothesis testing relative entropy $D_H^\varepsilon(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n))$. We write the optimal adaptive type II error rate for a given finite number of channel uses n as [44]

$$\sup_{\rho_1, \{\Lambda_i\}_i} \frac{1}{n} D_H^\varepsilon(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) \quad (24)$$

where the supremum is over every initial input state $\rho_1 = \sigma_1$ and all subsequent input preparation CPTP maps $\{\Lambda_i\}_{i=2}^n$. It was shown in [45, Prop. 3] that this is computable as a semi-definite program.

1) Asymptotic Equivalence of Adaptive and Parallel Strategies: It has been shown recently [20], [23] that asymptotically, when the number of channel uses goes to infinity, the best exponential decay rate of the type II error (per channel use) such that the type I error still goes to zero is given by the amortized Umegaki channel divergence, i.e.

$$\lim_{\varepsilon \rightarrow 0} \lim_{n \rightarrow \infty} \sup_{\rho_1, \{\Lambda_i\}_i} \frac{1}{n} D_H^\varepsilon(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) = D^A(\mathcal{E}\|\mathcal{F}) \quad (25)$$

where

$$D^A(\mathcal{E}\|\mathcal{F}) := \sup_{\rho, \sigma \in \mathcal{D}(R \otimes A)} [D(\mathcal{E}(\rho)\|\mathcal{F}(\sigma)) - D(\rho\|\sigma)]. \quad (26)$$

Note that the dimension of the reference system R in this last supremum can be arbitrarily large.

The chain rule from [22] states that this amortized divergence is in fact equal to the regularized channel divergence, i.e.

$$D^A(\mathcal{E}||\mathcal{F}) = D^{\text{reg}}(\mathcal{E}||\mathcal{F}) := \lim_{n \rightarrow \infty} \frac{1}{n} \sup_{\nu \in \mathcal{D}(R^{\otimes n} \otimes A^{\otimes n})} D(\mathcal{E}^{\otimes n}(\nu)||\mathcal{F}^{\otimes n}(\nu)), \quad (27)$$

the latter of which can be achieved by parallel protocols, as a consequence of [23, Theorem 3]. Note that the reference system R in the latter optimization can be chosen isomorphic to A . Hence, asymptotically (and in the regime in which the type I error goes to zero) adaptive strategies offer no advantage over parallel ones. We provide a finite n version of this statement, by giving explicit bounds on how much the error probabilities of adaptive and parallel strategies can differ for a finite number of channel uses.

III. PARALLELIZING AN n -SHOT ADAPTIVE PROTOCOL

We state our main result in two forms, first in a simple manner that illustrates the main idea and structure of the result, and secondly a more detailed theorem that gives a tighter bound, and which states in detail what one can choose as a parallel input state.

Remember that there is a tradeoff in minimizing the type I and type II errors. In the context of a strategy, for a given type I error α we will write the best achievable type II error as $\beta(\alpha)$. We are especially interested in the exponential decay rate of the type II error with the number of channel uses; i.e., if some strategy involving n uses of the channel has type II error $\beta(\alpha)$, we are interested in the quantity $-\frac{1}{n} \log(\beta(\alpha))$. Our main result compares this error decay rate per channel use between adaptive and parallel strategies:

Corollary 1 (Main Result, Simple Version): Let $\mathcal{E}, \mathcal{F} : A \rightarrow B$ be two quantum channels such that $D_{\max}(\mathcal{E}||\mathcal{F}) < \infty$. Let there be an adaptive discrimination protocol with n channel uses, that – for an arbitrary type I error $\alpha_a \in [0, 1]$ – achieves type II error $\beta_a(\alpha_a)$. Then, for all $\alpha_p \in (0, 1]$ there exists a parallel protocol with m channel uses and type II error $\beta_p(\alpha_p)$ such that for all $\alpha_a \in [0, 1]$ the type II error rates per channel use obey the following relation:

$$-\frac{1}{m} \log(\beta_p(\alpha_p)) \geq -\frac{1 - \alpha_a}{n} \log(\beta_a(\alpha_a)) - \frac{Cn}{\sqrt{m}} \log\left(\frac{8}{\alpha_p}\right) - \frac{1}{n}. \quad (28)$$

That is, the type II error rate of the parallel protocol is essentially at least as good as the adaptive one modulo an additional error term, which decays as $m \rightarrow \infty$. The constant C is given by

$$C := 7 \log(2^{D_2(\mathcal{E}||\mathcal{F})} + 2) \leq 7 \log(2^{D_{\max}(\mathcal{E}||\mathcal{F})} + 2). \quad (29)$$

Remark 2: If we take the limit $m \rightarrow \infty$ in (28), then $n \rightarrow \infty$, and finally $\alpha_a \rightarrow 0$ and $\alpha_p \rightarrow 0$, we find that asymptotically there exist parallel strategies with better or at least equal type II error decay than an adaptive one, and hence our result also implies the known result [22] that,

asymptotically (and with $\alpha \rightarrow 0$), adaptive strategies offer no advantage over parallel ones.

Remark 3: It is known that for small n and m , and suitably chosen α_a and α_p , the type II error rates for adaptive and parallel strategies can be arbitrarily far apart (see Section IV below for an example). From the asymptotic equivalence, we know that this difference has to vanish as n and m go to infinity, but the purely asymptotic statement does not tell us how exactly this vanishes, and what the required relationship between n and m is (in principle the required m to reach a similar rate as a given adaptive strategy with n channel uses could grow arbitrarily fast with n). Corollary 1 now tells us that the difference of type II error rates between an adaptive strategy and the corresponding parallel one will become arbitrarily small if $m = \omega(n^2)$ (i.e., m has to grow faster than n^2). Hence, given a sequence of adaptive strategies with n channel uses, we can convert these into parallel strategies using at most quadratically (or a little bit more than quadratically) as many channel uses each, and will achieve matching rates once n gets large enough. This quadratic relationship is universal in the sense that it holds for all pairs of channels $\mathcal{E}, \mathcal{F}$ with $D_{\max}(\mathcal{E}||\mathcal{F}) < \infty$ (where only the prefactor depends on the value of $D_{\max}(\mathcal{E}||\mathcal{F})$).

Note again that we are not comparing type II errors in Corollary 1, but rather decay rates of the type II error per channel use. When we say that the rates of a parallel strategy with $m = \omega(n^2)$ channel uses and an adaptive strategy with n channel uses are roughly equal, the parallel strategy will have much smaller type II error because it has many more channel uses.

The following is the more refined version of our result, with a tighter bound and a description of the parallel input state:

Theorem 1 (Main Result, Technical Version): Let $\mathcal{E}, \mathcal{F} : A \rightarrow B$ be quantum channels such that $D_{\max}(\mathcal{E}||\mathcal{F}) < \infty$. Given an arbitrary adaptive protocol with n channel uses, we write $\rho_i, \sigma_i \in \mathcal{D}(R_a \otimes A)$, $i \in \{1, \dots, n\}$ for the states that are input into the channel during the adaptive protocol (ρ_i if the channel is $\mathcal{E}$, and σ_i if the channel is $\mathcal{F}$; see Section II-D and Figure 1 above for a more detailed explanation of this notation). We define $\ell \in \{1, \dots, n\}$ as the step in the protocol where the distinguishability increases the most, i.e.,

$$\ell := \arg \max_{k \in \{1, \dots, n\}} [D(\mathcal{E}(\rho_k)||\mathcal{F}(\sigma_k)) - D(\rho_k||\sigma_k)]. \quad (30)$$

Then, for all $\alpha_p \in (0, 1]$, and $m \in \mathbb{N}$, there exists a state $\nu \in \mathcal{D}(R^{\otimes m} \otimes A^{\otimes m})$ such that for all $\alpha_a \in [0, 1]$:

$$\begin{aligned} \frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu)||\mathcal{F}^{\otimes m}(\nu)) &\geq \frac{1 - \alpha_a}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n)||\mathcal{F}(\sigma_n)) \\ &- \frac{c'_\ell}{\sqrt{m}} \left[\log\left(\frac{4}{\alpha_p}\right) + K \right] - \frac{1}{m} \left[\log\left(\frac{1}{\alpha_p}\right) - \log\left(1 - \frac{\alpha_p}{4}\right) \right] \\ &- \frac{h(\alpha_a)}{n}, \end{aligned} \quad (31)$$

where

$$K := \frac{\ln(2) \log^2(3)}{8} \cosh\left(\frac{\log(3)}{2}\right) \leq 0.29 \quad (32)$$

and c'_ℓ depends on the pair of channels $\mathcal{E}, \mathcal{F}$ and can be bounded as follows:

$$c'_\ell := \frac{4}{\log(3)} \inf_{\gamma_1, \gamma_2 \in (0,1]} [c_{\gamma_1}(\mathcal{E}(\rho_\ell) \|\mathcal{F}(\sigma_\ell)) + c_{\gamma_2}(\rho_\ell \|\sigma_\ell)] \quad (33)$$

$$\leq \frac{8\ell}{\log(3)} \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \|\mathcal{F}) \quad (34)$$

$$\leq \frac{8n}{\log(3)} \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \|\mathcal{F}). \quad (35)$$

where

$$c_\gamma(\rho \|\sigma) := \frac{1}{\gamma} \log \left(2^{\gamma D_{1+\gamma}(\rho \|\sigma)} + 2^{-\gamma D_{1-\gamma}(\rho \|\sigma)} + 1 \right), \quad (36)$$

$$\hat{c}_\gamma(\mathcal{E} \|\mathcal{F}) := \frac{1}{\gamma} \log \left(2^{\gamma \hat{D}_{1+\gamma}(\mathcal{E} \|\mathcal{F})} + 2 \right). \quad (37)$$

Moreover, if

$$m \geq \log \left(\frac{4}{\alpha_p} \right) \left(\frac{4}{\log(3) \sqrt{2 \ln(2)}} \right)^2, \quad (38)$$

then we have the following tighter bound

$$\begin{aligned} \frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu) \|\mathcal{F}^{\otimes m}(\nu)) &\geq \frac{1 - \alpha_a}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n) \|\mathcal{F}(\sigma_n)) \\ &- \frac{c_\ell}{\sqrt{m}} \sqrt{\log \left(\frac{4}{\alpha_p} \right)} - \frac{1}{m} \left[\log \left(\frac{1}{\alpha_p} \right) - \log \left(1 - \frac{\alpha_p}{4} \right) \right] \\ &- \frac{h(\alpha_a)}{n}, \end{aligned} \quad (39)$$

where c_ℓ is defined in a similar way as c'_ℓ but with different numerical constants K_1 and K_2 :

$$K_1 := 2 \sqrt{2 \ln(2) \cosh \left(\frac{\log(3)}{2} \right)} \leq 2.72, \quad (40)$$

$$K_2 := 2 \sqrt{2 \ln(2)} \leq 2.36, \quad (41)$$

$$c_\ell := \inf_{\gamma_1, \gamma_2 \in (0,1]} [K_1 c_{\gamma_1}(\mathcal{E}(\rho_\ell) \|\mathcal{F}(\sigma_\ell)) + K_2 c_{\gamma_2}(\rho_\ell \|\sigma_\ell)] \quad (42)$$

$$\leq \ell(K_1 + K_2) \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \|\mathcal{F}) \quad (43)$$

$$\leq n(K_1 + K_2) \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \|\mathcal{F}). \quad (44)$$

The parallel input state ν can be chosen as either:

- An optimizer in the smoothing of the max-divergence $D_{\max}^\varepsilon(\rho_\ell^{\otimes m} \|\sigma_\ell^{\otimes m})$, where $\varepsilon = \frac{1}{2} (1 - \sqrt{1 - \alpha_p})$, i.e.,

$$\nu = \tilde{\nu}_{R_a^m} := \arg \min_{\tilde{\rho} \in B_\varepsilon^\circ(\rho_\ell^{\otimes m})} D_{\max}(\tilde{\rho} \|\sigma_\ell^{\otimes m}). \quad (45)$$

Note that the reference system R_a depends on the adaptive protocol and can be arbitrarily large, and hence also $\tilde{\nu}$ might have an arbitrarily large reference system.

- The canonical (or any other) purification of the A^m -marginal $\tilde{\nu}_{A^m} = \text{Tr}_{R_a^m}(\tilde{\nu}_{R_a^m A^m})$. That is, we can choose $\nu = |\psi_{R^m A^m}\rangle\langle\psi_{R^m A^m}|$, with

$$|\psi_{R^m A^m}\rangle = \mathbb{1}_{R^m} \otimes \sqrt{\tilde{\nu}_{A^m}} |\Phi_{R^m A^m}\rangle \quad (46)$$

where R is isomorphic to A and $|\Phi_{R^m A^m}\rangle$ is an unnormalized maximally entangled state.

Remark 4: Even though the second choice for the parallel input state ν might seem preferable in most cases (due to the control over the size of the reference system R), it is not necessarily so. This is because even though the reference system of the first choice for ν could be very large, it is not always so. Since the overall state can be mixed, it might actually be smaller than the canonical purification of its marginal. Additionally, with the first choice for ν , one gets a bound that this parallel input state is ε -close to the input state of the adaptive strategy, which might be useful in cases where one wants to show that some property of the adaptive strategy (e.g., a satisfied energy constraint for the input states) is (approximately) satisfied also for the parallel strategy.

Remark 5: The constraint $D_{\max}(\mathcal{E} \|\mathcal{F}) < \infty$ is necessary in general, as the example of [16] (see also [21]) serves as a counterexample to our statement without this constraint. Specifically, [16] constructs two channels $\mathcal{E}, \mathcal{F}$ for which its authors then show that there exists an adaptive strategy with only two channel uses that achieves perfect discrimination, i.e., both $\alpha_a = 0$ and $\beta_a = 0$. In our terminology this implies that for all $\alpha_a \in [0, 1]$

$$D_H^{\alpha_a}(\mathcal{E}(\rho_2) \|\mathcal{F}(\sigma_2)) = \infty. \quad (47)$$

On the other hand, [16] shows for these two channels that with a parallel strategy, even with arbitrarily many channel uses, perfect discrimination can never be achieved (i.e., α_p and β_p cannot both be zero), which in our notation implies that for all m and for all $\nu \in \mathcal{D}(R \otimes A^{\otimes m})$

$$D_H^0(\mathcal{E}^{\otimes m}(\nu) \|\mathcal{F}^{\otimes m}(\nu)) < \infty, \quad (48)$$

Since it is well known (see, e.g., [26, Prop. 4.66]) that for all states ρ, σ

$$\lim_{\alpha_p \rightarrow 0} D_H^{\alpha_p}(\rho \|\sigma) = D_H^0(\rho \|\sigma) \quad (49)$$

this means that for all m , one can find a sufficiently small α_p such that for all $\nu \in \mathcal{D}(R \otimes A^{\otimes m})$

$$D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu) \|\mathcal{F}^{\otimes m}(\nu)) < \infty. \quad (50)$$

and thus a relation like (39) cannot hold for these two channels. Note that, even in this example, the Stein exponent, i.e., the optimal exponential decay rate of the type II error such that the type I error still goes to zero, is still identical for the parallel and adaptive strategies, as it is infinite also for the optimal parallel strategy. This follows from

$$D_{\max}(\mathcal{E} \|\mathcal{F}) = \infty \Rightarrow D(\mathcal{E} \|\mathcal{F}) = D^{\text{reg}}(\mathcal{E} \|\mathcal{F}) = \infty. \quad (51)$$

A. Computability

Remark 6: The usefulness of finite-length bounds often crucially depends on whether the quantities involved can actually be efficiently computed [46], [47], [48]. To demonstrate this for our bound, we consider the following two applications:

- 1) We are given an adaptive strategy, where we know (potentially only a lower bound on)

$D_H^{\alpha_a}(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n))$, and want to employ the theorem to get a lower bound on the performance of the best-possible parallel strategy (i.e., we want to know how much worse could a parallel strategy potentially be).

Such a lower bound can be computed using Theorem 1 in time $\mathcal{O}(1)$ (i.e., the computational complexity is independent of the number of channel uses n and m), by using the bound on c_ℓ from (44). A potentially much tighter bound, obtained by finding ℓ explicitly and then calculating c_ℓ , can be computed in time $\mathcal{O}(n)$.

- 2) We want to upper bound the performance of the best possible adaptive strategy with n channel uses, by finding the best possible parallel strategy and then using Theorem 1. We show in Lemma 3 below that the best possible parallel strategy with m channel uses can actually be calculated in $\mathcal{O}(\text{poly}(m))$ time, and hence by choosing $m = n^{2+\xi}$ for $\xi > 0$ we also obtain asymptotically tight upper bounds on any adaptive strategy with n channel uses in $\mathcal{O}(\text{poly}(n))$ time.

To our knowledge this is the first such poly-time bound obtained in the literature. Specifically, computing the best parallel strategy and then using our bound is exponentially faster than any currently known way of optimizing over all adaptive strategies. While the authors of [45] showed that optimizing over all adaptive strategies can be phrased as an SDP, the size of this SDP grows exponentially in n , and there is no obvious symmetry (like the permutation invariance in the parallel case) that allows one to reduce the number of variables.

Lemma 3: Let $\mathcal{E}, \mathcal{F} : A \rightarrow B$ be two quantum channels such that $D_{\max}(\mathcal{E}\|\mathcal{F}) < \infty$. Then, for a fixed $\varepsilon \in [0, 1)$, the quantity

$$\frac{1}{n} D_H^\varepsilon(\mathcal{E}^{\otimes n}\|\mathcal{F}^{\otimes n}) = \frac{1}{n} \sup_{\nu \in \mathcal{D}(R^{\otimes n} \otimes A^{\otimes n})} D_H^\varepsilon(\mathcal{E}^{\otimes n}(\nu)\|\mathcal{F}^{\otimes n}(\nu)) \quad (52)$$

can be computed up to an additive error $\delta > 0$ in time $\mathcal{O}(\text{poly}(n) \log(\frac{1}{\delta}))$ as $n \rightarrow \infty$ and $\delta \rightarrow 0$.

Proof: As shown in [23, Prop. 2], the quantity $2^{-D_H^\varepsilon(\mathcal{E}^{\otimes n}\|\mathcal{F}^{\otimes n})}$ can be expressed as the following semidefinite program (SDP) over operators $\Omega_{R^n B^n} \in \mathcal{B}(\mathcal{H}_R^{\otimes n} \otimes \mathcal{H}_B^{\otimes n})$ and $\rho_{R^n} \in \mathcal{B}(\mathcal{H}_R^{\otimes n})$

$$\begin{aligned} & \underset{\Omega_{R^n B^n}, \rho_{R^n}}{\text{minimize}} && \text{Tr}(\Omega_{R^n B^n} \Gamma_{R^n B^n}^{\mathcal{F}^{\otimes n}}) \\ & \text{subject to} && \text{Tr}(\Omega_{R^n B^n} \Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}}) \geq 1 - \varepsilon \\ & && 0 \leq \Omega_{R^n B^n} \leq \rho_{R^n} \otimes \mathbb{1}_{B^n} \\ & && \text{Tr}(\rho_{R^n}) = 1, \end{aligned} \quad (53)$$

where $\Gamma_{RB}^\mathcal{E}$ is the Choi matrix of $\mathcal{E}$ and it is easy to see that $\Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}} = (\Gamma_{RB}^\mathcal{E})^{\otimes n}$. The operators $\Omega_{R^n B^n}$ and ρ_{R^n} in this SDP have a number of parameters exponential in n , but we will show that we can use the permutation symmetry of this problem to rephrase it as an SDP polynomial in n . Throughout this proof we use $\mathcal{H}$ to denote a general Hilbert space, which we will then choose to be either $\mathcal{H}_R, \mathcal{H}_B$,

or $\mathcal{H}_R \otimes \mathcal{H}_B$ in different situations. We will denote any objects that depend on the chosen Hilbert space with a subscript or superscript $\mathcal{H}$ (such as $P_\mathcal{H}$ in the following paragraph), where we then replace the subscript or superscript with R, B , or RB whenever we choose a specific Hilbert space; e.g., we write P_R, P_B , or P_{RB} .

For any permutation $\pi \in \mathfrak{S}_n$ (where $\mathfrak{S}_n$ is the symmetric group) we write $P_\mathcal{H}(\pi)$ for the permutation matrix corresponding to the action of π on $\mathcal{H}^{\otimes n}$ by permuting the n copies of $\mathcal{H}$. It is then easy to see that these permutation matrices are unitary and $P_\mathcal{H}(\pi)^\dagger = P_\mathcal{H}(\pi^{-1})$. For any operator $X \in \mathcal{B}(\mathcal{H}^{\otimes n})$ we also write the group average as

$$\bar{X} := \frac{1}{|\mathfrak{S}_n|} \sum_{\pi \in \mathfrak{S}_n} P_\mathcal{H}(\pi) X P_\mathcal{H}(\pi)^\dagger, \quad (54)$$

and the set of all permutation invariant operators as

$$\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n}) := \{A \in \mathcal{B}(\mathcal{H}^{\otimes n}) \mid P_\mathcal{H}(\pi) A P_\mathcal{H}(\pi)^\dagger = A, \forall \pi \in \mathfrak{S}_n\}. \quad (55)$$

We start by showing that the minimum in our SDP (53) is always achieved by permutation invariant operators. Let $(\Omega_{R^n B^n}, \rho_{R^n})$ be feasible for the optimization problem; i.e., they satisfy the constraints of (53). Let $\pi \in \mathfrak{S}_n$ be any permutation; then

$$\text{Tr}(P_R(\pi) \rho_{R^n} P_R(\pi)^\dagger) = \text{Tr}(\rho_{R^n}) \quad (56)$$

since the $P_R(\pi)$ are unitary, and thus also $\text{Tr}(\rho_{R^n}) = \text{Tr}(\overline{\rho_{R^n}})$. Similarly, we get

$$0 \leq P_{RB}(\pi) \Omega_{R^n B^n} P_{RB}(\pi)^\dagger \leq P_R(\pi) \rho_{R^n} P_R(\pi)^\dagger \otimes \mathbb{1}_{B^n} \quad (57)$$

since positivity is preserved under unitary conjugation, and we also used

$$P_{RB}(\pi) = (\mathbb{1}_{R^n} \otimes P_B(\pi))(P_R(\pi) \otimes \mathbb{1}_{B^n}) \quad (58)$$

which is immediate from the definition. This again implies that

$$0 \leq \overline{\Omega_{R^n B^n}} \leq \overline{\rho_{R^n}} \otimes \mathbb{1}_{B^n}. \quad (59)$$

By the cyclicity of the trace, and the fact that $\Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}} = (\Gamma_{RB}^\mathcal{E})^{\otimes n}$, we also see that

$$\text{Tr}(\overline{\Omega_{R^n B^n}} \Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}}) = \text{Tr}(\Omega_{R^n B^n} \Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}}), \quad (60)$$

$$\text{Tr}(\overline{\Omega_{R^n B^n}} \Gamma_{R^n B^n}^{\mathcal{F}^{\otimes n}}) = \text{Tr}(\Omega_{R^n B^n} \Gamma_{R^n B^n}^{\mathcal{F}^{\otimes n}}). \quad (61)$$

Hence, also $(\overline{\Omega_{R^n B^n}}, \overline{\rho_{R^n}})$ are feasible for the optimization problem (i.e., they satisfy the constraints) and also achieve the exact same value as $(\Omega_{R^n B^n}, \rho_{R^n})$ does. The group averages are elements of $\text{End}^{\mathfrak{S}_n}(R^n B^n)$ and $\text{End}^{\mathfrak{S}_n}(R^n)$ respectively, and hence we can restrict the optimization in (53) to such permutation invariant operators.

While the dimension of the subspace of these permutation invariant operators is polynomial in n , this does not yet show computability in $\text{poly}(n)$ time using standard SDP solvers, as the problem is not phrased using matrices of size $\text{poly}(n)$. In order to rephrase our problem in this way we follow the

approach of [49], where a similar statement has been shown for the sharp Rényi divergence $D_\alpha^\#(\mathcal{E}^{\otimes n} \parallel \mathcal{F}^{\otimes n})$. The key idea is to construct a suitable basis of the permutation invariant subspaces and then rephrase the SDP as one in which we minimize over (now only $O(\text{poly}(n))$ many) basis coefficients. As explained in [49, page 7352] (see also [50], [51]), one can construct such an orthogonal basis $C_r^\mathcal{H}$, $r \in \{1, \dots, m_\mathcal{H}\}$ of $\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ (where orthogonality is with respect to the Hilbert-Schmidt inner product, and $m_\mathcal{H} = \dim \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$) as follows: Let $\{|i\rangle\}_{i=1}^{d_\mathcal{H}}$ be an orthonormal basis of $\mathcal{H}$. Then, for any multi-index $\mathbf{i} \in \{1, \dots, d_\mathcal{H}\}^n$ define the associated vector $|\mathbf{i}\rangle = \bigotimes_{k=1}^n |i_k\rangle$ on the tensor-product system $\mathcal{H}^{\otimes n}$, and it is immediate that the set of all these vectors forms an orthonormal basis of $\mathcal{H}^{\otimes n}$. For any pair of such multi-indices $(\mathbf{i}, \mathbf{j})$ – this pair should be thought of as indexing a matrix element of an operator on $\mathcal{H}^{\otimes n}$ – we write the group orbit of these indices under the action of $\mathfrak{S}_n$ as

$$O(\mathbf{i}, \mathbf{j}) := \{ (\pi(\mathbf{i}), \pi(\mathbf{j})) \mid \pi \in \mathfrak{S}_n \}, \quad (62)$$

where $\pi(\mathbf{i})$ permutes the components of $\mathbf{i}$, i.e., $\pi(\mathbf{i})_k = \mathbf{i}_{\pi^{-1}(k)}$ for $k \in \{1, \dots, n\}$. There are exactly $m_\mathcal{H} = \dim \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ such group orbits, which we will label as $O_r^\mathcal{H}$, $r \in \{1, \dots, m_\mathcal{H}\}$, and a representative element of each such orbit (i.e., a pair $(\mathbf{i}, \mathbf{j}) \in O_r^\mathcal{H}$) can be efficiently computed given r , as shown in [49, page 7352]. This corresponds to the intuition that for $A \in \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ we require $A_{\mathbf{i}\mathbf{j}} = A_{\pi(\mathbf{i})\pi(\mathbf{j})}$ for any $\pi \in \mathfrak{S}_n$; hence the matrix elements of A have to be constant on each group orbit, and so the number of group orbits is equal to the dimension of $\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$. The basis elements $C_r^\mathcal{H} \in \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ are now defined by specifying their matrix elements as

$$(C_r^\mathcal{H})_{\mathbf{i}\mathbf{j}} := \begin{cases} 1 & \text{if } (\mathbf{i}, \mathbf{j}) \in O_r^\mathcal{H} \\ 0 & \text{otherwise.} \end{cases} \quad r = 1, \dots, m_\mathcal{H}. \quad (63)$$

It follows immediately from the definition that the $C_r^\mathcal{H}$ are orthogonal, and since we have $m_\mathcal{H} = \dim \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ of them, they form a basis.

As explained in [49, page 7353], for any matrix $A^{\otimes n} \in \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$, its coefficients with respect to the basis $C_r^\mathcal{H}$ can be computed straightforwardly from a description of A by picking a representative of each orbit, and these coefficients are hence computable in $\mathcal{O}(\text{poly}(n))$ time. Specifically, one can show that for any r and any pair of indices in the group orbit $(\mathbf{i}, \mathbf{j}) \in O_r^\mathcal{H}$, the corresponding basis coefficient is given by

$$\gamma_r := \prod_{k=1}^n A_{\mathbf{i}_k \mathbf{j}_k}; \quad (64)$$

i.e., we get

$$A^{\otimes n} = \sum_{r=1}^{m_\mathcal{H}} \gamma_r C_r^\mathcal{H}. \quad (65)$$

This can be applied to the Choi matrix $\Gamma_{R^n B^n}^{\mathcal{E}^{\otimes n}} = (\Gamma_{RB}^\mathcal{E})^{\otimes n}$, and hence we write

$$(\Gamma_{RB}^\mathcal{E})^{\otimes n} = \sum_{r=1}^{m_{RB}} \gamma_r^\mathcal{E} C_r^{RB} \quad (66)$$

and similarly when $\mathcal{E}$ is replaced by $\mathcal{F}$.

Additionally, for any finite-dimensional Hilbert space $\mathcal{H}$, there exists a $*$ -algebra isomorphism from $\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ to block-diagonal matrices [52, Theorem 1]

$$\phi_\mathcal{H} : \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n}) \rightarrow \bigoplus_{i=1}^{t_\mathcal{H}} \mathbb{C}^{m_i \times m_i}, \quad (67)$$

where

$$t_\mathcal{H} \leq (n+1)^{d_\mathcal{H}}, \quad (68)$$

$$d_\mathcal{H} = \dim(\mathcal{H}), \quad (69)$$

$$\sum_{i=1}^{t_\mathcal{H}} m_i^2 = \dim(\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})) \leq (n+1)^{d_\mathcal{H}^2}. \quad (70)$$

Introducing the notation $M_\mathcal{H} := \sum_{i=1}^{t_\mathcal{H}} m_i$, for any $A \in \text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$ we have that $\phi_\mathcal{H}(A) \in \mathbb{C}^{M_\mathcal{H} \times M_\mathcal{H}}$, and we write $[\phi_\mathcal{H}(A)]_i$ for the i -th block of $\phi_\mathcal{H}(A)$. Crucially, by [49, Lemma 3.3], [53, Prop. 2.4.4],

$$A \geq 0 \Leftrightarrow \phi_\mathcal{H}(A) \geq 0 \Leftrightarrow [\phi_\mathcal{H}(A)]_i \geq 0 \quad \forall i \in \{1, \dots, t_\mathcal{H}\}. \quad (71)$$

From [52, Theorem 1] it also follows that $\phi_\mathcal{H}$ preserves orthogonality; i.e., if $\text{Tr}(A^\dagger B) = 0$ then also $\text{Tr}(\phi_\mathcal{H}(A)^\dagger \phi_\mathcal{H}(B)) = 0$. Remember that $\{C_r^\mathcal{H}\}_{r \in \{1, \dots, m_\mathcal{H}\}}$ is a basis for $\text{End}^{\mathfrak{S}_n}(\mathcal{H}^{\otimes n})$, and hence we can expand $\Omega_{R^n B^n} \in \text{End}^{\mathfrak{S}_n}(R^n B^n)$ and $\rho_{R^n} \in \text{End}^{\mathfrak{S}_n}(R^n)$ as follows:

$$\Omega_{R^n B^n} = \sum_{r=1}^{m_{RB}} y_r C_r^{RB}, \quad (72)$$

$$\rho_{R^n} = \sum_{r=1}^{m_R} z_r C_r^R, \quad (73)$$

where $\{y_r \in \mathbb{C}\}_{r=1}^{m_{RB}}$ and $\{z_r \in \mathbb{C}\}_{r=1}^{m_R}$ are the respective basis coefficients. Note that since the C_r are not necessarily Hermitian, the coefficients are not necessarily real. Since optimizing over elements is obviously equivalent to optimizing over their basis coefficients, we can rephrase our SDP as follows:

$$\begin{aligned} & \underset{\{y_r\}_{r=1}^{m_{RB}}, \{z_r\}_{r=1}^{m_R}}{\text{minimize}} \quad \sum_{r=1}^{m_{RB}} y_r (\gamma_r^\mathcal{F})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}) \\ & \text{subject to} \\ & \sum_{r=1}^{m_{RB}} y_r (\gamma_r^\mathcal{E})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}) \geq 1 - \varepsilon \\ & \sum_{r=1}^{t_R} z_r \text{Tr}(C_r^R) = 1 \\ & 0 \leq \sum_{r=1}^{m_{RB}} y_r [\phi_{RB}(C_r^{RB})]_i \\ & \leq \sum_{r=1}^{m_R} z_r [\phi_{RB}(C_r^R \otimes I_{B^n})]_i \quad \forall i \in \{1, \dots, t_{RB}\} \end{aligned} \quad (74)$$

where we also used (71) and (66).

We will show that this is an SDP in $\text{poly}(n)$ variables with $\text{poly}(n)$ constraints by casting it into standard form.

To simplify notation we will write the Hilbert-Schmidt inner product using $\langle \cdot, \cdot \rangle$, i.e., $\langle A, B \rangle := \text{Tr}(A^\dagger B)$ and also write $M = M_{RB}$ (i.e., $M = M_{\mathcal{H}}$ as below (70) with $\mathcal{H} = \mathcal{H}_R \otimes \mathcal{H}_B$). For $s \in \mathbb{C}$, consider the following $(2M+1) \times (2M+1)$ block-diagonal matrix

$$X := \left[\sum_{r=1}^{m_{RB}} y_r \phi_{RB}(C_r^{RB}) \right] \oplus \left[\sum_{r'=1}^{m_R} z_{r'} \phi_{RB}(C_{r'}^R \otimes I_{B^n}) - \sum_{r=1}^{m_{RB}} y_r \phi_{RB}(C_r^{RB}) \right] \oplus s \quad (75)$$

where s is a slack variable that turns the one inequality constraint into an equality constraint (see further below). From here on, we write $(\cdot) \oplus (\cdot) \oplus (\cdot)$ to specify a block-diagonal matrix with block sizes equal to the ones in (75). The constraint $X \geq 0$ now implies

$$0 \leq \sum_{r=1}^{m_{RB}} y_r \llbracket \phi_{RB}(C_r^{RB}) \rrbracket_i \leq \sum_{r=1}^{m_R} z_r \llbracket \phi_{RB}(C_r^R \otimes I_{B^n}) \rrbracket_i \quad (76)$$

for $i = 1, \dots, t_{RB}$. Additionally, since ϕ_{RB} preserves orthogonality, we can recover the coefficients y_r and z_r from X by taking inner products with suitable operators. Specifically, with

$$\tilde{Y}_r := \phi_{RB}(C_r^{RB}) \oplus 0 \oplus 0, \quad (77)$$

$$Y_r := \frac{\tilde{Y}_r}{\langle \tilde{Y}_r, \tilde{Y}_r \rangle}, \quad (78)$$

for $r \in \{1, \dots, m_{RB}\}$, and with

$$\tilde{Z}_r := \phi_{RB}(C_r^R \otimes \mathbb{1}_{B^n}) \oplus \phi_{RB}(C_r^R \otimes \mathbb{1}_{B^n}) \oplus 0, \quad (79)$$

$$Z_r := \frac{\tilde{Z}_r}{\langle \tilde{Z}_r, \tilde{Z}_r \rangle}, \quad (80)$$

for $r \in \{1, \dots, m_R\}$, we have

$$y_r = \langle Y_r, X \rangle, \quad (81)$$

$$z_r = \langle Z_r, X \rangle. \quad (82)$$

Hence we can rephrase the expressions in (74) as inner products with X . Specifically,

$$\sum_{r=1}^{m_{RB}} y_r (\gamma_r^{\mathcal{E}})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}) = \left\langle \sum_{r=1}^{m_{RB}} Y_r (\gamma_r^{\mathcal{E}})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}), X \right\rangle \quad (83)$$

and one can do the same with $\mathcal{E}$ replaced by $\mathcal{F}$.

We want to transform our SDP into one where we optimize over X , and for that we need to impose the necessary block-diagonal structure of X (including the block-diagonal substructure that comes from its parts being images of ϕ_{RB}). For this, consider the following linear space

$$\left\{ [\phi_{RB}(\Omega)] \oplus [\phi_{RB}(\rho \otimes \mathbb{1}_{B^n}) - \phi_{RB}(\Omega)] \oplus s \mid \Omega \in \text{End}^{\mathfrak{S}_n}(R^n B^n), \rho \in \text{End}^{\mathfrak{S}_n}(R^n), s \in \mathbb{C} \right\} \quad (84)$$

and define $\mathcal{A}$ to be a set of matrices that form a basis of the orthogonal complement of this linear space, where we take the orthogonal complement within $\mathbb{C}^{(2M+1) \times (2M+1)}$. It is then easy to see that $|\mathcal{A}| \leq \dim(\mathbb{C}^{(2M+1) \times (2M+1)}) = (2M+1)^2 = \mathcal{O}(\text{poly}(n))$.

With $S := (0 \oplus 0 \oplus 1)$, we can then introduce the following SDP

$$\begin{aligned} & \text{minimize}_X \left\langle \sum_{r=1}^{m_{RB}} Y_r (\gamma_r^{\mathcal{F}})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}), X \right\rangle \\ & \text{subject to} \\ & \left\langle \left(\sum_{r=1}^{m_{RB}} Y_r (\gamma_r^{\mathcal{E}})^* \text{Tr}((C_r^{RB})^\dagger C_r^{RB}) \right) - S, X \right\rangle = 1 - \varepsilon \end{aligned} \quad (85)$$

$$X \in \mathbb{C}^{(2M+1) \times (2M+1)}$$

$$X \geq 0$$

$$\left\langle \sum_{r=1}^{t_R} Z_r \text{Tr}(C_r^R), X \right\rangle = 1$$

$$\langle A, X \rangle = 0 \quad \forall A \in \mathcal{A}$$

which is in standard form with $\mathcal{O}(\text{poly}(n))$ many constraints and matrices of size $\mathcal{O}(\text{poly}(n))$.

This new SDP is equivalent to (74), which can be seen as follows: From (75) and the calculations thereafter it follows that for every feasible $(\{y_r\}_{r=1}^{m_{RB}}, \{z_r\}_{r=1}^{m_R})$ in (74) there is a corresponding feasible X in (85) which achieves the same value. Conversely, every X that satisfies the constraints of (85) has to lie in the subspace (84), and it is then immediate that it can be represented as in (75) for suitable y_r and z_r . These y_r and z_r then also achieve the same value in (74) as X did in (85).

It now only remains to show that we can also efficiently calculate all that is required to parameterize this SDP. It was shown in [49, pages 7352-7353] (see also [54]) that the $\llbracket \phi_{\mathcal{H}}(C_r^{\mathcal{H}}) \rrbracket_i$ (and hence also the $\phi_{\mathcal{H}}(C_r^{\mathcal{H}})$) can be computed in $\text{poly}(n)$ time. As $I_{B^n} = I_B^{\otimes n}$, the coefficients of I_{B^n} with respect to the $\{C_r^B\}$ can be computed efficiently just as for the Choi matrices above, and additionally, from the construction of the $C_r^{\mathcal{H}}$, there is an obvious one-to-one mapping $C_r^R \otimes C_{r'}^B = C_{r''}^{RB}$; hence also $\phi_{RB}(C_r^R \otimes I_{B^n})$ is efficiently computable. This gives a generating set of the subspace (84) and a basis for its orthogonal complement is then also computable in $\mathcal{O}(\text{poly}(n))$ time. Finally, from the construction of the $C_r^{\mathcal{H}}$, it is immediate that the norm $\text{Tr}((C_r^{\mathcal{H}})^\dagger C_r^{\mathcal{H}})$ is given by the size of the group orbit, which is equal to the number of distinct permutations of an element $(\mathbf{i}, \mathbf{j})$ in the orbit. Concretely, if $\ell \in \{1, \dots, d_{\mathcal{H}}^2\}$ indexes all pairs of single-system indices (i, j) , $i, j \in \{1, \dots, d_{\mathcal{H}}\}$, then given $(\mathbf{i}, \mathbf{j}) \in O_r^{\mathcal{H}}$, let K_ℓ denote the number of occurrences of $\ell = (i, j)$ in the multi-index $(\mathbf{i}, \mathbf{j})$, i.e., the number of different k values ($k \in \{1, \dots, n\}$) for which $\mathbf{i}_k = i$, $\mathbf{j}_k = j$. The size of the group orbit is then given by

$$\text{Tr}((C_r^{\mathcal{H}})^\dagger C_r^{\mathcal{H}}) = |O_r^{\mathcal{H}}| = \binom{n}{K_1, \dots, K_{d_{\mathcal{H}}^2}} \quad (86)$$

and this multinomial coefficient can be calculated in time $\mathcal{O}(d_H^2)$; see, e.g., [55].

Also, $\text{Tr}(C_r^R)$ can be efficiently computed by picking a representative (i, j) of the orbit O_r^R and counting the number of different k where $i_k = j_k$. Thus, we can compute all the coefficients in our SDP in $\text{poly}(n)$ time. An SDP with $\text{poly}(n)$ variables and $\text{poly}(n)$ constraints can be solved up to an additive error δ' in time $\mathcal{O}(\text{poly}(n) \log(1/\delta'))$; see, e.g., [56]. The final step in our proof is to pick a suitable δ' so that we can solve our original problem (52) up to an additive error δ . For any n we write the solution for our SDP as

$$R_n := 2^{-D_H^\varepsilon(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n})}. \quad (87)$$

Let us start by showing that $-\log(R_n) = \mathcal{O}(n)$ as $n \rightarrow \infty$. By Lemma 2, we have for all $\gamma \in (0, 1 - \varepsilon)$ that

$$\begin{aligned} D_H^\varepsilon(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n}) &\leq D_{\max}^{\sqrt{1-\varepsilon-\gamma}}(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n}) + \log\left(\frac{4(\varepsilon + \gamma)}{\gamma^2}\right) \\ &\leq D_{\max}(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n}) + \log\left(\frac{4(\varepsilon + \gamma)}{\gamma^2}\right) \end{aligned} \quad (88)$$

$$\leq D_{\max}(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n}) + \log\left(\frac{4(\varepsilon + \gamma)}{\gamma^2}\right) \quad (89)$$

$$= nD_{\max}(\mathcal{E} \| \mathcal{F}) + \log\left(\frac{4(\varepsilon + \gamma)}{\gamma^2}\right) = \mathcal{O}(n). \quad (90)$$

Note that the case $\varepsilon = 0$ (while not directly covered by Lemma 2) also follows immediately by first using that $D_H^0(\mathcal{E} \| \mathcal{F}) \leq D_H^\varepsilon(\mathcal{E} \| \mathcal{F})$ for all $\varepsilon' \in (0, 1)$.

Now, for any given $\delta > 0$, let us pick $\delta' := (2^\delta - 1)R_n$. Remember that δ' was the additive error we make when solving the SDP (74), and this error then propagates to our original problem via

$$-\frac{1}{n} \log(R_n + \delta') = -\frac{1}{n} \log(R_n 2^\delta) \quad (91)$$

$$= -\frac{1}{n} \log(R_n) - \frac{\delta}{n} \quad (92)$$

$$= \frac{1}{n} D_H^\varepsilon(\mathcal{E}^{\otimes n} \| \mathcal{F}^{\otimes n}) - \frac{\delta}{n} \quad (93)$$

and hence this leads to an additive error of $\delta/n \leq \delta$ for the original problem.¹ It is easy to see that $\log(2^\delta - 1) = \log(\delta) + \mathcal{O}(1)$ as $\delta \rightarrow 0$, which implies $\log(1/\delta') = -\log(R_n) - \log(2^\delta - 1) = \mathcal{O}(n) + \mathcal{O}(\log(1/\delta)) = \mathcal{O}(n \log(1/\delta))$ and hence our original problem can be calculated up to an error δ in time $\mathcal{O}(\text{poly}(n) \log(1/\delta')) = \mathcal{O}(\text{poly}(n) \log(1/\delta))$. $\square$

B. A Simple One-Shot Version of the Chain Rule

The remainder of this section proves Theorem 1. The general idea of the proof is the following: We will start by moving from the hypothesis testing relative entropy of the adaptive strategy to the Umegaki relative entropy, using Lemma 1. Then, we will see that by using an amortization argument (equations (124)–(127) below), we can bound the

¹This last inequality might seem very far from optimal, and if one wants to make further assumptions on how $n\delta$ behaves as $n \rightarrow \infty$, $\delta \rightarrow 0$ our runtime bounds can indeed be improved; however, this is not required for what we desire to show. Also, the signs here are chosen as the errors appear in practice: the SDP (74) is a minimization, so any numerical approximation will be larger than the true value and the additive error hence positive, which then translates to a value smaller than the true value for our original problem (52).

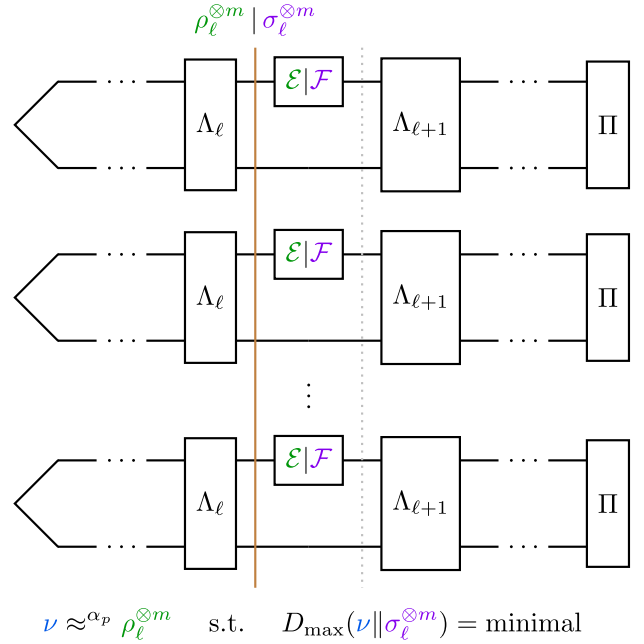


Fig. 2. Illustration of a key step in our proof, the construction of the parallel input state. We start by picking a single step $\ell \in \{1, \dots, n\}$ out of the adaptive protocol, where the distinguishability increase $D(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) - D(\rho_\ell \| \sigma_\ell)$ is maximal (this corresponds to the step from the orange to the dotted grey line in the diagram). Now consider m copies of the adaptive strategy in parallel. We construct our parallel input state ν starting from m copies of the input state of the adaptive strategy at this step ℓ if the channel was $\mathcal{E}$ (this is $\rho_\ell^{\otimes m}$). The state ν is then smoothed a bit to reduce its distance to $\sigma_\ell^{\otimes m}$ (which is the input state that we would have if the channel was $\mathcal{F}$). The degree to which we smooth depends on the type I error α_p we want to achieve with the parallel strategy. Having a small type I error means that the state ν is very close to $\rho_\ell^{\otimes m}$, whereas allowing for a larger type I error will move the state closer to $\sigma_\ell^{\otimes m}$.

performance of the adaptive strategy by the performance of just one of its steps, the step where the distinguishability of the two states (that occur if the channel is either $\mathcal{E}$ or $\mathcal{F}$) increases the most. We then consider m parallel copies of this step and construct a parallel input state using a chain rule for the smoothed max-relative entropy (Lemma 4); see Figure 2 below for an illustration of this step. The smoothed max-relative entropies can be related to the Umegaki relative entropies we used in the amortization argument by the non-asymptotic bounds presented in Lemma 6. Finally we connect to the hypothesis testing relative entropy of a parallel protocol using Lemma 2.

This subsection and the following one present lemmas we will use as part of our proof of Theorem 1.

The following is a variant of [22, Prop. 3.2], where our different convention of smoothing the max-divergence (smoothing only over normalized states) leads to a tighter and simpler bound, and restricting to the case of a single input system also makes things a bit simpler.

Lemma 4: Let $\mathcal{E}$ and $\mathcal{F}$ be arbitrary quantum channels from system A to system B , and let $\rho \in \mathcal{D}(A)$, $\sigma \in \mathcal{P}(A)$. Then for all $\varepsilon, \varepsilon' \in [0, 1]$ there exists a state $\nu \in B_\varepsilon^\circ(\rho)$ such that

$$D_{\max}^{\varepsilon+\varepsilon'}(\mathcal{E}(\rho) \| \mathcal{F}(\sigma)) \leq D_{\max}^\varepsilon(\rho \| \sigma) + D_{\max}^{\varepsilon'}(\mathcal{E}(\nu) \| \mathcal{F}(\nu)). \quad (94)$$

Moreover, ν can be chosen as

$$\nu = \arg \min_{\tilde{\rho} \in B_{\varepsilon}^{\circ}(\rho)} D_{\max}(\tilde{\rho} \parallel \sigma). \quad (95)$$

Proof: Let $\nu \in B_{\varepsilon}^{\circ}(\rho)$ be an optimal choice for $D_{\max}^{\varepsilon}(\rho \parallel \sigma)$; i.e.,

$$\nu \leq 2^{D_{\max}^{\varepsilon}(\rho \parallel \sigma)} \sigma. \quad (96)$$

Since $\mathcal{F}$ is a positive map, this implies that

$$\mathcal{F}(\nu) \leq 2^{D_{\max}^{\varepsilon}(\rho \parallel \sigma)} \mathcal{F}(\sigma). \quad (97)$$

Furthermore, let $\tau \in B_{\varepsilon'}^{\circ}(\mathcal{E}(\nu))$ be an optimal choice for $D_{\max}^{\varepsilon'}(\mathcal{E}(\nu) \parallel \mathcal{F}(\nu))$, so that

$$\tau \leq 2^{D_{\max}^{\varepsilon'}(\mathcal{E}(\nu) \parallel \mathcal{F}(\nu))} \mathcal{F}(\nu). \quad (98)$$

Combining the last two inequalities leads to

$$\tau \leq 2^{D_{\max}^{\varepsilon'}(\mathcal{E}(\nu) \parallel \mathcal{F}(\nu)) + D_{\max}^{\varepsilon}(\rho \parallel \sigma)} \mathcal{F}(\sigma). \quad (99)$$

It remains to show that $P(\tau, \mathcal{E}(\rho)) \leq \varepsilon + \varepsilon'$ (where P is the sine distance). This follows from

$$P(\tau, \mathcal{E}(\rho)) \leq P(\tau, \mathcal{E}(\nu)) + P(\mathcal{E}(\nu), \mathcal{E}(\rho)) \quad (100)$$

$$\leq \varepsilon' + P(\nu, \rho) \quad (101)$$

$$\leq \varepsilon' + \varepsilon, \quad (102)$$

where we used the triangle inequality and the data-processing inequality for the sine distance (see, e.g., [26]). $\square$

C. Non-Asymptotic Bounds for the Smoothed Max-Relative Entropy

The asymptotic equipartition property for the smoothed max-relative entropy states that [57]

$$\lim_{n \rightarrow \infty} \frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) = D(\rho \parallel \sigma) \quad \forall \varepsilon \in (0, 1). \quad (103)$$

For our proof of Theorem 1, we will require a bound on the speed of this convergence. Results in this direction have appeared before in the literature; however, none of the previous results turn out to be directly applicable for our purposes (see Remark 8 and Remark 9 below). To prove our result we follow the established path of using quantum Rényi divergences. One key ingredient is a bound on the distance between the Petz–Rényi divergence and relative entropy, where the following is a slight refinement of [58, Lemma 6.3]:

Lemma 5: Let $\rho, \sigma \in \mathcal{D}(\mathcal{H})$ be quantum states. For $\gamma \in (0, 1]$, define

$$c_{\gamma}(\rho \parallel \sigma) := \frac{1}{\gamma} \log \left(2^{\gamma D_{1+\gamma}(\rho \parallel \sigma)} + 2^{-\gamma D_{1-\gamma}(\rho \parallel \sigma)} + 1 \right). \quad (104)$$

Then, for all $\gamma \in (0, 1]$ and $\delta \in (0, \frac{\gamma}{2}]$:

$$D_{1+\delta}(\rho \parallel \sigma) \leq D(\rho \parallel \sigma) + \ln(2) \delta (c_{\gamma}(\rho \parallel \sigma))^2 \quad (105)$$

$$\leq D(\rho \parallel \sigma) + \delta (c_{\gamma}(\rho \parallel \sigma))^2. \quad (106)$$

Furthermore, if $D(\rho \parallel \sigma) < \infty$, then for all $\gamma \in (0, 1]$ and $\delta \in (0, \frac{\gamma}{2}]$

$$D_{1-\delta}(\rho \parallel \sigma) \geq D(\rho \parallel \sigma) - \ln(2) \delta (c_{\gamma}(\rho \parallel \sigma))^2 \cosh(\ln(2) \delta c_{\gamma}(\rho \parallel \sigma)). \quad (107)$$

and for all $\delta \in (0, \frac{\log 3}{2c_{\gamma}(\rho \parallel \sigma)}]$:

$$D_{1-\delta}(\rho \parallel \sigma) \geq D(\rho \parallel \sigma) - \ln(2) \cosh(\log(3)/2) \delta (c_{\gamma}(\rho \parallel \sigma))^2 \quad (108)$$

$$\geq D(\rho \parallel \sigma) - \delta (c_{\gamma}(\rho \parallel \sigma))^2. \quad (109)$$

The proof of this lemma appears in Appendix A.

Remark 7: The previously known bound [58, Lemma 6.3] states only an analogue of (106) and follows from Lemma 5 after setting $\gamma = 1/2$. Note that in its analogue of (106), [58, Lemma 6.3] also has a stronger constraint on the range of δ , which turns out not to be necessary.

Using Lemma 5 we can then establish a bound on the convergence speed in the asymptotic equipartition property:

Lemma 6: Let $\rho, \sigma \in \mathcal{D}(\mathcal{H})$ be quantum states, and for $\gamma \in (0, 1]$, take $c_{\gamma}(\rho \parallel \sigma)$ as in (104). Then, for all $\varepsilon \in (0, 1)$ and $n \in \mathbb{N}$:

$$\begin{aligned} \frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) &\geq D(\rho \parallel \sigma) \\ &\quad - \frac{c_{\gamma}(\rho \parallel \sigma)}{\sqrt{n}} \left[\frac{\ln(2) \log(3)}{2} \cosh\left(\frac{\log(3)}{2}\right) \right. \\ &\quad \left. + \frac{4}{\log(3)} \log\left(\frac{1}{1-\varepsilon}\right) \right], \end{aligned} \quad (110)$$

$$\begin{aligned} \frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) &\leq D(\rho \parallel \sigma) \\ &\quad + \frac{c_{\gamma}(\rho \parallel \sigma)}{\sqrt{n}} \left[\frac{\ln(2) \log(3)}{2} + \frac{4}{\log(3)} \log\left(\frac{1}{\varepsilon}\right) \right] \\ &\quad + \frac{1}{n} \log\left(\frac{1}{1-\varepsilon^2}\right), \end{aligned} \quad (111)$$

which implies

$$\frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) \geq D(\rho \parallel \sigma) - \frac{4c_{\gamma}(\rho \parallel \sigma)}{\sqrt{n}} \log\left(\frac{2}{1-\varepsilon}\right), \quad (112)$$

$$\begin{aligned} \frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) &\leq D(\rho \parallel \sigma) + \frac{4c_{\gamma}(\rho \parallel \sigma)}{\sqrt{n}} \log\left(\frac{2}{\varepsilon}\right) \\ &\quad + \frac{1}{n} \log\left(\frac{1}{1-\varepsilon^2}\right). \end{aligned} \quad (113)$$

These bounds can be tightened by adding a condition on n to be large enough. Define first the numerical constants

$$K_1 := 2\sqrt{2 \ln(2) \cosh\left(\frac{\log(3)}{2}\right)} \leq 2.72, \quad (114)$$

$$K_2 := 2\sqrt{2 \ln(2)} \leq 2.36. \quad (115)$$

Then, if

$$n \geq \log\left(\frac{1}{1-\varepsilon}\right) \left(\frac{8}{\log(3) K_1}\right)^2, \quad (116)$$

we have the stronger bound

$$\begin{aligned} \frac{1}{n} D_{\max}^{\varepsilon}(\rho^{\otimes n} \parallel \sigma^{\otimes n}) &\geq D(\rho \parallel \sigma) - \frac{K_1}{\sqrt{n}} c_{\gamma}(\rho \parallel \sigma) \sqrt{\log\left(\frac{1}{1-\varepsilon}\right)}, \\ &\quad (117) \end{aligned}$$

and similarly, if

$$n \geq \log\left(\frac{1}{\varepsilon}\right) \left(\frac{8}{\gamma c_\gamma(\rho\|\sigma) K_2}\right)^2, \quad (118)$$

it holds that

$$\begin{aligned} \frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) &\leq D(\rho\|\sigma) + \frac{K_2}{\sqrt{n}} c_\gamma(\rho\|\sigma) \sqrt{\log\left(\frac{1}{\varepsilon}\right)} \\ &\quad + \frac{1}{n} \log\left(\frac{1}{1-\varepsilon^2}\right). \end{aligned} \quad (119)$$

Remember that $\gamma c_\gamma(\rho\|\sigma) \geq \log(3)$, and hence (118) is also always satisfied if

$$n \geq \log\left(\frac{1}{\varepsilon}\right) \left(\frac{8}{\log(3) K_2}\right)^2. \quad (120)$$

The proof of this lemma appears in Appendix B.

Remark 8: Our Lemma 6 can be seen an extension of [58, Theorem 6.4], where a similar bound to (119) is shown (however with a worse constant and different smoothing convention). Note that [58] uses the notation $S(\rho\|\sigma) := -D(\rho\|\sigma)$ and $S_{\min}^\varepsilon(\rho\|\sigma) := -D_{\max}^\varepsilon(\rho\|\sigma)$, and hence [58, Theorem 6.4], while looking like a lower bound, actually is an upper bound on $D_{\max}^\varepsilon$. An equivalent of (117) is shown in [58] only for the smoothed conditional min-entropy and not for the (more general) smoothed max-relative entropy.

Remark 9: Another already existing bound for the AEP convergence is the so-called second-order expansion [59], which gives a tight asymptotic characterization also of the second-order $\sqrt{n}$ term in the convergence to the relative entropy. There, the second-order term is shown to be proportional to the square root of the relative entropy variance

$$V(\rho\|\sigma) := \text{Tr}(\rho(\log(\rho) - \log(\sigma) - D(\rho\|\sigma))^2). \quad (121)$$

Later in the proof of Theorem 1, we want to apply these convergence bounds to the case in which $\rho = \rho_n$ and $\sigma = \sigma_n$ are the states in an adaptive protocol, and we would like to obtain a bound on the convergence parameter in n . We will see that by using chain rules for the geometric relative Rényi entropy (or the max-relative entropy) we will be able to show that $c_\gamma(\rho_n\|\sigma_n) = \mathcal{O}(n)$, while we are not aware of any way of obtaining such a bound for $V(\rho_n\|\sigma_n)$, and hence cannot directly use second-order asymptotics.

D. Proof of Our Main Result (Theorem 1)

Proof: We start by applying Lemma 1 to $D_H^{\alpha_a}(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n))$:

$$\begin{aligned} \frac{1}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) \\ \leq \frac{1}{n} \frac{1}{1-\alpha_a} (D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) + h(\alpha_a)). \end{aligned} \quad (122)$$

Note that a classical version of this equation in the context of channel discrimination was previously obtained in [15, Eq. (33)]. Note now that we can write

$$D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) \quad (123)$$

$$= D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) - D(\rho_n\|\sigma_n) + D(\rho_n\|\sigma_n) \quad (124)$$

$$\begin{aligned} &= D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) - D(\rho_n\|\sigma_n) \\ &\quad + D(\Lambda_n(\mathcal{E}(\rho_{n-1}))\|\Lambda_n(\mathcal{F}(\sigma_{n-1}))) \end{aligned} \quad (125)$$

$$\leq D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) - D(\rho_n\|\sigma_n) + D(\mathcal{E}(\rho_{n-1})\|\mathcal{F}(\sigma_{n-1})) \quad (126)$$

$$\leq \dots \leq \sum_{k=1}^n [D(\mathcal{E}(\rho_k)\|\mathcal{F}(\sigma_k)) - D(\rho_k\|\sigma_k)], \quad (127)$$

where we used the definition of ρ_k and σ_k , the data-processing inequality, and the fact that $\rho_1 = \sigma_1$. Let us use the index ℓ for the step in the adaptive protocol where this amortized difference is the largest, i.e.

$$\ell := \arg \max_{k \in \{1, \dots, n\}} [D(\mathcal{E}(\rho_k)\|\mathcal{F}(\sigma_k)) - D(\rho_k\|\sigma_k)]. \quad (128)$$

Then,

$$\frac{1}{n} D(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) \leq D(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) - D(\rho_\ell\|\sigma_\ell). \quad (129)$$

We can convert this to smoothed max-relative entropies by using Lemma 6. We will proceed with the bound requiring a condition on m (or n as it is called in Lemma 6); the m -independent bound is achieved in complete analogy by just taking the alternative statements from Lemma 6. We get:

$$\begin{aligned} D(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) - D(\rho_\ell\|\sigma_\ell) &\leq \\ \frac{1}{m} \left(D_{\max}^{\varepsilon_1}(\mathcal{E}^{\otimes m}(\rho_\ell^{\otimes m})\|\mathcal{F}^{\otimes m}(\sigma_\ell^{\otimes m})) - D_{\max}^{\varepsilon_2}(\rho_\ell^{\otimes m}\|\sigma_\ell^{\otimes m}) \right) \\ &\quad + \frac{1}{\sqrt{m}} \left[K_1 c_{\gamma_1}(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) \sqrt{\log\left(\frac{1}{1-\varepsilon_1}\right)} \right. \\ &\quad \left. + K_2 c_{\gamma_2}(\rho_\ell\|\sigma_\ell) \sqrt{\log\left(\frac{1}{1-\varepsilon_2}\right)} \right] + \frac{1}{m} \log\left(\frac{1}{1-\varepsilon_2^2}\right) \end{aligned} \quad (130)$$

where $\gamma_1, \gamma_2 \in (0, 1]$ and $\varepsilon_1, \varepsilon_2 \in (0, 1)$ are arbitrary. It will be very convenient to choose $1-\varepsilon_1 = \varepsilon_2 =: \varepsilon$, which is almost optimal as $K_1 \approx K_2$ and $c_\gamma(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) \approx c_\gamma(\rho_\ell\|\sigma_\ell)$ for large ℓ (which is the regime we are most interested in). Then,

$$\begin{aligned} D(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) - D(\rho_\ell\|\sigma_\ell) &\leq \\ \frac{1}{m} \left(D_{\max}^{1-\varepsilon}(\mathcal{E}^{\otimes m}(\rho_\ell^{\otimes m})\|\mathcal{F}^{\otimes m}(\sigma_\ell^{\otimes m})) - D_{\max}^\varepsilon(\rho_\ell^{\otimes m}\|\sigma_\ell^{\otimes m}) \right) \\ &\quad + \frac{c_\ell}{\sqrt{m}} \sqrt{\log\left(\frac{1}{\varepsilon}\right)} + \frac{1}{m} \log\left(\frac{1}{1-\varepsilon^2}\right), \end{aligned} \quad (131)$$

where

$$c_\ell := \inf_{\gamma_1, \gamma_2 \in (0, 1]} [K_1 c_{\gamma_1}(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) + K_2 c_{\gamma_2}(\rho_\ell\|\sigma_\ell)], \quad (132)$$

and the conditions on m from Lemma 6 will be satisfied if

$$m \geq \log\left(\frac{1}{\varepsilon}\right) \left(\frac{8}{\log(3) K_2}\right)^2. \quad (133)$$

Taking $\varepsilon \leq 1/2$, we can apply Lemma 4 to (131). We then get a state $\tilde{\nu} \in B_\varepsilon^\circ(\rho_\ell^{\otimes m})$ such that

$$D(\mathcal{E}(\rho_\ell)\|\mathcal{F}(\sigma_\ell)) - D(\rho_\ell\|\sigma_\ell)$$

$$\leq \frac{1}{m} D_{\max}^{1-2\varepsilon}(\mathcal{E}^{\otimes m}(\tilde{\nu}) \| \mathcal{F}^{\otimes m}(\tilde{\nu})) + \frac{c_\ell}{\sqrt{m}} \sqrt{\log\left(\frac{1}{\varepsilon}\right)} + \frac{1}{m} \log\left(\frac{1}{1-\varepsilon^2}\right). \quad (134)$$

Note that $\tilde{\nu} = \tilde{\nu}_{R_a^m A^m} \in \mathcal{D}(R_a^m \otimes A^m)$ is the first of the two possible choices for ν given in Theorem 1. For the other one, let $\tilde{\nu}_{R' R_a^m A^m}$ be a purification of $\tilde{\nu}_{R_a^m A^m}$, and hence also a purification of $\tilde{\nu}_{A^m}$. As the smoothed max-divergence satisfies the data-processing inequality (see, e.g., [26, Prop. 4.60]) we have

$$D_{\max}^{1-2\varepsilon}(\mathcal{E}^{\otimes m}(\tilde{\nu}_{R_a^m A^m}) \| \mathcal{F}^{\otimes m}(\tilde{\nu}_{R_a^m A^m})) \leq D_{\max}^{1-2\varepsilon}(\mathcal{E}^{\otimes m}(\tilde{\nu}_{R' R_a^m A^m}) \| \mathcal{F}^{\otimes m}(\tilde{\nu}_{R' R_a^m A^m})). \quad (135)$$

Now, let ν be the canonical purification of $\tilde{\nu}_{A^m}$, as specified in Theorem 1. Since all purifications are equivalent up to an isometry on the purifying system (see, e.g., [26]) on which the channels $\mathcal{E}^{\otimes m}$ and $\mathcal{F}^{\otimes m}$ act as an identity, and since the smoothed max-divergence is invariant under isometries, we have

$$D_{\max}^{1-2\varepsilon}(\mathcal{E}^{\otimes m}(\tilde{\nu}_{R' R_a^m A^m}) \| \mathcal{F}^{\otimes m}(\tilde{\nu}_{R' R_a^m A^m})) = D_{\max}^{1-2\varepsilon}(\mathcal{E}^{\otimes m}(\nu) \| \mathcal{F}^{\otimes m}(\nu)). \quad (136)$$

We will proceed with ν as the chosen state, but note that everything works analogously by choosing $\tilde{\nu}$. We now further apply the upper bound from Lemma 2 to find:

$$\begin{aligned} D(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) - D(\rho_\ell \| \sigma_\ell) &\leq \frac{1}{m} D_H^{1-(1-2\varepsilon)^2}(\mathcal{E}^{\otimes m}(\nu) \| \mathcal{F}^{\otimes m}(\nu)) + \frac{c_\ell}{\sqrt{m}} \sqrt{\log\left(\frac{1}{\varepsilon}\right)} \\ &\quad + \frac{1}{m} \left[\log\left(\frac{1}{1-\varepsilon^2}\right) + \log\left(\frac{1}{1-(1-2\varepsilon)^2}\right) \right]. \end{aligned} \quad (137)$$

To get our desired expression we set $\alpha_p := 1 - (1 - 2\varepsilon)^2$, or equivalently $\varepsilon = \frac{1}{2}(1 - \sqrt{1 - \alpha_p})$. Using the estimates

$$\frac{1 - \sqrt{1 - \alpha_p}}{2} \geq \frac{\alpha_p}{4} \quad (138)$$

and

$$1 - \varepsilon^2 = \frac{1 + \sqrt{1 - \alpha_p}}{2} + \frac{\alpha_p}{4} \geq 1 - \frac{\alpha_p}{4}, \quad (139)$$

we arrive at the expression

$$\begin{aligned} D(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) - D(\rho_\ell \| \sigma_\ell) &\leq \frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu) \| \mathcal{F}^{\otimes m}(\nu)) \\ &\quad + \frac{c_\ell}{\sqrt{m}} \sqrt{\log\left(\frac{4}{\alpha_p}\right)} + \frac{1}{m} \left[\log\left(\frac{1}{\alpha_p}\right) - \log\left(1 - \frac{\alpha_p}{4}\right) \right], \end{aligned} \quad (140)$$

which we can insert into (129) and then (122) to get

$$\begin{aligned} &\frac{1}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n) \| \mathcal{F}(\sigma_n)) \\ &\leq \frac{1}{1 - \alpha_a} \left[\frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu) \| \mathcal{F}^{\otimes m}(\nu)) + \frac{c_\ell}{\sqrt{m}} \sqrt{\log\left(\frac{4}{\alpha_p}\right)} \right] \end{aligned}$$

$$+ \frac{1}{m} \left[\log\left(\frac{1}{\alpha_p}\right) - \log\left(1 - \frac{\alpha_p}{4}\right) \right] + \frac{h(\alpha_a)}{n}, \quad (141)$$

which leads to the desired statement in (39).

For the bounds on c_ℓ , we start with

$$\begin{aligned} c_\gamma(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) &= \frac{1}{\gamma} \log\left(2^{\gamma D_{1+\gamma}(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell))} + 2^{-\gamma D_{1-\gamma}(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell))} + 1\right) \\ &\leq \frac{1}{\gamma} \log\left(2^{\gamma \hat{D}_{1+\gamma}(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell))} + 2\right), \end{aligned} \quad (142)$$

where we used (12) for the inequality. Now, by repeated use of the chain rule for the geometric Rényi divergence [60, Thm 3.4] we get

$$\begin{aligned} \hat{D}_{1+\gamma}(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) &\leq \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F}) + \hat{D}_{1+\gamma}(\rho_\ell \| \sigma_\ell) \end{aligned} \quad (143)$$

$$= \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F}) + \hat{D}_{1+\gamma}(\Lambda_\ell(\mathcal{E}(\rho_{\ell-1})) \| \Lambda_\ell(\mathcal{F}(\sigma_{\ell-1}))) \quad (144)$$

$$\leq \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F}) + \hat{D}_{1+\gamma}(\mathcal{E}(\rho_{\ell-1}) \| \mathcal{F}(\sigma_{\ell-1})) \quad (145)$$

$$\leq \dots \leq \ell \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F}). \quad (146)$$

Defining the corresponding channel quantity

$$\hat{c}_\gamma(\mathcal{E} \| \mathcal{F}) := \frac{1}{\gamma} \left(2^{\gamma \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F})} + 2 \right), \quad (147)$$

we find that

$$c_\gamma(\mathcal{E}(\rho_\ell) \| \mathcal{F}(\sigma_\ell)) \leq \frac{1}{\gamma} \log\left(2^{\ell \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F})} + 2\right) \quad (148)$$

$$\leq \frac{\ell}{\gamma} \log\left(2^{\hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F})} + 2\right) \quad (149)$$

$$= \ell \hat{c}_\gamma(\mathcal{E} \| \mathcal{F}). \quad (150)$$

Using the same argument, we find that also $\hat{D}_{1+\gamma}(\rho_\ell \| \sigma_\ell) \leq \ell \hat{D}_{1+\gamma}(\mathcal{E} \| \mathcal{F})$ and hence also

$$c_\gamma(\rho_\ell \| \sigma_\ell) \leq \ell \hat{c}_\gamma(\mathcal{E} \| \mathcal{F}). \quad (151)$$

Thus,

$$c_\ell \leq \ell(K_1 + K_2) \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \| \mathcal{F}) \quad (152)$$

$$\leq n(K_1 + K_2) \inf_{\gamma \in (0,1]} \hat{c}_\gamma(\mathcal{E} \| \mathcal{F}), \quad (153)$$

concluding the proof. $\square$

1) *Proof of Corollary 1:*

Proof: Corollary 1 follows from (31) by making the following estimates

$$h(\alpha_a) \leq 1, \quad (154)$$

$$K \leq 1, \quad (155)$$

$$\begin{aligned} c'_\ell &\leq \frac{8n}{\log(3)} \hat{c}_1(\mathcal{E} \| \mathcal{F}) \approx 5.05 n \hat{c}_1(\mathcal{E} \| \mathcal{F}) \\ &\leq 6 n \hat{c}_1(\mathcal{E} \| \mathcal{F}), \end{aligned} \quad (156)$$

$$-\log\left(1 - \frac{\alpha_p}{4}\right) \leq 1. \quad (157)$$

Combining all of these we can bound the error term by

$$\frac{6n}{\sqrt{m}} \hat{c}_1(\mathcal{E} \| \mathcal{F}) \log\left(\frac{8}{\alpha_p}\right) + \frac{1}{m} \log\left(\frac{2}{\alpha_p}\right) + \frac{1}{n}$$

$$\leq \frac{7n}{\sqrt{m}} \hat{c}_1(\mathcal{E} \parallel \mathcal{F}) \log\left(\frac{8}{\alpha_p}\right) + \frac{1}{n}, \quad (158)$$

where we used that $\hat{c}_\gamma(\mathcal{E} \parallel \mathcal{F}) \geq 1$ for all $\gamma \in (0, 1]$. Note finally that $\hat{D}_2(\mathcal{E} \parallel \mathcal{F}) = D_2(\mathcal{E} \parallel \mathcal{F}) \leq D_{\max}(\mathcal{E} \parallel \mathcal{F})$, which leads to the desired expression of C in (28). $\square$

IV. AN EXAMPLE

In this section we provide an example that illustrates how adaptive and parallel strategies can differ in the finite-length regime, and how our result bounds this difference. Specifically, we construct an example where the number of channel uses required by a parallel strategy to match a specific adaptive strategy turns out to be arbitrarily large. This demonstrates that the relation between adaptive and parallel strategies in the finite-length regime is in general complex, and one cannot expect a substantially simpler relationship between adaptive and parallel strategies than what we obtain in Corollary 1 and Theorem 1.

Our example is inspired by the already mentioned example from [16] and [21] that shows a separation between the asymptotic error decay rate of adaptive and parallel strategies in the *symmetric* setting. For a specific pair of channels $(\mathcal{E}, \mathcal{F})$, they construct an adaptive strategy with two channel uses that achieves perfect discrimination (i.e., type I and type II error are equal to zero), whereas for parallel strategies they upper bound the symmetric error exponent by a finite value (i.e., they upper bound the rate at which both errors can simultaneously go to zero); so in particular, perfect discrimination is never possible with parallel strategies. Looking at this example in the asymmetric setting, one finds that there exists an input state $\nu \in \mathcal{D}(A)$ such that for any arbitrary type I error α_p there exists an m such that $D_H^{\alpha_p}((\mathcal{E}(\nu))^{\otimes m} \parallel (\mathcal{F}(\nu))^{\otimes m}) = \infty$; i.e., already with a parallel strategy with product inputs one can achieve zero type II error with an arbitrary small type I error if one only makes m large enough. Hence, unlike the symmetric setting, in the asymmetric setting there is no asymptotic gap between adaptive and parallel strategies, but there is still a significant difference for finite m , as the adaptive strategy achieves $\alpha_a = 0, \beta_a = 0$ with only two channel uses, whereas the parallel strategy requires a large m to achieve $\alpha_p \leq \varepsilon, \beta_p = 0$. As mentioned in Remark 5, the two channels $\mathcal{E}$ and $\mathcal{F}$ used in this example have $D_{\max}(\mathcal{E} \parallel \mathcal{F}) = \infty$ and hence Theorem 1 does not immediately apply. What we are going to do though, is use a slightly noisy version of this channel $\mathcal{F}$, which makes $D_{\max}(\mathcal{E} \parallel \mathcal{F})$ finite.

Define the channels $\mathcal{E}, \mathcal{F} : \mathcal{D}(\mathbb{C}^2 \otimes \mathbb{C}^2) \rightarrow \mathcal{D}(\mathbb{C}^2)$ for $\kappa \in [0, 1]$ as follows:

$$\begin{aligned} \mathcal{E}(\rho \otimes \omega) &:= |0\rangle\langle 0| \langle 0|\omega|0\rangle + |0\rangle\langle 0| \langle 0|\rho|0\rangle \langle 1|\omega|1\rangle \\ &\quad + \frac{1}{2} \langle 11|\rho \otimes \omega|11\rangle \end{aligned} \quad (159)$$

$$\begin{aligned} \mathcal{F}(\rho \otimes \omega) &:= (1 - \kappa) \left[|+\rangle\langle +| \langle 0|\omega|0\rangle \right. \\ &\quad + |1\rangle\langle 1| \langle +|\rho|+\rangle \langle 1|\omega|1\rangle \\ &\quad \left. + \frac{1}{2} \langle -1|\rho \otimes \omega|-1\rangle \right] + \kappa \frac{\mathbb{1}}{2}. \end{aligned} \quad (160)$$

It is easy to see that

$$\mathcal{E}(|0\rangle\langle 0| \otimes |1\rangle\langle 1|) = |0\rangle\langle 0|, \quad (161)$$

$$\mathcal{F}(|0\rangle\langle 0| \otimes |1\rangle\langle 1|) = (1 - \delta(\kappa)) |1\rangle\langle 1| + \delta(\kappa) |0\rangle\langle 0|. \quad (162)$$

where $\delta(\kappa) = (3\kappa - \kappa^2)/4$. For $\kappa = 0$ (which corresponds to the original example in [16]) we find $\delta(0) = 0$ and hence the above gives an adaptive strategy that makes the channels perfectly distinguishable with just two channel uses. The exact same strategy will become arbitrarily good if κ is nonzero but small, specifically

$$\frac{1}{2} D_H^0(\mathcal{E}(\rho_2) \parallel \mathcal{F}(\sigma_2)) = -\frac{1}{2} \log(\delta(\kappa)), \quad (163)$$

where the adaptive strategy has $\rho_1 = \sigma_1 = |00\rangle\langle 00|$ and $\rho_2 = |01\rangle\langle 01|$, $\sigma_2 = (1 - \kappa/2) |11\rangle\langle 11| + \kappa/2 |1\rangle\langle 1|$. We also find:

$$\begin{aligned} D(\mathcal{E}(\rho_1) \parallel \mathcal{F}(\sigma_1)) - D(\rho_1 \parallel \sigma_1) &= D(\mathcal{E}(\rho_1) \parallel \mathcal{F}(\rho_1)) \\ &= D(|0\rangle\langle 0| \parallel (1 - \kappa/2) |+\rangle\langle +| + \kappa/2 |-\rangle\langle -|) \\ &= -\frac{1}{2} \left[\log\left(\frac{\kappa}{2}\right) + \log\left(1 - \frac{\kappa}{2}\right) \right], \end{aligned} \quad (164)$$

$$\begin{aligned} D(\mathcal{E}(\rho_2) \parallel \mathcal{F}(\sigma_2)) - D(\rho_2 \parallel \sigma_2) \\ &= \log(\delta(\kappa)) + \frac{1}{2} \left[\log\left(\frac{\kappa}{2}\right) + \log\left(1 - \frac{\kappa}{2}\right) \right], \end{aligned} \quad (165)$$

where interestingly (164) is always larger than (165), and so it is the first step that increases the distinguishability the most (when measured in terms of relative entropy). This also implies that this adaptive strategy is not asymptotically optimal, as

$$\frac{1}{2} D(\mathcal{E}(\rho_2) \parallel \mathcal{F}(\sigma_2)) < D(\mathcal{E}(\rho_1) \parallel \mathcal{F}(\rho_1)), \quad (166)$$

and hence it is asymptotically better to just use a parallel strategy with tensor copies of ρ_1 as an input state. Nevertheless, we will see that the adaptive strategy is still far superior in a regime where the number of channel uses m is not too large: It is well known (see, e.g., [26, Prop. 4.66]) that for all states ρ, σ

$$\lim_{\alpha_p \rightarrow 0} D_H^{\alpha_p}(\rho \parallel \sigma) = D_H^0(\rho \parallel \sigma) = -\log(\text{Tr}(\rho^0 \sigma)). \quad (167)$$

In this specific example, for all input states $\nu \in \mathcal{D}(\mathcal{H}_R \otimes \mathbb{C}^2 \otimes \mathbb{C}^2)$, it can be shown that

$$D_H^0(\mathcal{E}(\nu) \parallel \mathcal{F}(\nu)) \leq 2. \quad (168)$$

To see this, we start by observing that the second input system of our channels (previously called ω) can be treated as classical, and hence the maximum is attained at either $\omega = |0\rangle\langle 0|$ or $\omega = |1\rangle\langle 1|$ (this follows from joint convexity). If we choose the former, the channel output does not depend on the remaining input state and one finds

$$\begin{aligned} D_H^0(\mathcal{E}(\rho_{RA} \otimes |0\rangle\langle 0|) \parallel \mathcal{F}(\rho_{RA} \otimes |0\rangle\langle 0|)) \\ = -\log \text{Tr}(|0\rangle\langle 0| (|+\rangle\langle +| (1 - \kappa/2) + |-\rangle\langle -| (\kappa/2))) = 1. \end{aligned} \quad (169)$$

For the second choice of ω one finds that $D_H^0(\mathcal{E}(\nu)\|\mathcal{F}(\nu)) = 0$ unless $\nu = \rho_R \otimes |0\rangle\langle 0| \otimes |1\rangle\langle 1|$ and then

$$D_H^0(\mathcal{E}(|0\rangle\langle 0| \otimes |1\rangle\langle 1|)\|\mathcal{F}(|0\rangle\langle 0| \otimes |1\rangle\langle 1|)) \\ = -\log \text{Tr}(|0\rangle\langle 0|(\mathbb{1}/4(1+\kappa))) \leq 2. \quad (170)$$

An analogous argument also immediately yields $D_H^0(\mathcal{E}^{\otimes m}(\nu)\|\mathcal{F}^{\otimes m}(\nu)) \leq 2m$ for any (potentially entangled) ν . Hence, for fixed m , the rate

$$\frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu)\|\mathcal{F}^{\otimes m}(\nu)) \quad (171)$$

of any parallel strategy can be brought down to 2 by making the type I error threshold α_p small enough, whereas the mentioned adaptive strategy achieves zero type I error and a type II error rate that becomes arbitrarily large as $\kappa \rightarrow 0$.

To actually calculate the performance of a parallel strategy where the input ρ_1 is used m times, we can use the second-order asymptotics of hypothesis testing relative entropy to relative entropy [61] (this only works here because our input state is a product state). Define $\rho = \mathcal{E}(\rho_1)$, $\sigma = \mathcal{F}(\rho_1)$. Then [61, Thm. 5] implies

$$\frac{1}{m} D_H^{\alpha_p}(\rho^{\otimes m}\|\sigma^{\otimes m}) \\ \geq D(\rho\|\sigma) + \sqrt{\frac{V}{m}} \Phi^{-1}\left(\alpha_p - \frac{1}{\sqrt{m}} \frac{CT^3}{\sqrt{V^3}}\right), \quad (172)$$

$$\frac{1}{m} D_H^{\alpha_p}(\rho^{\otimes m}\|\sigma^{\otimes m}) \leq D(\rho\|\sigma) \\ + \sqrt{\frac{V}{m}} \Phi^{-1}\left(\alpha_p + \frac{1}{\sqrt{m}} \left(\frac{CT^3}{\sqrt{V^3}} + 2\right)\right), \quad (173)$$

where Φ^{-1} is the inverse of the cumulative distribution function of the standard normal distribution, $C \leq 0.4784$, and

$$V \equiv V(\rho\|\sigma) := \text{Tr}[\rho(\log \rho - \log \sigma - D(\rho\|\sigma))^2], \quad (174)$$

$$T^3 \equiv T^3(\rho\|\sigma) \quad (175)$$

$$:= \sum_{i,j} \lambda_i |\langle x_i | y_j \rangle|^2 |\log(\lambda_i) - \log(\mu_j) - D(\rho\|\sigma)|^3 \quad (176)$$

for spectral decompositions $\rho = \sum_i \lambda_i |x_i\rangle\langle x_i|$ and $\sigma = \sum_j \mu_j |y_j\rangle\langle y_j|$.

Note that, instead of comparing the type II error decay rate of a parallel strategy with m channel uses to the corresponding rate of an adaptive strategy with two channel uses, it might be more intuitive to compare the parallel strategy to a setup where the adaptive strategy with two channel uses is repeated $m/2$ times in parallel (so that the adaptive and parallel strategies have the same number of channel uses, and comparing rates is the same as comparing type II errors). For this example, since the type I error of the adaptive strategy was chosen to be zero, this is actually equivalent, as

$$\frac{1}{2k} D_H^0((\mathcal{E}(\rho_2))^{\otimes k}\|(\mathcal{F}(\sigma_2))^{\otimes k}) \\ = \frac{1}{2} D_H^0(\mathcal{E}(\rho_2)\|\mathcal{F}(\sigma_2)) = -\frac{1}{2} \log(\delta(\kappa)). \quad (177)$$

Our theorem (Theorem 1) gives an upper bound on the extent to which all finite-length parallel strategies can have worse type II errors compared to the adaptive strategy,

or equivalently how large m has to be chosen to achieve similar performances. For this example specifically, due to (164) and (165) we can choose $\ell = 1$ in Theorem 1, and hence also the parallel input state ν in Theorem 1 will just be $\nu = \rho_1^{\otimes m}$.

Figure 3 depicts our lower bound (from Theorem 1) on the performance of a parallel strategy for the given adaptive strategy, together with the actual performance of the parallel strategy choosing $\rho_1^{\otimes m}$ as the input state. For the figure we chose $\kappa = 2^{-50}$, $\alpha_a = 0$, and $\alpha_p = 2^{-5}$. The parameters are chosen to make the following features nicely visible simultaneously in one plot: (i) the range of m where the parallel strategy is worse, (ii) the range of m where it surpasses the adaptive strategy, and (iii) our bound. For c_ℓ in Theorem 1 we used (42) with a numerical optimization over γ_1 . One finds that there is a range of values for m for which the adaptive strategy is better, and the parallel strategy is lower bounded fairly tightly by our bound. As the given adaptive strategy is not asymptotically optimal it is eventually surpassed by the parallel strategy.

V. DISCUSSION AND OUTLOOK

In this final section, we discuss several pathways through which one could hope to improve or extend our result, together with a discussion of obstacles we encountered on these pathways, and their relation to different open problems in quantum channel discrimination.

First of all, one might hope to be able to remove the factor of $1 - \alpha_a$ appearing in (39), perhaps at the cost of an additional error term proportional to $\log(1 - \alpha_a)$. If this additional error term decays in m and n (say, as long as α_a is bounded away from one), this would prove the strong converse property for quantum channel discrimination. To see this, suppose we could show something along the lines of the following inequality:

$$\frac{1}{m} D_H^{\alpha_p}(\mathcal{E}^{\otimes m}(\nu)\|\mathcal{F}^{\otimes m}(\nu)) \stackrel{?}{\geq} \frac{1}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)) \\ - \frac{C_1 n}{\sqrt{m}} \log\left(\frac{8}{\alpha_p}\right) - \frac{C_2}{n} \log\left(\frac{1}{1 - \alpha_a}\right). \quad (178)$$

Now, fixing any $\alpha_a \in (0, 1)$ and taking limits (in this order) $m \rightarrow \infty, n \rightarrow \infty, \alpha_p \rightarrow 0$, we would find that for an arbitrary $\alpha_a \in (0, 1)$

$$D^{\text{reg}}(\mathcal{E}\|\mathcal{F}) = \lim_{\alpha_p \rightarrow 0} \lim_{n \rightarrow \infty} D_H^{\alpha_p}(\mathcal{E}^{\otimes n}\|\mathcal{F}^{\otimes n}) \quad (179)$$

$$\stackrel{?}{\geq} \limsup_{n \rightarrow \infty} \frac{1}{n} D_H^{\alpha_a}(\mathcal{E}(\rho_n)\|\mathcal{F}(\sigma_n)), \quad (180)$$

i.e., allowing a finite type I error $\alpha_a \in (0, 1)$ does not improve the best achievable type II error rate. This is precisely the strong converse property. Establishing the strong converse property for quantum channel discrimination is an important and very interesting open problem, which is still unsolved despite serious efforts. Significant progress has been made in [62], and another recent attempt was made in [63]; see also [64].

One might hope to obtain such a variant of our bound without the factor $1 - \alpha_a$, by, instead of transitioning from hypothesis testing entropy to relative entropy in (122),

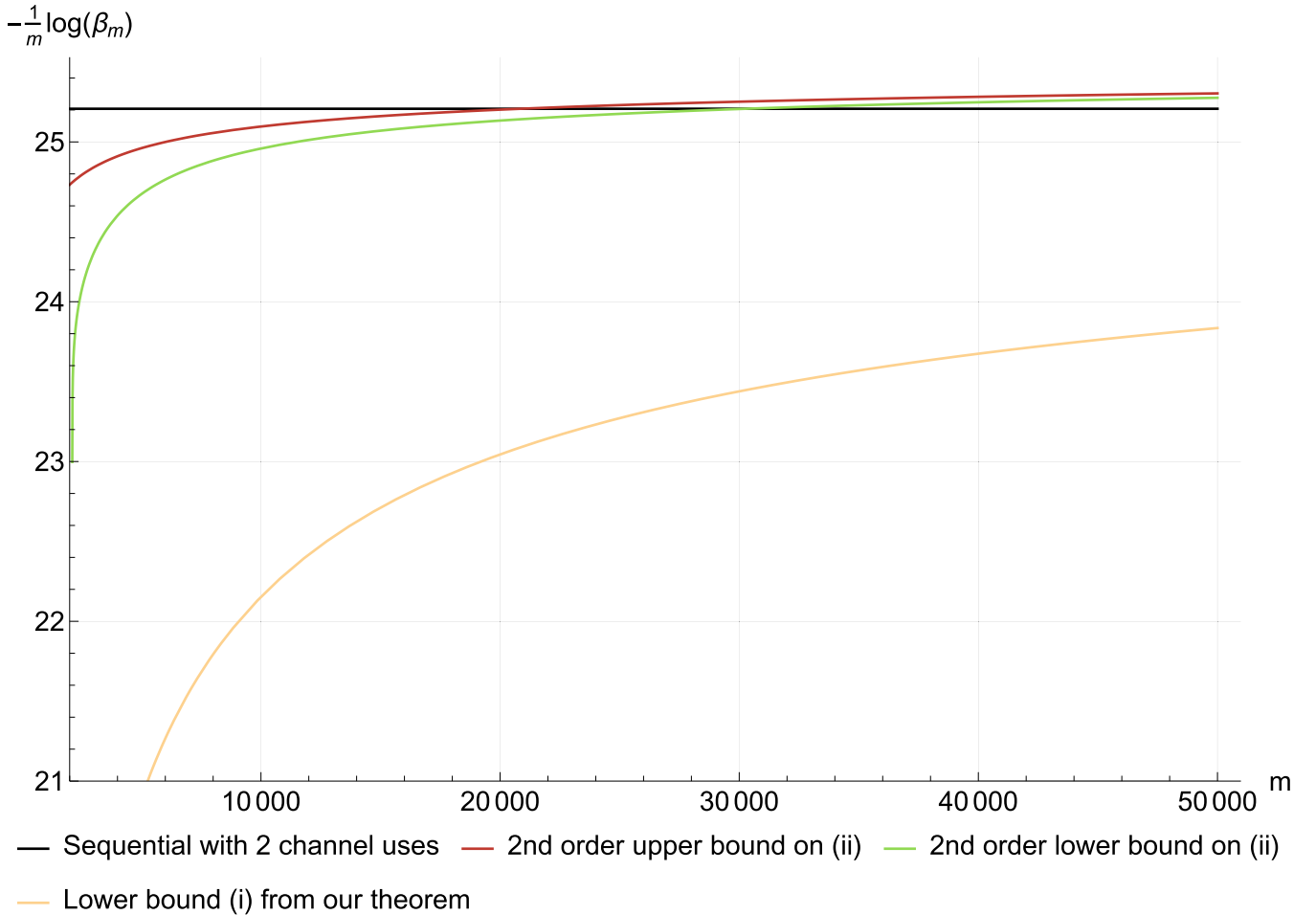


Fig. 3. Illustration of the type II error decay rate per channel use of a simple adaptive and parallel strategy for a specific pair of channels (see (159) and (160) for the definitions of the channels $\mathcal{E}$ and $\mathcal{F}$, respectively, where $\kappa = 2^{-50}$). We compare a fixed adaptive strategy with two channel uses (constant black line) to (i) our lower bound on the performance of a parallel strategy (yellow line) and (ii) the actual performance of a parallel strategy (red and green lines), which are plotted as functions of the number of parallel channel uses m . The black line shows the value of (163), i.e., the type II error exponent for the given adaptive strategy with two channel uses and type I error $\alpha_a = 0$. This can alternatively be thought of as the rate of repeating the two-step adaptive strategy $m/2$ times in parallel. The yellow line shows the lower bound on the parallel strategy from our theorem (i.e., the right-hand side of (39)), choosing $\alpha_p = 2^{-5}$. For this specific example we can calculate the parallel input state ν of our theorem, and while we cannot explicitly find the optimal POVM and corresponding type II error (i.e., we cannot explicitly calculate the left-hand side of (39)), we can bound it from above and below using the second-order asymptotics of the hypothesis testing relative entropy, which is shown in the red and green lines, corresponding to the values of (173) and (172). We see that for small m there is a gap between the adaptive and parallel strategies; i.e., the adaptive strategy offers an advantage. This advantage disappears once m gets larger and in this specific example the chosen adaptive strategy even eventually gets surpassed by the parallel strategy, as the adaptive strategy turns out not to be asymptotically optimal.

moving to an α -Rényi relative entropy instead and subsequently employing the known relations between Rényi relative entropies and smoothed max-relative entropies. It turns out that this will at some point require bounding the difference between Rényi relative entropies of order $1 - \alpha$ and $1 + \alpha$. This is possible using Lemma 5, however at the cost of an error term $(c_\gamma(\rho_n \|\sigma_n))^2$, which will scale quadratically in n . As this does not pick up any dependence in m , such an approach will not lead to anything useful in the asymptotic limit, and is hence unsuccessful.

One further interesting question deals with the relation between ℓ and n in Theorem 1. One might hope that an optimal adaptive strategy achieves the maximum possible distinguishability gain per channel use, i.e.

$$D^A(\mathcal{E} \|\mathcal{F}) = \sup_{\rho, \sigma \in \mathcal{D}(R \otimes A)} D(\mathcal{E}(\rho) \|\mathcal{F}(\sigma)) - D(\rho \|\sigma), \quad (181)$$

after some finite number of channel uses, or at least comes very close to it. If that was the case, one could bound ℓ by a constant and would remove the dependence on n in (28). However, note that the supremum in (181) goes over arbitrarily large reference systems R and hence does not need to be achieved, and we are not aware of any bounds on the system size R to get close to the optimum value. Alternatively, since $D^A(\mathcal{E} \|\mathcal{F}) = D^{\text{reg}}(\mathcal{E} \|\mathcal{F})$ also a bound on the speed of convergence of the regularized channel divergence would do, which we are however also not aware of (Ref. [62] proves such a bound for the sandwiched Rényi divergence, which unfortunately cannot be extended straightforwardly to relative entropy). Hence, for now we have to assume that ℓ can in general be n , and that to adequately simulate an adaptive strategy with n channel uses, one might need to employ more than n^2 parallel channel uses to keep the error term in our bound small.

While our result becomes quite powerful in the asymmetric asymptotic setting, Theorem 1 can in principle be applied for any combination of type I and II errors. Hence, it could be interesting to apply it to scenarios where the type I error decays with m , say as $\alpha_p = 2^{-km}$ for some constant k . In this case, one would want to apply (39), to have at least a chance of the error term being bounded; however – depending on k – the corresponding condition (38) might not always be fulfilled. It would be interesting to see whether this condition on m could perhaps be relaxed, to allow for a wider range of k in this specific scenario.

As mentioned already in Remark 9, our bounds could be significantly tightened if we were able to employ second-order expansions, which however requires controlling the variance of the relative entropy $V(\rho_n|\sigma_n)$ in n , which we are unable to do. This seems to be again related to the strong converse problem, as second-order expansions for the hypothesis testing relative entropy imply the strong converse property. Already for the parallel case, if one was able to show

$$V(\mathcal{E}^{\otimes n}(\nu)\|\mathcal{F}^{\otimes n}(\nu)) = \mathcal{O}(n), \quad (182)$$

where $\nu \in \mathcal{D}(R^{\otimes n} \otimes A^{\otimes n})$ is the optimal joint input state, this would be a very significant step towards proving the strong converse for quantum channel discrimination. Conversely, examples where this scales faster than linearly in n are quite likely to lend themselves to counterexamples for the strong converse property. We are not aware of any such examples: we leave this question open for further studies.

Finally, we believe that it should be possible to generalize our result to the infinite-dimensional setting with only minor modifications, which we also leave open to future investigations.

APPENDIX

A. Proof of Lemma 5 (Petz–Rényi Continuity Bound in α)

Proof of Lemma 5: Note first that for (105), where we did not explicitly require $D(\rho|\sigma) < \infty$, we can restrict to that case, as otherwise also $D_{1+\delta}(\rho|\sigma) = \infty$ and the statement is trivial. Hence, we can assume σ to be invertible (otherwise restrict to the subspace where σ is supported). Define the operator $X = \rho \otimes (\sigma^{-1})^T$, and the canonical purification of ρ on $\mathcal{H} \otimes \mathcal{H}$:

$$|\phi\rangle = \sum_i \sqrt{\rho} |i\rangle \otimes |i\rangle, \quad (183)$$

where $|i\rangle$ is an orthonormal basis of $\mathcal{H}$. Then, for all $\delta \in [-1, 1]$:

$$D_{1+\delta}(\rho|\sigma) = \frac{1}{\delta} \log(\langle \phi | X^\delta | \phi \rangle) \quad (184)$$

and

$$D(\rho|\sigma) = \langle \phi | \log(X) | \phi \rangle. \quad (185)$$

For all $t > 0$ and $\delta \in \mathbb{R}$, the first term in the expansion of t^δ around $\delta = 0$ is $\delta \ln(t)$. Hence let us write $t^\delta = 1 + \delta \ln(t) + r_\delta(t)$ where $r_\delta(t) := t^\delta - \delta \ln(t) - 1$. Since $1 + x \leq e^x$ for all $x \in \mathbb{R}$ we see that

$$r_\delta(t) = t^\delta - \delta \ln(t) - 1 \quad (186)$$

$$\leq t^\delta + e^{-\delta \ln(t)} - 2 \quad (187)$$

$$= e^{\delta \ln(t)} + e^{-\delta \ln(t)} - 2 \quad (188)$$

$$= 2(\cosh(\delta \ln(t)) - 1) =: s_\delta(t). \quad (189)$$

It is easy to see that $s_{-\delta}(t) = s_\delta(t)$ and $s_\delta(t) = s_{\gamma\delta}(t^{1/\gamma})$ for all $\gamma \in \mathbb{R}$. Also, it is easy to verify that $s_\delta(t)$ is monotonically increasing in t for $t \geq 1$ and concave in t if $\delta \leq \frac{1}{2}$ and $t \geq 3$. For all $t \geq 0$, either t or $1/t$ will be larger than or equal to one, and so we can always use the monotonicity to write

$$s_\delta(t) = s_\delta\left(\frac{1}{t}\right) \leq s_\delta\left(t + \frac{1}{t}\right). \quad (190)$$

Using this, we get for all $t \geq 0$ and $\gamma \in (0, 1]$

$$s_\delta(t) = s_{\delta/\gamma}(t^\gamma) \leq s_{\delta/\gamma}(t^\gamma + t^{-\gamma}) \quad (191)$$

$$\leq s_{\delta/\gamma}(t^\gamma + t^{-\gamma} + 1). \quad (192)$$

It is easy to see that the real function $x + 1/x$ has a global minimum for $x > 0$ at $x = 1$ and hence the argument at the right hand side is guaranteed to be larger than 3. Hence we can use the concavity of $s_{\delta/\gamma}$ to write

$$\langle \phi | s_\delta(X) | \phi \rangle \leq \langle \phi | s_{\delta/\gamma}(X^\gamma + X^{-\gamma} + \mathbb{1}) | \phi \rangle \quad (193)$$

$$\leq s_{\delta/\gamma}(\langle \phi | X^\gamma + X^{-\gamma} + \mathbb{1} | \phi \rangle) \quad (194)$$

$$= s_{\delta/\gamma}(2^{\gamma c_\gamma(\rho\|\sigma)}), \quad (195)$$

where the second inequality is a well-known property of concave functions (Jensen's inequality) and c_γ is defined in (104). Finally, we use Taylor's theorem with the Lagrange form of the remainder to bound

$$s_\delta(t) = s_0(t) + \frac{d}{d\delta} s_\delta(t) \Big|_{\delta=0} \delta + \frac{1}{2} \frac{d^2}{d\delta^2} s_\delta(t) \Big|_{\delta=\xi} \delta^2 \quad (196)$$

$$= \delta^2 (\ln(t))^2 \cosh(\xi \ln(t)) \quad (197)$$

$$\leq \delta^2 (\ln(t))^2 \cosh(\delta \ln(t)), \quad (198)$$

for all $t \geq 0$ and some $\xi \in (0, \delta)$, where we have used that $s_0(t) = \frac{d}{d\delta} s_\delta(t) \Big|_{\delta=0} = 0$. Hence,

$$\begin{aligned} \langle \phi | s_\delta(X) | \phi \rangle &\leq s_{\delta/\gamma}(2^{\gamma c_\gamma(\rho\|\sigma)}) \\ &\leq (\delta \ln(2) c_\gamma(\rho\|\sigma))^2 \cosh(\delta \ln(2) c_\gamma(\rho\|\sigma)). \end{aligned} \quad (199)$$

To finally prove our claims, let us start with the case where $\delta > 0$. Then,

$$D_{1+\delta}(\rho|\sigma) = \frac{1}{\delta} \log(\langle \phi | X^\delta | \phi \rangle) \quad (200)$$

$$= \frac{1}{\delta} \log(1 + \delta \ln(2) \langle \phi | \log(X) | \phi \rangle + \langle \phi | r_\delta(X) | \phi \rangle) \quad (201)$$

$$= \frac{1}{\delta} \log(1 + \delta \ln(2) D(\rho|\sigma) + \langle \phi | r_\delta(X) | \phi \rangle) \quad (202)$$

$$\leq \frac{1}{\delta} \log(1 + \delta \ln(2) D(\rho|\sigma) + \langle \phi | s_\delta(X) | \phi \rangle) \quad (203)$$

$$= \frac{1}{\delta} \log(1 + \delta \ln(2) D(\rho|\sigma)) \quad (204)$$

$$+ \frac{1}{\delta} \log\left(1 + \frac{\langle \phi | s_\delta(X) | \phi \rangle}{1 + \delta \ln(2) D(\rho|\sigma)}\right)$$

$$\leq D(\rho|\sigma) + \frac{1}{\delta} \log(1 + \langle \phi | s_\delta(X) | \phi \rangle) \quad (205)$$

$$\leq D(\rho|\sigma)$$

$$+ \frac{1}{\delta} \log(1 + (\delta \ln(2) c_\gamma(\rho\|\sigma))^2 \cosh(\delta \ln(2) c_\gamma(\rho\|\sigma))) + \frac{1}{n} \frac{2}{\alpha - 1} \log\left(\frac{1}{1 - \varepsilon}\right), \quad (215)$$

$$\leq D(\rho\|\sigma) + \delta \ln(2) (c_\gamma(\rho\|\sigma))^2, \quad (207)$$

where the third-to-last inequality uses $\log(1+x) \leq \frac{x}{\ln(2)}$ and (for the second term) the fact that δ and $D(\rho\|\sigma)$ are non-negative. The final inequality follows from the fact that $k^2 - \ln(1 + k^2 \cosh(k))$ is monotonically increasing in k and hence positive.

If $\delta < 0$, the issue arises that $1 + \delta \ln(2) D(\rho\|\sigma)$ no longer has to be greater than 1 (it might in fact be negative) and so we have to apply a slightly different argument. Starting analogously, we get

$$D_{1+\delta}(\rho\|\sigma) = \frac{1}{\delta} \log(1 + \delta \ln(2) D(\rho\|\sigma) + \langle \phi | r_\delta(X) | \phi \rangle) \quad (208)$$

$$\geq \frac{1}{\delta} \log(1 + \delta \ln(2) D(\rho\|\sigma) + \langle \phi | s_\delta(X) | \phi \rangle) \quad (209)$$

$$\geq D(\rho\|\sigma) + \frac{1}{\delta \ln(2)} \langle \phi | s_\delta(X) | \phi \rangle \quad (210)$$

$$\geq D(\rho\|\sigma) + \delta \ln(2) (c_\gamma(\rho\|\sigma))^2 \cosh(\delta \ln(2) c_\gamma(\rho\|\sigma)), \quad (211)$$

where we used $\log(1+x) \leq \frac{x}{\ln(2)}$ in the second inequality. Finally, if we assume that $|\delta| \leq \frac{\log 3}{2c_\gamma(\rho\|\sigma)}$, then $\ln(2) \cosh(\ln(3)/2) < 1$ and so we get

$$D_{1+\delta}(\rho\|\sigma) \geq D(\rho\|\sigma) + \delta (c_\gamma(\rho\|\sigma))^2. \quad (212)$$

Note that since $c_\gamma(\rho\|\sigma) \geq \frac{\log(3)}{\gamma}$ the condition $|\delta| \leq \gamma/2$ is automatically fulfilled. $\square$

B. Proof of Lemma 6 (AEP Convergence Bound)

Proof of Lemma 6: We start with the proof of (117). We use [41, Prop. 4], which states that for all $\alpha \in [0, 1)$ and all $\varepsilon \in [0, 1)$

$$D_{\max}^\varepsilon(\rho\|\sigma) \geq D_\alpha(\rho\|\sigma) + \frac{2}{\alpha - 1} \log\left(\frac{1}{1 - \varepsilon}\right). \quad (213)$$

Note that the statement in [41] is given for smoothing in trace-distance, but since the trace distance is always less than the sine distance, for fixed ε the trace-distance-smoothed max-divergence is smoothed over a larger ball, and hence also always smaller, which is why the statement for the purified-distance-smoothed max-divergence is implied. We apply this to $\rho^{\otimes n}$ and $\sigma^{\otimes n}$ and use the additivity of the Petz-Rényi relative entropy to get:

$$\frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) \geq D_\alpha(\rho\|\sigma) + \frac{1}{n} \frac{2}{\alpha - 1} \log\left(\frac{1}{1 - \varepsilon}\right). \quad (214)$$

Now, we combine this with (108) of Lemma 5 to get

$$\begin{aligned} \frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) &\geq D(\rho\|\sigma) \\ &+ \ln(2) \cosh\left(\frac{\log(3)}{2}\right) c_\gamma(\rho\|\sigma)^2 (\alpha - 1) \end{aligned}$$

together with the condition $0 \leq 1 - \alpha \leq \frac{\log 3}{2c_\gamma(\rho\|\sigma)}$. We are now free to choose α to optimize the right-hand side. It is easy to see that the right-hand side will be maximal if both terms are equal, which is achieved for

$$\begin{aligned} 1 - \alpha &= \sqrt{\frac{2 \log\left(\frac{1}{1 - \varepsilon}\right)}{n \ln(2) \cosh\left(\frac{\log(3)}{2}\right) c_\gamma(\rho\|\sigma)^2}} \\ &= \frac{4}{K_1 c_\gamma(\rho\|\sigma)} \sqrt{\frac{\log\left(\frac{1}{1 - \varepsilon}\right)}{n}}. \end{aligned} \quad (216)$$

Hence we get

$$\frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) \geq D(\rho\|\sigma) - \frac{K_1}{\sqrt{n}} c_\gamma(\rho\|\sigma) \sqrt{\log\left(\frac{1}{1 - \varepsilon}\right)}, \quad (217)$$

together with the condition

$$n \geq \log\left(\frac{1}{1 - \varepsilon}\right) \left(\frac{4}{K_1 \log(3)}\right)^2. \quad (218)$$

In order to get an expression without a condition on n we have to choose a different α . Choosing

$$1 - \alpha = \frac{\log(3)}{2c_\gamma(\rho\|\sigma)\sqrt{n}} \quad (219)$$

satisfies the condition on $1 - \alpha$ for all n and gives:

$$\begin{aligned} \frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) &\geq D(\rho\|\sigma) \\ &- \frac{c_\gamma(\rho\|\sigma)}{\sqrt{n}} \left[\frac{\ln(2) \log(3)}{2} \cosh\left(\frac{\log(3)}{2}\right) \right. \\ &\quad \left. + \frac{4}{\log(3)} \log\left(\frac{1}{1 - \varepsilon}\right) \right]. \end{aligned} \quad (220)$$

The simplification (112) follows from $\log(3) > 1$ and

$$\frac{\ln(2) \log(3)}{2} \cosh\left(\frac{\log(3)}{2}\right) < 1. \quad (221)$$

The second equation (119) is proved very analogously. We start again with a relation to a Rényi relative entropy for $\alpha > 1$ [26, Prop. 4.61]:

$$D_{\max}^\varepsilon(\rho\|\sigma) \leq D_\alpha(\rho\|\sigma) + \frac{2}{\alpha - 1} \log\left(\frac{1}{\varepsilon}\right) + \log\left(\frac{1}{1 - \varepsilon^2}\right). \quad (222)$$

Note that [26, Prop. 4.61] uses the sandwiched Rényi divergence, which however is always less than the Petz, and hence the above statement is implied. We again apply this to n tensor powers of ρ and σ and combine it with (105) from Lemma 5 to obtain:

$$\begin{aligned} \frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n}\|\sigma^{\otimes n}) &\leq D(\rho\|\sigma) + \ln(2) c_\gamma(\rho\|\sigma)^2 (\alpha - 1) \\ &+ \frac{2}{n(\alpha - 1)} \log\left(\frac{1}{\varepsilon}\right) + \frac{1}{n} \log\left(\frac{1}{1 - \varepsilon^2}\right), \end{aligned} \quad (223)$$

together with the condition

$$\alpha - 1 \leq \frac{\gamma}{2}. \quad (224)$$

The right hand side is minimized again for

$$\alpha - 1 = \sqrt{\frac{2 \ln(2) \log\left(\frac{1}{\varepsilon}\right)}{n c_\gamma(\rho \parallel \sigma)^2}} = \frac{4}{K_2 c_\gamma(\rho \parallel \sigma)} \sqrt{\frac{\log\left(\frac{1}{\varepsilon}\right)}{n}}, \quad (225)$$

which gives

$$\begin{aligned} \frac{1}{n} D_{\max}^\varepsilon(\rho^{\otimes n} \parallel \sigma^{\otimes n}) &\leq D(\rho \parallel \sigma) \\ &+ \frac{K_2}{\sqrt{n}} c_\gamma(\rho \parallel \sigma) \sqrt{\log\left(\frac{1}{\varepsilon}\right)} + \frac{1}{n} \log\left(\frac{1}{1 - \varepsilon^2}\right), \end{aligned} \quad (226)$$

as well as the condition

$$n \geq \log\left(\frac{1}{\varepsilon}\right) \left(\frac{8}{\gamma c_\gamma(\rho \parallel \sigma) K_2}\right)^2. \quad (227)$$

Alternatively, choosing

$$\alpha - 1 = \frac{\log(3)}{2 c_\gamma(\rho \parallel \sigma) \sqrt{n}} \quad (228)$$

will again give (111), and the simplified form (113) follows by the same argument that $\log(3) > 1$ and

$$\frac{\ln(2) \log(3)}{2} < 1. \quad (229)$$

This concludes the proof. $\square$

ACKNOWLEDGMENT

The authors would like to thank the anonymous referees for various helpful comments and suggestions. For the purpose of open access, the authors have applied a creative commons attribution (CC BY) license to any author accepted manuscript version arising from this submission.

REFERENCES

- [1] C. W. Helstrom, "Quantum detection and estimation theory," *J. Stat. Phys.*, vol. 1, no. 2, pp. 231–252, Jun. 1969.
- [2] A. S. Holevo, "An analogue of statistical decision theory and noncommutative probability theory," *Trudy Moskovskogo Matematicheskogo Obščestva*, vol. 26, pp. 133–149, 1972. [Online]. Available: <https://mathscinet.ams.org/mathscinet-getitem?mr=0365809>
- [3] A. S. Kholevo, "On asymptotically optimal hypothesis testing in quantum statistics," *Theory Probab. Appl.*, vol. 23, no. 2, pp. 411–415, Mar. 1979.
- [4] F. Hiai and D. Petz, "The proper formula for relative entropy and its asymptotics in quantum probability," *Commun. Math. Phys.*, vol. 143, no. 1, pp. 99–114, Dec. 1991.
- [5] T. Ogawa and H. Nagaoka, "Strong converse and Stein's lemma in quantum hypothesis testing," *IEEE Trans. Inf. Theory*, vol. 46, no. 7, pp. 2428–2433, Nov. 2000.
- [6] H. Nagaoka, "The converse part of the theorem for quantum Hoeffding bound," Nov. 2006, *arXiv:quant-ph/0611289*.
- [7] K. M. R. Audenaert et al., "Discriminating states: The quantum Chernoff bound," *Phys. Rev. Lett.*, vol. 98, no. 16, Apr. 2007, Art. no. 160501.
- [8] M. Hayashi, "Error exponent in asymmetric quantum hypothesis testing and its application to classical-quantum channel coding," *Phys. Rev. A, Gen. Phys.*, vol. 76, no. 6, Dec. 2007, Art. no. 062301.
- [9] K. M. R. Audenaert, M. Nussbaum, A. Szkoła, and F. Verstraete, "Asymptotic error rates in quantum hypothesis testing," *Commun. Math. Phys.*, vol. 279, no. 1, pp. 251–283, Apr. 2008.
- [10] M. Nussbaum and A. Szkoła, "The Chernoff lower bound for symmetric quantum hypothesis testing," *Ann. Statist.*, vol. 37, no. 2, Apr. 2009.
- [11] K. M. R. Audenaert, M. Mosonyi, and F. Verstraete, "Quantum state discrimination bounds for finite sample size," *J. Math. Phys.*, vol. 53, no. 12, Dec. 2012, Art. no. 122205.
- [12] J. Bae and L.-C. Kwek, "Quantum state discrimination and its applications," *J. Phys. A, Math. Theor.*, vol. 48, no. 8, Feb. 2015, Art. no. 083001.
- [13] G. Chiribella, G. M. D'Ariano, and P. Perinotti, "Memory effects in quantum channel discrimination," *Phys. Rev. Lett.*, vol. 101, no. 18, Oct. 2008, Art. no. 180501.
- [14] R. Duan, Y. Feng, and M. Ying, "Perfect distinguishability of quantum operations," *Phys. Rev. Lett.*, vol. 103, no. 21, Nov. 2009, Art. no. 210501.
- [15] M. Hayashi, "Discrimination of two channels by adaptive methods and its application to quantum system," *IEEE Trans. Inf. Theory*, vol. 55, no. 8, pp. 3807–3820, Aug. 2009.
- [16] A. W. Harrow, A. Hassidim, D. W. Leung, and J. Watrous, "Adaptive versus nonadaptive strategies for quantum channel discrimination," *Phys. Rev. A, Gen. Phys.*, vol. 81, no. 3, Mar. 2010, Art. no. 032339.
- [17] E. Martínez Vargas et al., "Quantum sequential hypothesis testing," *Phys. Rev. Lett.*, vol. 126, no. 18, May 2021, Art. no. 180502.
- [18] Y. Li, V. Y. F. Tan, and M. Tomamichel, "Optimal adaptive strategies for sequential quantum hypothesis testing," *Commun. Math. Phys.*, vol. 392, no. 3, pp. 993–1027, Jun. 2022.
- [19] Y. Li, C. Hirche, and M. Tomamichel, "Sequential quantum channel discrimination," in *Proc. IEEE Int. Symp. Inf. Theory (ISIT)*, Espoo, Finland, Jun. 2022, pp. 270–275.
- [20] M. M. Wilde, M. Berta, C. Hirche, and E. Kaur, "Amortized channel divergence for asymptotic quantum channel discrimination," *Lett. Math. Phys.*, vol. 110, no. 8, pp. 2277–2336, Aug. 2020.
- [21] F. Salek, M. Hayashi, and A. Winter, "Usefulness of adaptive strategies in asymptotic quantum channel discrimination," *Phys. Rev. A, Gen. Phys.*, vol. 105, no. 2, Feb. 2022, Art. no. 022419.
- [22] K. Fang, O. Fawzi, R. Renner, and D. Sutter, "Chain rule for the quantum relative entropy," *Phys. Rev. Lett.*, vol. 124, no. 10, Mar. 2020, Art. no. 100501.
- [23] X. Wang and M. M. Wilde, "Resource theory of asymmetric distinguishability for quantum channels," *Phys. Rev. Res.*, vol. 1, no. 3, Dec. 2019, Art. no. 033169.
- [24] H. Umegaki, "Conditional expectation in an operator algebra, IV (entropy and information)," *Kodai Math. Seminar Rep.*, vol. 14, no. 2, pp. 59–85, Jan. 1962.
- [25] G. Lindblad, "Completely positive maps and entropy inequalities," *Commun. Math. Phys.*, vol. 40, no. 2, pp. 147–151, Jun. 1975.
- [26] S. Khatri and M. M. Wilde, "Principles of quantum communication theory: A modern approach," 2020, *arXiv:2011.04672*.
- [27] A. Uhlmann, "The 'transition probability' in the state space of a *-algebra," *Rep. Math. Phys.*, vol. 9, no. 2, pp. 273–279, Apr. 1976.
- [28] A. E. Rastegin, "Relative error of state-dependent cloning," *Phys. Rev. A, Gen. Phys.*, vol. 66, no. 4, Oct. 2002, Art. no. 042304.
- [29] A. E. Rastegin, "A lower bound on the relative error of mixed-state cloning and related operations," *J. Opt. B, Quantum Semiclass. Opt.*, vol. 5, no. 6, pp. S647–S650, Oct. 2003.
- [30] A. Gilchrist, N. K. Langford, and M. A. Nielsen, "Distance measures to compare real and ideal quantum processes," *Phys. Rev. A, Gen. Phys.*, vol. 71, no. 6, Jun. 2005, Art. no. 062310.
- [31] A. E. Rastegin, "Sine distance for quantum states," Feb. 2006, *arXiv:quant-ph/0602112*.
- [32] N. Datta, "Min- and max- relative entropies and a new entanglement monotone," *IEEE Trans. Inf. Theory*, vol. 55, no. 6, pp. 2816–2826, Jun. 2009.
- [33] D. Petz, "Quasi-entropies for finite quantum systems," *Rep. Math. Phys.*, vol. 23, no. 1, pp. 57–65, Feb. 1986.
- [34] K. Matsumoto, "A new quantum version of f-divergence," in *Reality and Measurement in Algebraic Quantum Theory* (Springer Proceedings in Mathematics & Statistics), M. Ozawa, J. Butterfield, H. Halvorson, M. Rédei, Y. Kitajima, and F. Buscemi, Eds. Singapore: Springer, 2018, pp. 229–273.
- [35] F. Leditzky, E. Kaur, N. Datta, and M. M. Wilde, "Approaches for approximate additivity of the Holevo information of quantum channels," *Phys. Rev. A, Gen. Phys.*, vol. 97, no. 1, Jan. 2018, Art. no. 012332.
- [36] M. G. Díaz et al., "Using and reusing coherence to realize quantum processes," *Quantum*, vol. 2, p. 100, Oct. 2018.
- [37] F. Buscemi and N. Datta, "The quantum capacity of channels with arbitrarily correlated noise," *IEEE Trans. Inf. Theory*, vol. 56, no. 3, pp. 1447–1460, Mar. 2010.

- [38] F. G. S. L. Brandao and N. Datta, "One-shot rates for entanglement manipulation under non-entangling maps," *IEEE Trans. Inf. Theory*, vol. 57, no. 3, pp. 1754–1760, Mar. 2011.
- [39] L. Wang and R. Renner, "One-shot classical-quantum capacity and hypothesis testing," *Phys. Rev. Lett.*, vol. 108, no. 20, May 2012, Art. no. 200501.
- [40] A. Anshu, R. Jain, and N. A. Warsi, "Building blocks for communication over noisy quantum networks," *IEEE Trans. Inf. Theory*, vol. 65, no. 2, pp. 1287–1306, Feb. 2019.
- [41] X. Wang and M. M. Wilde, "Resource theory of asymmetric distinguishability," *Phys. Rev. Res.*, vol. 1, no. 3, Dec. 2019, Art. no. 033170.
- [42] M. Hayashi, *Quantum Information: An Introduction*. Berlin, Germany: Springer, 2006.
- [43] A. Anshu, M. Berta, R. Jain, and M. Tomamichel, "A minimax approach to one-shot entropy inequalities," *J. Math. Phys.*, vol. 60, no. 12, Dec. 2019, Art. no. 122201.
- [44] T. Cooney, M. Mosonyi, and M. M. Wilde, "Strong converse exponents for a quantum channel discrimination problem and quantum-feedback-assisted communication," *Commun. Math. Phys.*, vol. 344, no. 3, pp. 797–829, Jun. 2016.
- [45] V. Katariya and M. M. Wilde, "Evaluating the advantage of adaptive strategies for quantum channel distinguishability," *Phys. Rev. A, Gen. Phys.*, vol. 104, no. 5, Nov. 2021, Art. no. 052406.
- [46] M. Hayashi and S. Watanabe, "Finite-length analyses for source and channel coding on Markov chains," *Entropy*, vol. 22, no. 4, p. 460, Apr. 2020.
- [47] S. Watanabe and M. Hayashi, "Finite-length analysis on tail probability for Markov chain and application to simple hypothesis testing," *Ann. Appl. Probab.*, vol. 27, no. 2, pp. 811–845, Apr. 2017.
- [48] M. Hayashi and S. Watanabe, "Uniform random number generation from Markov chains: Non-asymptotic and asymptotic analyses," *IEEE Trans. Inf. Theory*, vol. 62, no. 4, pp. 1795–1822, Apr. 2016.
- [49] O. Fawzi, A. Shayeghi, and H. Ta, "A hierarchy of efficient bounds on quantum capacities exploiting symmetry," *IEEE Trans. Inf. Theory*, vol. 68, no. 11, pp. 7346–7360, Nov. 2022.
- [50] E. De Klerk, D. V. Pasechnik, and A. Schrijver, "Reduction of symmetric semidefinite programs using the regular *-representation," *Math. Program.*, vol. 109, no. 2, pp. 613–624, Mar. 2007.
- [51] M. F. Anjos and J. B. Lasserre, Eds., *Handbook on Semidefinite, Conic and Polynomial Optimization* (International Series in Operations Research & Management Science), vol. 166. Boston, MA, USA: Springer, 2012.
- [52] D. C. Gijswijt, "Matrix algebras and semidefinite programming techniques for codes," Ph.D. dissertation, Fac. Sci. (FNWI), Korteweg-de Vries Inst. Math. (KdVI), Univ. Amsterdam, Amsterdam, The Netherlands, 2005. [Online]. Available: <https://dare.uva.nl/search?identifier=032593d8-9f15-4f57-a288-627fb55effae>
- [53] S. Polak, "New methods in coding theory: Error-correcting codes and the Shannon capacity," Ph.D. dissertation, Fac. Sci. (FNWI), Korteweg-de Vries Inst. Math. (KdVI), Univ. Amsterdam, Amsterdam, The Netherlands, Sep. 2019. [Online]. Available: <http://arxiv.org/abs/2005.02945>
- [54] B. Litjens, S. Polak, and A. Schrijver, "Semidefinite bounds for non-binary codes based on quadruples," *Des., Codes Cryptogr.*, vol. 84, nos. 1–2, pp. 87–100, Jul. 2017.
- [55] L. C. Araujo, J. P. H. Sansão, and A. S. Vale-Cardoso, "Fast computation of multinomial coefficients," *Numer. Algorithms*, vol. 88, no. 2, pp. 837–851, Oct. 2021.
- [56] H. Jiang, T. Kathuria, Y. T. Lee, S. Padmanabhan, and Z. Song, "A faster interior point method for semidefinite programming," in *Proc. IEEE 61st Annu. Symp. Found. Comput. Sci. (FOCS)*, Nov. 2020, pp. 910–918.
- [57] M. Tomamichel, R. Colbeck, and R. Renner, "A fully quantum asymptotic equipartition property," *IEEE Trans. Inf. Theory*, vol. 55, no. 12, pp. 5840–5847, Dec. 2009.
- [58] M. Tomamichel, "A framework for non-asymptotic quantum information theory," Ph.D. thesis, Dept. Phys., ETH Zürich, Zürich, Switzerland, 2012.
- [59] M. Tomamichel and M. Hayashi, "A hierarchy of information quantities for finite block length analysis of quantum tasks," *IEEE Trans. Inf. Theory*, vol. 59, no. 11, pp. 7693–7710, Nov. 2013.
- [60] K. Fang and H. Fawzi, "Geometric Rényi divergence and its applications in quantum channel capacities," *Commun. Math. Phys.*, vol. 384, no. 3, pp. 1615–1677, Jun. 2021.
- [61] K. Li, "Second-order asymptotics for quantum hypothesis testing," *Ann. Statist.*, vol. 42, no. 1, Feb. 2014.
- [62] H. Fawzi and O. Fawzi, "Defining quantum divergences via convex optimization," *Quantum*, vol. 5, p. 387, Jan. 2021.
- [63] K. Fang, G. Gour, and X. Wang, "Towards the ultimate limits of quantum channel discrimination," 2021, *arXiv:2110.14842*.
- [64] M. Berta et al., "On a gap in the proof of the generalised quantum Stein's lemma and its consequences for the reversibility of quantum resources," *Quantum*, vol. 7, p. 1103, Sep. 2023.

Bjarne Bergh received the bachelor's and master's degrees in physics from the University of Heidelberg, Germany, in 2018 and 2020, respectively, and the master's degree in applied mathematics from the University of Cambridge, U.K., in 2019, where he is currently pursuing the Ph.D. degree working on problems in quantum information theory.

Nilanjana Datta received the Ph.D. degree in mathematical physics from ETH Zürich, Switzerland, in 1996. From 1997 to 2000, she was a Post-Doctoral Researcher with the Dublin Institute of Advanced Studies, CNRS Marseille, and EPFL, Lausanne. In 2001, she joined the University of Cambridge, U.K., as a Lecturer in mathematics with the Pembroke College and a member of the Statistical Laboratory, Centre for Mathematical Sciences. She is currently a Professor in quantum information theory with the Department of Applied Mathematics and Theoretical Physics, University of Cambridge, and a fellow of the Pembroke College. Her research interests include quantum information theory and mathematical physics.

Robert Salzmann received the Ph.D. degree in mathematics from the University of Cambridge, U.K., in 2023. He is currently a Post-Doctoral Researcher with École Normale Supérieure de Lyon working on problems in quantum information theory and mathematical physics.

Mark M. Wilde (Fellow, IEEE) received the Ph.D. degree in electrical engineering from the University of Southern California, Los Angeles, CA, USA. He is currently an Associate Professor in electrical and computer engineering with Cornell University. His current research interests are in quantum Shannon theory, quantum computation, quantum optical communication, quantum computational complexity theory, and quantum error correction. He was a recipient of the National Science Foundation Career Development Award. He was also a co-recipient of the 2018 AHP-Birkhäuser Prize, awarded to "The most remarkable contribution" published in the journal *Annales Henri Poincaré*. He is an Outstanding Referee of the American Physical Society.