

Predicting quantum channels over general product distributions

Sitan Chen^{*} Jaume de Dios Pont[†] Jun-Ting Hsieh[‡] Hsin-Yuan Huang[§]
 Jane Lange[¶] Jerry Li^{||}

Abstract

We investigate the problem of predicting the output behavior of unknown quantum channels. Given query access to an n -qubit channel $\mathcal{E}$ and an observable $\mathcal{O}$, we aim to learn the mapping

$$\rho \mapsto \text{Tr}(\mathcal{O}\mathcal{E}[\rho])$$

to within a small error for most ρ sampled from a distribution $\mathcal{D}$. Previously, Huang, Chen, and Preskill [HCP23] proved a surprising result that even if $\mathcal{E}$ is arbitrary, this task can be solved in time roughly $n^{O(\log(1/\varepsilon))}$, where ε is the target prediction error. However, their guarantee applied only to input distributions $\mathcal{D}$ invariant under all single-qubit Clifford gates, and their algorithm fails for important cases such as general product distributions over product states ρ .

In this work, we propose a new approach that achieves accurate prediction over essentially any product distribution $\mathcal{D}$, provided it is not “classical” in which case there is a trivial exponential lower bound. Our method employs a “biased Pauli analysis,” analogous to classical biased Fourier analysis. Implementing this approach requires overcoming several challenges unique to the quantum setting, including the lack of a basis with appropriate orthogonality properties. The techniques we develop to address these issues may have broader applications in quantum information.

1 Introduction

When is it possible to learn to predict the outputs of a quantum channel $\mathcal{E}$? Such questions arise naturally in a variety of settings, such as the experimental study of complex quantum dynamics [HBC⁺22, HKP21], and in fast-forwarding simulations of Hamiltonian evolutions [CHI⁺20, GHC⁺24]. However, in the worst case this problem is intractable, as it generalizes the classical problem of learning an arbitrary Boolean function over the uniform distribution from black-box access. To circumvent this, our goal is to understand families of natural restrictions on the problem under which efficient estimation is possible.

One way to avoid this exponential scaling would be to posit further structure on the channel, e.g. by assuming it is given by a shallow quantum circuit [NPVY23, HLB⁺24] or a structured Pauli channel [HFW20, FO21, HYF21, VDBMKT23, ADGP24]. However, there are settings where the evolutions may be quite complicated — e.g. the channel might correspond to the time evolution of an evaporating black

^{*}Harvard University. Email: sitan@seas.harvard.edu.

[†]ETH Zurich. Email: jaume.dediospont@math.ethz.ch.

[‡]Carnegie Mellon University. Email: juntingh@cs.cmu.edu.

[§]Google Quantum AI, Caltech. Email: hsinyuan@google.com, hsinyuan@caltech.edu.

[¶]MIT. Email: jlange@mit.edu.

^{||}Microsoft Research. Email: jerryzli@u.washington.edu.

hole [HP07, HP19, PSSY22, YE23] — and where it is advantageous to avoid such strong structural assumptions on the underlying channel.

Recently, [HCP23] considered an alternative workaround in which one only attempts to learn a complicated n -qubit channel in an *average-case* sense. Given query access to $\mathcal{E}$, and given an observable $\mathcal{O}$, the goal is to learn the mapping

$$\rho \mapsto \text{Tr}(\mathcal{O} \mathcal{E}[\rho])$$

accurately on average over input states ρ drawn from some n -qubit distribution $\mathcal{D}$, rather than over worst-case input states. The authors of [HCP23] came to the surprising conclusion that this average-case task is tractable even for *arbitrary* channels $\mathcal{E}$, provided $\mathcal{D}$ comes from a certain class of “locally flat” distributions. Their key observation was that the Heisenberg-evolved observable $\mathcal{E}^\dagger[\mathcal{O}]$ admits a *low-degree approximation* in the Pauli basis, where the quality of approximation is defined in an average-case sense over input state ρ .

Another interesting feature of this result is that their learning algorithm only needs to query $\mathcal{E}$ on random product states, regardless of the choice of locally flat distribution $\mathcal{D}$. This is both an advantage and a shortcoming. On one hand, if one is certain that the states ρ one wants to predict on are samples from a locally flat distribution, no further information about $\mathcal{D}$ is needed to implement the learning protocol in [HCP23]. On the other hand, locally flat distributions are quite specialized: they are constrained to be invariant under any single-qubit Clifford gate. In particular, almost all product distributions over product states fall outside this class. Worse yet, the general approach of low-degree approximation in the Pauli basis can be shown to fail when local flatness does not hold (see [Section 5.2](#)). We therefore ask:

Are there more general families of distributions $\mathcal{D}$ under which one can learn to predict arbitrary quantum dynamics?

Identifying rich settings where it is possible to characterize the average-case behavior of such dynamics, while making minimal assumptions on the dynamics, is of intense practical interest. Unfortunately, our understanding of this remains limited: even for general product distributions, known techniques break down. In this work we take an important first step towards this goal by completely characterizing the complexity of learning to predict arbitrary quantum dynamics in the product setting. Informally stated, our main result is that learning is possible *so long as the distribution is not classical*. That is, for this problem there is a “blessing of quantum-ness”: as long as the distribution displays any quantitative level of quantum behavior, there is an efficient algorithm for predicting arbitrary quantum dynamics under this distribution.

More formally, note that if D is the uniform distribution over the computational basis states $|0\rangle$ and $|1\rangle$, then the task of predicting $\text{Tr}(\mathcal{O} \mathcal{E}[\rho])$ on average over $\rho \sim \mathcal{D} \triangleq D^{\otimes n}$ for an arbitrary channel $\mathcal{E}$ is equivalent to the task of learning an arbitrary Boolean function from random labeled examples, which trivially requires exponentially many samples. This logic naturally extends to any “two-point” distribution in which D is supported on two diametrically opposite points on the Bloch sphere. Note that any such distribution, up to a rotation, is an embedding of a classical distribution onto the Bloch sphere.

A natural way of quantifying closeness to such distributions is in terms of the second moment matrix $\mathcal{S} \in \mathbb{R}^{3 \times 3}$ of the distribution D , when D is viewed as a distribution over the Bloch sphere (see [Section 2.1](#) for formal definitions). We refer to this matrix as the *Pauli second moment matrix* of D . For the purposes of this discussion, the key property of this matrix is that $\|\mathcal{S}\|_{\text{op}} \leq 1$ for all D , and moreover, $\|\mathcal{S}\|_{\text{op}} = 1$

if and only if D is one of the aforementioned two-point distributions. With this, we can now state our main result:

Theorem 1.1 (Learning an unknown quantum channel). *Let $\varepsilon, \delta, \eta \in (0, 1)$. Let D be an unknown distribution over the Bloch sphere with Pauli second moment matrix $\mathcal{S}$ such that $\|\mathcal{S}\|_{\text{op}} \leq 1 - \eta$. Let $\mathcal{E}$ be an unknown n -qubit quantum channel, and let $\mathcal{O}$ be a known n -qubit observable. There exists an algorithm with time and sample complexity $\min(2^{O(n)}/\varepsilon^2, n^{O(\log(1/\varepsilon)/\log(1/(1-\eta)))}) \cdot \log(1/\delta)$ that outputs an efficiently computable map f' such that*

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [(\text{Tr}(\mathcal{O}\mathcal{E}[\rho]) - \text{Tr}(f'(\rho)))^2] \leq \varepsilon$$

with probability at least $1 - \delta$.

Note that the only condition on D we require is a quantitative bound on the spectral norm of its Pauli second moment matrix. In other words, so long as the distribution D is far from any two-point distribution, i.e., it is far from any classical distribution, we demonstrate that there is an efficient algorithm for learning to predict general quantum dynamics under this distribution. Previously it was only known how to achieve the above guarantee in the special case where D has mean zero. Indeed, as soon as one deviates from the mean zero case, the Pauli decomposition approach of [HCP23] breaks down. In contrast, our guarantee works for any product distribution whose marginal second moment matrices have operator norm bounded away from 1.

Remark 1.2. *We note that our techniques generalize to the case where the distribution is the product of different distributions over qubits, so long as each distribution has second moment with operator norm bounded by $1 - \eta$. However, for readability we will primarily focus on the case where all of the distributions are the same. See Remark 4.7 for a discussion of how to easily generalize our techniques to this setting.*

Beyond low-degree concentration in an orthonormal basis. Here we briefly highlight the key conceptual novelties of our analysis, which may be of independent interest. We begin by recalling the analysis in [HCP23] in greater detail. They considered the decomposition of $\mathcal{O}' \triangleq \mathcal{E}^\dagger[\mathcal{O}]$ into the basis of n -qubit Pauli operators, i.e. $\mathcal{O}' = \sum_{P \in \{I, X, Y, Z\}^n} \alpha_P \cdot P$ and argued that this is well-approximated by the *low-degree truncation* $\mathcal{O}'_{\text{low}} = \sum_{|P| < t} \alpha_P \cdot P$. This can be readily seen from the following calculation. By rotating the distribution D , we may assume the covariance is diagonal, with entries bounded by $1 - \eta$. Then the error achieved by the low-degree truncation is given by

$$\mathbb{E}[\text{Tr}((\mathcal{O}' - \mathcal{O}'_{\text{low}})\rho)^2] = \mathbb{E}\left[\left(\sum_{|P| \geq t} \alpha_P \cdot \text{Tr}(P\rho)\right)^2\right] \leq \sum_{|P| \geq t} (1 - \eta)^{|P|} \cdot \alpha_P^2 \leq (1 - \eta)^t \cdot \frac{1}{2^n} \|\mathcal{O}'\|_F^2 \leq (1 - \eta)^t,$$

where the second step follows from the fact that $\mathbb{E}[\text{Tr}(P\rho) \text{Tr}(Q\rho)] = 0$ if $P \neq Q$ and is at most $(1 - \eta)^{|P|}$ if $P = Q$ (since D is mean zero and its covariance is diagonal with entries bounded by $1 - \eta$), and the last step follows by the assumption that $\|\mathcal{O}'\|_{\text{op}} \leq 1$.

Note that when D is not mean zero, this step breaks, and we do not have this nice exponential decay in t . In fact, in Section 5.2 we construct examples of operators which are not well-approximated by their low-degree truncations in the Pauli basis when D has mean bounded away from zero.

A natural attempt at a workaround would be to change the basis under which we truncate. At least classically, biased product distributions over the Boolean hypercube still admit suitable orthonormal bases of functions, namely the *biased Fourier characters*. As we show in Section 3, this idea can be used to give the following learning guarantee in the classical case where D is only supported along the Z direction in the Bloch sphere:

Theorem 1.3 (PAC learning over a concentrated product distribution). *Let $\varepsilon, \delta, \eta \in (0, 1)$. Let D be an unknown distribution over the interval $[-(1 - \eta), 1 - \eta]$. Let $f : [-1, 1]^n \rightarrow [-1, 1]$ be an unknown bounded, multilinear function. There exists an algorithm with time and sample complexity $n^{O(\log(1/\varepsilon)/\log(1/(1-\eta)))}$. $\log(1/\delta)$ that outputs a hypothesis f' such that*

$$\mathbb{E}_{x \sim D^{\otimes n}} [(f(x) - f'(x))^2] \leq \varepsilon$$

with probability at least $1 - \delta$.

Unfortunately, when we move beyond the classical setting, the picture becomes trickier. In particular, it is not immediately clear what the suitable analogue of the biased Fourier basis should be in the quantum setting. We could certainly try to consider single-qubit operators of the form $\tilde{P} = P - \mu_P \cdot I$ for $P \in \{X, Y, Z\}$, where μ_P denotes the P -th coordinate of the mean of D regarded as a distribution over the Bloch sphere. We could then extend naturally to give a basis over n qubits, and the functions $\rho \mapsto \text{Tr}(\tilde{P}\rho)$ would by design be orthogonal to each other with respect to the distribution $D^{\otimes n}$. Writing $\mathcal{O}' = \sum_P \tilde{\alpha}_P \cdot \tilde{P}$ and defining $\tilde{\mathcal{O}}_{\text{low}} \triangleq \sum_{|P| < t} \tilde{\alpha}_P \cdot \tilde{P}$, we can mimic the calculation above and obtain

$$\mathbb{E}[\text{Tr}((\mathcal{O}' - \tilde{\mathcal{O}}_{\text{low}})\rho)^2] = \mathbb{E}\left[\left(\sum_{|P| \geq t} \tilde{\alpha}_P \cdot \text{Tr}(\tilde{P}\rho)\right)^2\right] \leq (1 - \eta)^{|P|} \sum_{|P| \geq t} \tilde{\alpha}_P^2.$$

Unfortunately, at this juncture the above naive approach hits a snag. In the mean zero case, we could easily relate $\sum_P \alpha_P^2$ to $\frac{1}{2^n} \|\mathcal{O}'\|_F^2$ because of a fortuitous peculiarity of the mean-zero setting. In that setting, we implicitly exploited both that the Pauli operators P are orthogonal to each other with respect to the trace inner product, and also that the functions $\rho \mapsto \text{Tr}(P\rho)$ are orthogonal to each other with respect to $D^{\otimes n}$. In the above approach for nonzero mean, we achieved the latter condition by shifting the Pauli operators to define $\tilde{P}$, but these shifted operators are no longer orthogonal to each other with respect to the trace inner product.

Circumventing this issue is the technical heart of our proof. As we will see in [Section 4](#), several key technical moves are needed. First, instead of assuming that the covariance of D is diagonalized, we will fix a rotation that simplifies the mean, $\mathbb{E}[\rho]$. Second, instead of shifting the basis operators $\{X, Y, Z\}$ so that the resulting functions $\rho \mapsto \text{Tr}(\tilde{P}\rho)$ are orthogonal with respect to $D^{\otimes n}$, we shift them so that $\mathbb{E}[\rho]$ is orthogonal to them, and define $\mathcal{O}'_{\text{low}}$ by truncating in this new basis instead. Finally, instead of directly bounding the truncation error $\mathbb{E}[\text{Tr}((\mathcal{O}' - \mathcal{O}'_{\text{low}})\rho)^2]$ using the above sequence of steps, we crucially relate it to the quantity

$$\text{Tr}((\mathcal{O}')^2 \mathbb{E}[\rho])$$

in order to establish exponential decay. Note that $\text{Tr}((\mathcal{O}')^2 \mathbb{E}[\rho]) \leq \|\mathcal{O}'\|_{\text{op}}^2 \cdot \text{Tr}(\mathbb{E}[\rho]) \leq 1$. To our knowledge, all three of these components are new to our analysis. We leave it as an intriguing open question to find other applications of these ingredients to domains where “biased Pauli analysis” arises.

Remark 1.4 (Predicting multiple observables). Just as in [\[HCP23\]](#), we can also easily extend our guarantee to the setting where we wish to learn the joint mapping

$$(\rho, \mathcal{O}) \mapsto \text{Tr}(\mathcal{O} \mathcal{E}[\rho]). \quad (1)$$

This is the natural channel learning analogue of the question of classical shadows for state learning [\[HKP20\]](#) – recall that in the latter setting, one would like to perform measurements on copies of ρ and obviously produce a classical description of the state that can then be used to compute some collection of observable values. We sketch the argument for extending to learning the joint mapping in Eq. (1) in [Remark 4.6](#).

Impossibility for general concentrated distributions. It is natural to wonder to what extent our results can be generalized, especially to states that are entangled. Could it be that all one needs is some kind of global covariance bound? Unfortunately, we show in [Section 5](#) that even in the classical setting, this is not the case. Since classical distributions can be encoded by distributions over qubits, this implies hardness for learning in the quantum setting as well.

Theorem 1.5 (Hardness of learning over general concentrated distributions). *There exists a distribution D over $[-(1 - \eta), 1 - \eta]^n$ and a concept class C such that no algorithm PAC-learns the class C over D in subexponential time.*

There is a wide spectrum of distributional assumptions that interpolates between fully product distributions and general concentrated distributions — for instance, products of k -dimensional qudit distributions, output states of small quantum circuits, or distributions over negatively associated variables. As discussed earlier, our understanding of when it is possible to predict the average-case behavior of arbitrary quantum dynamics is still nascent, and understanding learnability with respect to these more expressive distributional assumptions remains an important open question.

Organization. In [Section 2](#), we state some preliminaries. In [Section 3](#), we prove [Theorem 1.3](#), which is the classical setting and can be viewed as a warm-up to the quantum setting. In [Section 4](#), we prove [Theorem 1.1](#), our main result. Finally, in [Section 5](#), we show impossibility results including [Theorem 1.5](#) and the failure of low-degree truncation in the standard Pauli basis.

1.1 Related work

Our work is part of a growing literature bridging classical computational learning theory and its quantum counterpart. Its motivation can be thought of as coming from the general area of quantum process tomography [[MRL08](#)], but as this is an incredibly extensive research direction, here we only focus our attention on surveying directly relevant works.

Quantum analysis of Boolean functions. In [[MO08](#)], it was proposed to study Pauli decompositions of *Hermitian* unitaries as the natural quantum analogue of Boolean functions. One notable follow-up work [[RWZ22](#)] proved various quantum versions of classical Fourier analytic results like Talagrand’s variance inequality and the KKL theorem in this setting, and also obtained corollaries about learning Hermitian unitaries in Frobenius norm given oracle access (see also [[CNY23](#), [BY23](#)] for the non-Hermitian case). Recently, [[NPVY23](#)] considered the Pauli spectrum of the *Choi representation* of quantum channels and proved low-degree concentration for channels implemented by QAC⁰. One technical difference with our work is that these notions of Pauli decomposition are specific to the channel, whereas the object whose Pauli decomposition we consider is specific to the Heisenberg-evolved operator $\mathcal{E}^\dagger[\mathcal{O}]$. Additionally, we note that all of the above mentioned works focus on questions more akin to learning a full description of the channel and thus are inherently tied to channels with specific structure.

In contrast, our focus is on learning certain *properties* of the channel, and only in an average-case sense over input states. As mentioned previously, this specific question was first studied in [[HCP23](#)]. There have been two direct follow-up works to this paper which are somewhat orthogonal to the thrust of our contributions. The first [[VZ23](#)] establishes refined versions of the so-called *non-commutative Bohnenblust-Hille inequality* which was developed and leveraged by [[HCP23](#)] to obtain *logarithmic* sample complexity bounds. In this work, we did not pursue this avenue of improvement but leave it

as an interesting open question to improve our sample complexity guarantees accordingly. The second follow-up [KSVZ23] to [HCP23] studies the natural qudit generalization of the original question where the distribution over qudits is similarly closed under a certain family of single-site transformations.

Finally, we note the recent work of [ADGP24] which studied the learnability of quantum channels with only low-degree Pauli coefficients. Their focus is incomparable to ours as they target a stronger metric for learning, namely ℓ_2 -distance for channels, but need to make a strong assumption on the complexity of the channel being learned. In contrast, we target a weaker metric, namely average-case error for predicting observables, but our guarantee applies for arbitrary channels.

Classical low-degree learning. The general technique of low-degree approximation in classical learning theory is too prevalent to do full justice to in this section. This idea of learning Boolean functions by approximating their low-degree Fourier truncation was first introduced in the seminal work of [LMN93]. Fourier-analytic techniques have been used to obtain new classical learning results for various concept classes like decision trees [KM91], linear threshold functions [KKMS08], Boolean formulas [KS01], low-degree polynomials [EI22], and more.

While Fourier analysis over biased distributions dates back to early work of Margulis and Russo [Mar74, Rus81, Rus82], it was first applied in a learning-theoretic context in [FJS91], extending the aforementioned result of [LMN93].

2 Preliminaries

2.1 Bloch sphere and Pauli covariance matrices

The Pauli matrices I, X, Y, Z provide the basis for 2×2 Hermitian matrices. This is captured by the following standard fact on expanding a single-qubit state using Pauli matrices.

Fact 2.1 (Pauli expansion of states). *Any single-qubit mixed state ρ can be written as*

$$\rho = \frac{1}{2}(I + \alpha_x X + \alpha_y Y + \alpha_z Z),$$

where $\vec{\alpha} \in \mathbb{R}^3$, $\|\vec{\alpha}\|_2 \leq 1$, and X, Y, Z are the standard Pauli matrices:

$$X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

The set of all such $\vec{\alpha}$ of unit norm is the Bloch sphere.

Any single-qubit distribution D can be viewed as a distribution over the Bloch sphere. We use $\mathbb{E}_D[\rho] \in \mathbb{C}^{2 \times 2}$ and $\vec{\mu} \in \mathbb{R}^3$ to refer to the expected state and the expected Bloch vector respectively.

By taking tensor products of Pauli matrices, we obtain the collection of 4^n Pauli observables $\{I, X, Y, Z\}^{\otimes n}$, which form a basis for the space of $2^n \times 2^n$ Hermitian matrices:

Fact 2.2 (Pauli expansion of observables). *Let $\mathcal{O}$ be an n -qubit observable. Then $\mathcal{O}$ can be written in the following form:*

$$\mathcal{O} = \sum_{P \in \{I, X, Y, Z\}^{\otimes n}} \widehat{\mathcal{O}}(P) \cdot P,$$

where $\widehat{\mathcal{O}}(P) \triangleq \text{Tr}(\mathcal{O}P)/2^n$.

Next, we define the Pauli covariance and second moment matrices associated to any distribution D over the single-qubit Bloch sphere.

Definition 2.1 (Pauli covariance and second moment matrices). *Let D be a distribution over the Bloch sphere. We will associate with D the second moment matrix $S \in \mathbb{R}^{3 \times 3}$ and covariance matrix $\Sigma \in \mathbb{R}^{3 \times 3}$, indexed by the non-identity Pauli components X, Y , and Z . We define S such that $S_{P,Q} = \mathbb{E}_{\rho \sim D} [\text{Tr}(P\rho)\text{Tr}(Q\rho)]$, and Σ such that $\Sigma_{P,Q} = \mathbb{E}_{\rho \sim D} [\text{Tr}(P\rho)\text{Tr}(Q\rho) - \text{Tr}(P \mathbb{E}[\rho])\text{Tr}(Q \mathbb{E}[\rho])]$.*

The following is a consequence of [Fact 2.1](#):

Fact 2.3. *For any distribution over the Bloch sphere, the Pauli second moment matrix S satisfies $\text{Tr}(S) = 1$.*

2.2 Access model

In this paper we consider the following, standard access model, see e.g. [\[HCP23\]](#). We assume that we can interact with the unknown channel $\mathcal{E}$ by preparing any input state, passing it through $\mathcal{E}$, and performing a measurement on the output. Additionally, we are given access to training examples from some distribution $\mathcal{D} = D^{\otimes n}$ over product states, and their corresponding *classical descriptions*. Because product states over n qubits can be efficiently represent efficiently using $O(n)$ bits, the training set can be stored efficiently on classical computers. The standard approach to represent a product state on a classical computer is as follows. For each state $\rho = \otimes_{i=1}^n |\psi_i\rangle$ sampled from $\mathcal{D}$, the classical description can be given by their 1-qubit Pauli expectation values: $\text{Tr}(P|\psi_i\rangle\langle\psi_i|)$ for all $i \in [n]$ ranging over each qubit and $P \in \{X, Y, Z\}$.

Given these classical samples from $\mathcal{D}$ and the ability to query $\mathcal{E}$, the learning goal is to produce a hypothesis f' which takes as input the classical description of a product state ρ and outputs an estimate for $\text{Tr}(\mathcal{O} \mathcal{E}[\rho])$. Formally, we want this hypothesis to have small test loss in the sense that $\mathbb{E}_{\rho \sim \mathcal{D}} [(\text{Tr}(\mathcal{O} \mathcal{E}[\rho]) - \text{Tr}(f'(\rho)))^2] \leq \epsilon$ with probability at least $1 - \delta$ over the randomness of the learning algorithm and the training examples from $\mathcal{D}$.

2.3 Generalization bounds for learning

For our learning protocol, we will use the following elementary results about linear and polynomial regression:

Fact 2.4 (Rademacher complexity generalization bound [\[MRT18\]](#)). *Let $\mathcal{F}$ be the class of bounded linear functions $[-1, 1]^d \rightarrow [-B, B]$, and let ℓ be a loss function with Lipschitz constant L and a uniform upper bound of c . With probability $1 - \delta$ over the choice of a training set S of size m drawn i.i.d. from distribution $\mathcal{D}$,*

$$\mathbb{E}_{(x,y) \sim \mathcal{D}} [\ell(f(x), y)] \leq \mathbb{E}_{(x,y) \sim S} [\ell(f(x), y)] + 4LB\sqrt{\frac{d}{m}} + 2c\sqrt{\frac{\log 1/\delta}{2m}}.$$

Corollary 2.5. *Let f be a function that is ϵ -close to a degree- $\leq d$ polynomial f^* :*

$$\mathbb{E}_{x \sim \mathcal{D}} [(f(x) - f^*(x))^2] \leq \epsilon.$$

Then linear regression over the set of degree- d polynomials with coefficients in $[-1, 1]$ has time and sample complexity $\text{poly}(n^d, \log 1/\epsilon) \cdot \log 1/\delta$ and finds h such that

$$\mathbb{E}_{x \sim \mathcal{D}} [(f(x) - h(x))^2] \leq O(\epsilon)$$

with probability $1 - \delta$.

3 Warm-up: the classical case

In this section we will prove [Theorem 1.3](#), which is a special case of [Theorem 1.1](#) where the distribution is classical, i.e. supported only on the Z component.

The key ingredient in the proof is to show that for any function which is L^2 -integrable with respect to a product distribution over $[-(1-\eta), 1-\eta]^n$ and whose extension to the hypercube is bounded, the function admits a “low-degree” approximation under an appropriate orthonormal basis. Roughly, the intuition is that the space of linear functions over a distribution D on $[-(1-\eta), 1-\eta]$ has an orthonormal basis which is a $(1-\eta)$ -scaling of a basis for a distribution on $\{-1, 1\}$. Therefore, the space of multilinear functions over the corresponding product distribution $D^{\otimes n}$ has a basis whose degree- d components are scaled by $(1-\eta)^d$. This, combined with the assumption that the function is bounded on the hypercube, allows us to conclude that the contribution of the degree- d component to the variance of f over $D^{\otimes n}$ is at most $(1-\eta)^{2d}$.

We prove this structural result in [Section 3.1](#) and conclude the proof of [Theorem 1.3](#) in [Section 3.2](#).

3.1 Existence of low-degree approximation

We first review some basic facts about classical biased Fourier analysis. For a more extensive overview of this topic, we refer the reader to [\[O’D14, Chapter 8\]](#).

Given a measure μ , we let $L^2(\mu)$ denote the space of L^2 -integrable functions with respect to μ .

Fact 3.1 (Biased Fourier basis). *Let D be a distribution over $\{-1, 1\}$ with mean $\mu \in (-1, 1)$. Given $f \in L_2(D^{\otimes n})$, the μ -biased Fourier expansion of $f : \{-1, 1\}^n \rightarrow \mathbb{R}$ is*

$$f(x) = \sum_{S \subseteq [n]} \widehat{f}(S) \phi_S(x),$$

where $\phi_S(x) = \prod_{i \in S} \frac{x_i - \mu}{\sqrt{1 - \mu^2}}$ and $\widehat{f}(S) = \mathbb{E}_{x \sim D^{\otimes n}} [\phi_S(x) f(x)]$.

The functions ϕ_S provide an orthonormal basis for the space of functions $L^2(D^{\otimes n})$, where D is the distribution over $\{1, -1\}$ with mean μ . We can naturally extend this to arbitrary product distributions over $\mathbb{R}^d$ as follows:

Fact 3.2 (Basis for an arbitrary product distribution). *Let D be a distribution over $\mathbb{R}$ with mean μ and variance $\sigma^2 > 0$. Then $\{1, \frac{x - \mu}{\sigma}\}$ is an orthonormal basis for $L^2(D)$, and thus $\{1, \frac{x - \mu}{\sigma}\}^{\otimes n}$ is an orthonormal basis for $L^2(D^{\otimes n})$.*

The orthonormality of the basis immediately implies the following simple fact:

Fact 3.3 (Parseval’s Theorem). *For any function f expressed as $f(x) = \sum_{S \subseteq [n]} \widehat{f}(S) \phi_S(x)$, we have $\mathbb{E}_{x \sim D^{\otimes n}} [f(x)^2] = \sum_{S \subseteq [n]} \widehat{f}(S)^2$.*

The following is the crucial structural result in the classical setting that gives rise to [Theorem 1.3](#). Roughly speaking, it ensures that for the “concentrated product distributions” D considered therein, any bounded multilinear function has decaying coefficients when expanded in the orthonormal basis for $L^2(D^{\otimes n})$.

Lemma 3.4. Let $f : [-1, 1]^n \rightarrow [-1, 1]$ be a multilinear function and let D be a distribution over $\mathbb{R}$ with mean μ and variance $\sigma^2 < 1 - \mu^2$. Then there exists a function $f^{\leq d}$ such that

$$\mathbb{E}_{x \sim D^{\otimes n}} [(f(x) - f^{\leq d}(x))^2] \leq \left(\frac{\sigma^2}{1 - \mu^2} \right)^d.$$

Proof. Let f be expressed in the basis $B_{\text{hypercube}} = \{1, \frac{x-\mu}{\sqrt{1-\mu^2}}\}^{\otimes n}$, i.e. as

$$f(x) = \sum_{S \subseteq [n]} \hat{f}(S) \psi_S(x)$$

where $\psi_S = \prod_{i \in S} \frac{x_i - \mu}{\sqrt{1 - \mu^2}}$. Note that $\{\psi_S\}_{S \subseteq [n]}$ is orthonormal with respect to the distribution $\tilde{D}^{\otimes n}$ where $\tilde{D}$ is supported on $\{\pm 1\}$ with mean μ . Since $|f(x)| \leq 1$ for $x \in [-1, 1]^n$, it follows that

$$1 \geq \mathbb{E}_{x \sim \tilde{D}^{\otimes n}} [f(x)^2] = \sum_{S \subseteq [n]} \hat{f}(S)^2$$

via [Fact 3.3](#).

Now, consider the basis $\phi_S := \prod_{i \in S} \frac{x_i - \mu}{\sigma} = \left(\frac{\sqrt{1 - \mu^2}}{\sigma} \right)^{|S|} \cdot \psi_S$. By [Fact 3.2](#), we know that $\{\phi_S\}_{S \subseteq [n]}$ is orthonormal with respect to D . Let $f^{> d} := \sum_{|S| > d} \hat{f}(S) \psi_S$. We have

$$\begin{aligned} \mathbb{E}_{x \sim D^{\otimes n}} [f^{> d}(x)^2] &= \mathbb{E}_{x \sim D^{\otimes n}} \left[\left(\sum_{|S| > d} \hat{f}(S) \psi_S \right)^2 \right] \\ &= \mathbb{E}_{x \sim D^{\otimes n}} \left[\left(\sum_{|S| > d} \hat{f}(S) \left(\frac{\sigma}{\sqrt{1 - \mu^2}} \right)^{|S|} \phi_S \right)^2 \right] \\ &= \sum_{|S| > d} \hat{f}(S)^2 \left(\frac{\sigma^2}{1 - \mu^2} \right)^{|S|} \\ &\leq \left(\frac{\sigma^2}{1 - \mu^2} \right)^d \sum_{|S| > d} \hat{f}(S)^2, \end{aligned}$$

again using [Fact 3.3](#). Since $\sum_S \hat{f}(S)^2 \leq 1$, we have

$$\mathbb{E}_{x \sim D^{\otimes n}} [(f(x) - f^{\leq d}(x))^2] \leq \left(\frac{\sigma^2}{1 - \mu^2} \right)^d$$

as desired. □

3.2 Sample complexity and error analysis

In light of [Lemma 3.4](#), the proof of [Theorem 1.3](#) is straightforward given the following elementary fact:

Fact 3.5 (Bernoulli maximizes variance). Let D be a distribution over the interval $[-(1 - \eta), 1 - \eta]$ with mean μ . Then $\text{Var}_{x \sim D}(x) \leq (1 - \eta)^2 (1 - \mu^2)$.

We can now conclude the proof of [Theorem 1.3](#):

Proof of Theorem 1.3. By [Fact 3.5](#), we have $\text{Var}_D(x) \leq (1 - \eta)^2(1 - \mu^2)$. Then [Lemma 3.4](#) gives us a degree- d approximation $f^{\leq d}$ to f such that

$$\mathbb{E}_{x \sim D^{\otimes n}} [(f(x) - f^{\leq d}(x))^2] \leq (1 - \eta)^{2d}.$$

Taking $d := \log(1/\epsilon)/\log(1/(1 - \eta))$ gives an approximation with error $\leq \epsilon$. Then by [Corollary 2.5](#), linear regression on the space of polynomials of degree $\log(1/\epsilon)/\log(1/(1 - \eta))$ finds an $O(\epsilon)$ -error hypothesis in $n^{O(\log(1/\epsilon)/\log(1/(1 - \eta)))} \cdot \log(1/\delta)$ time and samples. $\square$

Remark 3.6. In the classical setting linear regression is not required, as we can estimate the mean, and the coefficients satisfy

$$\hat{f}(S) = \mathbb{E}_{x \sim D^{\otimes n}} [f(x)\phi_S(x)],$$

so they can be estimated directly. We give the guarantee in terms of linear regression because the approach of directly estimating the coefficients does not generalize to the quantum setting.

4 Learning an unknown quantum channel

In this section we will prove [Theorem 1.1](#). First we will show that under any product distribution with second moment S such that $\|S\|_{\text{op}} \leq 1 - \eta$ for some $\eta \in (0, 1)$, every observable has a low-degree approximation. A distribution has $\|S\|_{\text{op}} = 1$ only if it is effectively a classical distribution; i.e. it is supported on two antipodal points in the Bloch sphere. So we are showing that any product distribution which is “spread out” within the Bloch sphere behaves well with low-degree approximation.

To do this, we cannot use exactly the same argument as in [Section 3](#), because there is not necessarily an orthonormal basis for our product distribution $D^{\otimes n}$ that is a “stretched” basis for some other distribution over the Bloch sphere. Instead, we compare the variance of the observable under $D^{\otimes n}$ to the quantity $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^2 \rho)]$. This allows us to use the boundedness of $\mathcal{O}$ to derive bounds for the contribution of the degree- d part to the variance of $\mathcal{O}$ under $D^{\otimes n}$. The learning algorithm will find the low-degree approximation by linear regression over the degree- $\log(1/\epsilon)$ Pauli coefficients. The notion of low-degreeness will be with respect to a basis adapted to $D^{\otimes n}$.

Definition 4.1. Let D be a distribution over the Bloch sphere. Let $U^\dagger A U$ be the eigendecomposition of $\bar{\rho} = \mathbb{E}_D[\rho]$. Let $\tilde{X}, \tilde{Y}, \tilde{Z} := U^\dagger X U, U^\dagger Y U, \frac{U^\dagger Z U - \text{Tr}(\bar{\rho} U^\dagger Z U) I}{\sqrt{1 - \text{Tr}(\bar{\rho} U^\dagger Z U)^2}}$.

Let $B = \{I, \tilde{X}, \tilde{Y}, \tilde{Z}\}^{\otimes n}$. The degree of $P \in B$ is the number of non-identity elements in the product. The degree of a linear combination $\sum_{P \in B} \alpha_P P$ of elements in B is the largest degree of $P \in B$ such that $\alpha_P \neq 0$.

The fact that the degree is defined for any observable (which is equivalent to $\tilde{X}, \tilde{Y}, \tilde{Z}$ forming a basis of the space of operators) is the content of [Lemma 4.2](#). The existence of low-degree approximation is guaranteed by the following lemma.

Lemma 4.1 (Low-degree approximation). Let $\mathcal{O}$ be a bounded n -qubit observable and let D be a distribution over the Bloch sphere with mean $\bar{\mu}$ and Pauli second moment matrix S such that $\|S\| \leq 1 - \eta$ for some $\eta \in (0, 1)$. Then there exists a degree- d observable $\mathcal{O}^{\leq d}$ and a constant $\eta' \in (0, 1)$ such that

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [(\text{Tr}(\mathcal{O} \rho) - \text{Tr}(\mathcal{O}^{\leq d} \rho))^2] \leq (1 - \eta')^d,$$

where η' is a function of η .

Once [Lemma 4.1](#) is established, [Theorem 1.1](#) follows readily from an application of [Corollary 2.5](#).

Proof of Theorem 1.1, using Lemma 4.1. Let $1 - \eta$ be a known upper bound on $\|S\|_{\text{op}}$. We assume the access model of [Section 2.2](#), where we get a set S of examples $[\text{Tr}(P\rho)]$ for 1-local P . Our algorithm is as follows:

1. Compute η' as in the last line of [Lemma 4.5](#). Let $d := O(\log(1/\varepsilon)/\log(1/(1 - \eta')))$.
2. Draw a set S of size $n^d \cdot \log(1/\delta)$, and initialize S' to be empty.
3. For each $x \in S$, prepare a set T of $\log(|S| + 1/\varepsilon^2 + 1/\delta)$ copies of the state ρ that matches the 1-local expectations of x . Let $\text{est}(\text{Tr}(\mathcal{O}\mathcal{E}[\rho]))$ be the estimate of $\text{Tr}(\mathcal{O}\mathcal{E}[\rho])$, where for each $\rho \in T$, we measure with respect to $\{\mathcal{E}^\dagger[\mathcal{O}], I - \mathcal{E}^\dagger[\mathcal{O}]\}$, and $\text{est}(\text{Tr}(\mathcal{O}\mathcal{E}[\rho]))$ is the empirical probability of measuring the first outcome. Add

$$(x^{\otimes \log(1/\varepsilon)/\log(1/(1-\eta'))}, \text{est}(\text{Tr}(\mathcal{O}\mathcal{E}[\rho])))$$

to the set S' .

4. Run linear regression on S' and output the returned hypothesis h .

By Hoeffding's inequality, each estimate of $\text{Tr}(\mathcal{O}\mathcal{E}[\rho])$ is within ε of its expectation with probability $1 - \exp(-\Omega(|T|\varepsilon^2))$. By union bound over the $|S|$ estimates, with probability $\geq 1 - \delta$, all estimates are within ε of $\text{Tr}(\mathcal{O}\mathcal{E}[\rho])$.

The time, sample, and error bounds follow from [Lemma 4.1](#), which guarantees that $\mathcal{E}^\dagger[\mathcal{O}]$ is ε -close to some degree- d polynomial. This implies the labels of our sample set are 2ε -close to such a polynomial. The dimension of the linear regression problem is $\leq \binom{n}{d} \cdot 3^d \leq n^{O(d)}$, as there are 3 choices for each non-identity component. Then by [Corollary 2.5](#), linear regression has time and sample complexity

$$n^{O(\log(1/\varepsilon)/\log(1/(1-\eta')))} \cdot \log(1/\delta)$$

and outputs a hypothesis h such that

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [(\text{Tr}(\mathcal{O}\mathcal{E}[\rho]) - \text{Tr}(h\rho))^2] \leq O(\varepsilon). \quad \square$$

The proof of [Lemma 4.1](#) follows from three technical Lemmas. The first gives, for any product distribution over n -qubit states, a (non-orthogonal) decomposition of any observable into operators which are centered and bounded in variance with respect to that distribution.

Lemma 4.2. *Let D be a distribution over the Bloch sphere. Let $U^\dagger AU$ be the eigendecomposition of $\bar{\rho} = \mathbb{E}_D[\rho]$, and let $\tilde{X}, \tilde{Y}, \tilde{Z} := U^\dagger XU, U^\dagger YU, \frac{U^\dagger ZU - \text{Tr}(\bar{\rho}U^\dagger ZU)I}{\sqrt{1 - \text{Tr}(\bar{\rho}U^\dagger ZU)^2}}$. Then $\{I, \tilde{X}, \tilde{Y}, \tilde{Z}\}$ is a basis for the set of 2×2 unitary Hermitian matrices, and thus every n -qubit observable $\mathcal{O}$ can be written as*

$$\mathcal{O} = \sum_{P \in B} \hat{\mathcal{O}}(P)P$$

for $B = \{I, \tilde{X}, \tilde{Y}, \tilde{Z}\}^{\otimes n}$. Furthermore, each non-identity $P \in B$ satisfies $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(P\rho)] = 0$ and $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(P\rho)^2] \leq 1$.

Proof. It is clear that $\{I, \tilde{X}, \tilde{Y}, \tilde{Z}\}$ is linearly independent and thus B forms a basis for any n -qubit observable. Note that $\text{Tr}(P\rho) = \prod_{i=1}^n \text{Tr}(P_i \rho_i)$ as ρ is a product state, so we can restrict the analysis to single qubit states drawn from D .

For $P = \tilde{X}$ or $\tilde{Y}$, it is clear that $\mathbb{E}_{\rho \sim D}[\text{Tr}(P\rho)] = 0$ and $\mathbb{E}_{\rho \sim D}[\text{Tr}(P\rho)^2] = 1$. For $\tilde{Z}$, we have

$$\mathbb{E}_{\rho \sim D}[\text{Tr}(\tilde{Z}\rho)] = \frac{\mathbb{E}[\text{Tr}(\rho U^\dagger Z U)] - \text{Tr}(\bar{\rho} U^\dagger Z U)}{\sqrt{1 - \text{Tr}(\bar{\rho} U^\dagger Z U)^2}} = 0.$$

Moreover,

$$\mathbb{E}_{\rho \sim D}[\text{Tr}(\tilde{Z}\rho)^2] = \frac{\text{Var}[\text{Tr}(\rho U^\dagger Z U)]}{1 - \text{Tr}(\bar{\rho} U^\dagger Z U)^2} \leq 1,$$

because $\text{Var}[\text{Tr}(\rho U^\dagger Z U)] + \mathbb{E}[\text{Tr}(\rho U^\dagger Z U)]^2 = \mathbb{E}[\text{Tr}(\rho U^\dagger Z U)^2] \leq 1$. $\square$

Remark 4.3. Going forward, we will assume w.l.o.g. that the mean of the distribution $\mathbb{E}_D[\rho]$ is diagonal because we can always transform the basis according to U . Thus, we will assume that the mean state $\mathbb{E}_D[\rho] = \frac{1}{2}(I + \mu Z)$ and the mean Bloch vector $\vec{\mu} = (0, 0, \mu)$ for some $\mu \in [-1, 1]$,

We next prove two important components required in the proof of [Lemma 4.1](#). The first lemma gives an eigenvalue lower bound on a Hermitian matrix that arises when we expand $\text{Tr}(\mathcal{O}^2 \mathbb{E}[\rho])$.

Lemma 4.4. Let $\mu \in (-1, 1)$ and

$$\tilde{M} = \begin{pmatrix} 1 & i\mu \\ -i\mu & 1 \end{pmatrix}.$$

Then $\lambda_{\min}(\text{Re}(\tilde{M}^{\otimes k})) \geq (1 - \mu^2)^{k/2}$ for any $k \in \mathbb{N}$.

Proof. Note that $\tilde{M} = I + \mu Y$ and Y is imaginary. Thus,

$$\text{Re}((I + \mu Y)^{\otimes k}) = \sum_{P \in \{I, Y\}^k : |P| \text{ even}} \mu^{|P|} \bigotimes_{i=1}^k P_i.$$

Note also that the eigenvalues of the above remain the same if we replace the Y 's with Z 's. Thus, the eigenvalues of the above can be indexed by $x \in \{\pm 1\}^k$, and can be expressed as follows:

$$\lambda(x) = \sum_{S \subseteq [k] : |S| \text{ even}} \mu^{|S|} x^S.$$

Let $f : \mathbb{R}^k \rightarrow \mathbb{R}$ be the function $f(x) = \prod_{i=1}^k (1 + x_i) = \sum_{S \subseteq [k]} x^S$. Then, we have $\lambda(x) = \frac{1}{2}(f(\mu x) + f(-\mu x))$. By the AM-GM inequality,

$$\lambda(x) \geq \sqrt{f(\mu x)f(-\mu x)} = \prod_{i=1}^k \sqrt{(1 + \mu x_i)(1 - \mu x_i)} = (1 - \mu^2)^{k/2}. \quad \square$$

The next lemma gives an upper bound on the Pauli covariance matrix Σ (recall [Definition 2.1](#)) scaled by a specific diagonal matrix.

Lemma 4.5. *Let D be a distribution over the Bloch sphere with mean $\vec{\mu} = (0, 0, \mu)$, Pauli second moment matrix $\mathcal{S}$ such that $\|\mathcal{S}\|_{\text{op}} \leq 1 - \eta$ for some $\eta \in (0, 1)$, and Pauli covariance matrix Σ . Let $\Delta = \text{diag}((1 - \mu^2)^{-1/4}, (1 - \mu^2)^{-1/4}, (1 - \mu^2)^{-1/2})$. Then there exists $\eta' \in (0, 1)$ such that $\|\Delta\Sigma\Delta\|_{\text{op}} \leq 1 - \eta'$, where η' is a function of η .*

Proof. We split into two cases based on whether $\mu^2 \geq \eta/2$. The case where $\mu^2 < \eta/2$ is the simpler case: since $\|\Sigma_{\text{op}}\| \leq \|\mathcal{S}\|_{\text{op}} \leq 1 - \eta$ and $\|\Delta\|_{\text{op}}^2 \leq (1 - \mu^2)^{-1}$, we have

$$\|\Delta\Sigma\Delta\|_{\text{op}} \leq \frac{\|\Sigma\|_{\text{op}}}{1 - \mu^2} \leq \frac{1 - \eta}{1 - \mu^2} \leq \frac{1 - \eta}{1 - \eta/2} \leq 1 - \eta/2.$$

Now we consider the case where the inequality is not satisfied. We note that $\Sigma = \mathcal{S} - \vec{\mu}\vec{\mu}^\top$ and $\vec{\mu} = (0, 0, \mu)$, thus we have that $\Delta\Sigma\Delta$ has the following block structure:

$$\Delta\Sigma\Delta = \begin{pmatrix} \frac{\mathcal{S}_{2 \times 2}}{\sqrt{1 - \mu^2}} & b \\ b^\dagger & \frac{\mathcal{S}_{zz} - \mu^2}{1 - \mu^2} \end{pmatrix},$$

where $\mathcal{S}_{2 \times 2}$ is the top left 2×2 block of $\mathcal{S}$, $\mathcal{S}_{zz}$ is the bottom right entry, and b is the remaining part (its values will not be directly relevant). Note also that $\Delta\Sigma\Delta \geq 0$, and it is well-known that any PSD matrix of the form $\begin{bmatrix} A & b \\ b^\dagger & c \end{bmatrix} \geq 0$ has operator norm at most $\|A\|_{\text{op}} + c$. We thus know that

$$\|\Delta\Sigma\Delta\|_{\text{op}} \leq \frac{\mathcal{S}_{zz} - \mu^2}{1 - \mu^2} + \frac{\|\mathcal{S}_{2 \times 2}\|_{\text{op}}}{\sqrt{1 - \mu^2}} \leq \frac{\mathcal{S}_{zz} - \mu^2}{1 - \mu^2} + \frac{\text{Tr}(\mathcal{S}_{2 \times 2})}{\sqrt{1 - \mu^2}} \leq \frac{\mathcal{S}_{zz} - \mu^2}{1 - \mu^2} + \frac{1 - \mathcal{S}_{zz}}{\sqrt{1 - \mu^2}},$$

where we use the fact that $1 = \text{Tr}(\mathcal{S}) = \text{Tr}(\mathcal{S}_{2 \times 2}) + \mathcal{S}_{zz}$. Then we have

$$\begin{aligned} \|\Delta\Sigma\Delta\|_{\text{op}} &\leq \frac{\mathcal{S}_{zz} - \mu^2}{1 - \mu^2} + \frac{1 - \mathcal{S}_{zz}}{\sqrt{1 - \mu^2}} \\ &= \mathcal{S}_{zz} \left(\frac{1}{1 - \mu^2} - \frac{1}{\sqrt{1 - \mu^2}} \right) + \frac{1}{\sqrt{1 - \mu^2}} - \frac{\mu^2}{1 - \mu^2} \\ &\leq (1 - \eta) \left(\frac{1}{1 - \mu^2} - \frac{1}{\sqrt{1 - \mu^2}} \right) + \frac{1}{\sqrt{1 - \mu^2}} - \frac{\mu^2}{1 - \mu^2} \\ &= 1 - \eta \left(\frac{1}{1 - \mu^2} - \frac{1}{\sqrt{1 - \mu^2}} \right) \\ &\leq 1 - \eta \frac{1 - \sqrt{1 - \eta/2}}{1 - \eta/2}. \end{aligned}$$

The third line is by the fact that $\mathcal{S}_{zz} \leq \|\mathcal{S}\|_{\text{op}} \leq 1 - \eta$, and the last inequality is because the function $\frac{1}{1 - \mu^2} - \frac{1}{\sqrt{1 - \mu^2}}$ is increasing with μ^2 and that we are in the $\mu^2 \geq \eta/2$ case.

Thus, combining the two cases, there always exists η' such that $\|\Delta\Sigma\Delta\|_{\text{op}} \leq 1 - \eta'$. Specifically, we have

$$\eta' \geq \min \left\{ \eta \frac{1 - \sqrt{1 - \eta/2}}{1 - \eta/2}, \eta/2 \right\} > 0. \quad \square$$

Remark 4.6. As mentioned in [Remark 1.4](#), our techniques can be extended to learn not just the mapping $\rho \mapsto \text{Tr}(O\mathcal{E}[\rho])$, but also the joint mapping $(O, \rho) \mapsto \text{Tr}(O\mathcal{E}[\rho])$. As this is standard, here we only briefly sketch the main ideas.

The general strategy is to produce a classical description of the channel that we can then use to make predictions about properties of output states. To do this, we draw many input states $\rho_1, \dots, \rho_N$ from $\mathcal{D}$, query the channel on each them, and for each of the output states $\mathcal{E}[\rho_j]$, we apply a randomized Pauli measurement to each of them and use these to form unbiased estimators for the output state. Concretely, given an output state $\mathcal{E}[\rho_j]$, a randomized Pauli measurement will result in a stabilizer state $|\psi^{(j)}\rangle = \otimes_{i=1}^n |s_i^{(j)}\rangle \in \{|0\rangle, |1\rangle, |+\rangle, |-\rangle, |y+\rangle, |y-\rangle\}^{\otimes n}$, and the expectation of $\otimes_{i=1}^n (3|s_i^{(j)}\rangle\langle s_i^{(j)}| - I)$ is $\mathcal{E}[\rho_j]$. The classical description of the channel is given by the $O(\log(1/\epsilon))$ -body reduced density matrices of the input states $\rho_1, \dots, \rho_N$, together with the classical encodings of $|\psi^{(1)}\rangle, \dots, |\psi^{(N)}\rangle$.

Given an observable O , we can then perform regression to predict the labels $\text{Tr}(O \otimes_{i=1}^n (3|s_i^{(j)}\rangle\langle s_i^{(j)}| - I))$ given the features $\{\text{Tr}(P\rho_j)\}_{|P| \leq O(\log(1/\epsilon))}$. Because the labels are unbiased estimates of $\text{Tr}(O\mathcal{E}[\rho_j])$, the resulting estimator will be an accurate approximation to $\rho \mapsto \text{Tr}(O\mathcal{E}[\rho])$ for $\rho \sim \mathcal{D}$.

In this work, we do not belabor these details as they are already investigated in depth in [\[HCP23\]](#). Instead, we focus on the single observable case as this is where the main difficulty lies in extending the results of [\[HCP23\]](#) to more general input distributions.

Now we prove [Lemma 4.1](#) which shows the existence of a low-degree approximator.

Proof of Lemma 4.1. Assume w.l.o.g., as in [Remark 4.3](#), that $\mathbb{E}_D[\rho] = \frac{1}{2}(I + \mu Z)$ and $\vec{\mu} = (0, 0, \mu)$. Let $\mathcal{O}$ be expressed in the basis $B = \{I, X, Y, \frac{Z - \mu I}{\sqrt{1 - \mu^2}}\}^{\otimes n}$ as in [Lemma 4.2](#). For a subset $S \subseteq [n]$, we will denote by $\{P \in B : P \sim S\}$ the set of basis elements with non-identity components at indices in S and identity components at indices in $\bar{S}$. We will also denote by $|P|$ the number of non-identity components in P .

Let $\mathcal{O}^{>d} := \sum_{|S|>d} \sum_{P \sim S} \widehat{\mathcal{O}}(S)P$. We will show that $\mathcal{O}^{>d}$ satisfies $\mathbb{E}_{\rho \sim D^{\otimes n}}[\text{Tr}(\mathcal{O}^{>d}\rho)^2] \leq (1 - \eta')^d$ where $\eta' \in (0, 1)$ depends on η .

We first note a bound on the related quantity $\mathbb{E}_{\rho \sim D^{\otimes n}}[\text{Tr}(\mathcal{O}^2\rho)]$. Because $\|\mathcal{O}\|_{\text{op}} \leq 1$ and $\text{Tr}(\mathbb{E}[\rho]) \leq 1$, $\text{Tr}(\mathcal{O}^2 \mathbb{E}[\rho]) \leq 1$ as well. We will expand this quantity as a quadratic form.

$$\begin{aligned} \mathbb{E}_{\rho \sim D^{\otimes n}}[\text{Tr}(\mathcal{O}^2\rho)] &= \sum_{P, Q \in B} \widehat{\mathcal{O}}(P)\widehat{\mathcal{O}}(Q)\text{Tr}(PQ \mathbb{E}[\rho]) \\ &= \sum_{P, Q \in B} \widehat{\mathcal{O}}(P)\widehat{\mathcal{O}}(Q)\text{Tr}(\otimes_{i=1}^n P_i Q_i \mathbb{E}[\rho_i]) \\ &= \sum_{P, Q \in B} \widehat{\mathcal{O}}(P)\widehat{\mathcal{O}}(Q) \prod_{i=1}^n \text{Tr}(P_i Q_i \mathbb{E}[\rho_i]). \end{aligned}$$

Note that the product is 0 whenever exactly one of P_i, Q_i is I for any $i \in [n]$ (as $\text{Tr}((Z - \mu I) \cdot (I + \mu Z)) = 0$). Therefore, we can partition the terms into groups that share a subset of identity variables:

$$\begin{aligned} \mathbb{E}_{\rho \sim D^{\otimes n}}[\text{Tr}(\mathcal{O}^2\rho)] &= \sum_{S \subseteq [n]} \sum_{P, Q \sim S} \widehat{\mathcal{O}}(P)\widehat{\mathcal{O}}(Q) \prod_{i=1}^n \text{Tr}(P_i Q_i \mathbb{E}[\rho_i]) \\ &= \sum_{S \subseteq [n]} \widehat{\mathcal{O}}_S^\dagger M_S^{\otimes |S|} \widehat{\mathcal{O}}_S, \end{aligned}$$

where $\widehat{\mathcal{O}}_S$ is the vector of coefficients for the set $\{P : P \sim S\}$, and M is the 3×3 matrix such that $M_{P,Q} = \text{Tr}(PQ \mathbb{E}_D[\rho])$:

$$M = \begin{pmatrix} 1 & i\mu & 0 \\ -i\mu & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Here, the entry $i\mu$ arises because $XY = iZ$ and $\text{Tr}(XY \cdot \frac{1}{2}(I + \mu Z)) = i\mu$. Since M is positive semidefinite, we have $\widehat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \widehat{\mathcal{O}}_S \geq 0$ for all S , and therefore we have

$$1 \geq \mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^2 \rho)] = \sum_{S \subseteq [n]} \widehat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \widehat{\mathcal{O}}_S \geq \sum_{S \subseteq [n]: |S| > d} \widehat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \widehat{\mathcal{O}}_S. \quad (2)$$

Now we will write the desired quantity $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^{>d} \rho)^2]$ as a quadratic form as well:

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^{>d} \rho)^2] = \sum_{P, Q \subseteq B: |P|, |Q| > d} \widehat{\mathcal{O}}(P) \widehat{\mathcal{O}}(Q) \prod_{i \in [n]} \mathbb{E}[\text{Tr}(P_i \rho) \text{Tr}(Q_i \rho)].$$

Similarly, the product is 0 whenever exactly one of P_i and Q_i is identity (by the guarantee of [Lemma 4.2](#) that $\mathbb{E}[\text{Tr}(P_i \rho)] = 0$ for $P_i \neq I$), so we can make the same partition:

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^{>d} \rho)^2] = \sum_{|S| > d} \widehat{\mathcal{O}}_S^\dagger M'{}^{\otimes |S|} \widehat{\mathcal{O}}_S. \quad (3)$$

Here M' is 3×3 matrix such that $M'_{PQ} = \mathbb{E}_{\rho \sim D} [\text{Tr}(P\rho) \text{Tr}(Q\rho)]$ for $P, Q \in \{X, Y, \frac{Z - \mu I}{\sqrt{1 - \mu^2}}\}$; in other words, it is the second moment matrix of the non-identity elements of our biased Pauli basis. Below, we show the entries of M' in terms of the Pauli covariance matrix Σ :

$$M' = \begin{pmatrix} \Sigma_{xx} & \Sigma_{xy} & \frac{\Sigma_{xz}}{\sqrt{1 - \mu^2}} \\ \Sigma_{xy} & \Sigma_{yy} & \frac{\Sigma_{yz}}{\sqrt{1 - \mu^2}} \\ \frac{\Sigma_{xz}}{\sqrt{1 - \mu^2}} & \frac{\Sigma_{yz}}{\sqrt{1 - \mu^2}} & \frac{\Sigma_{zz}}{1 - \mu^2} \end{pmatrix}.$$

Our aim is to bound M' in terms of M in order to show the existence of some $\eta' \in (0, 1)$ such that

$$\mathbb{E}[\text{Tr}(\mathcal{O}^{>d} \rho)^2] \leq (1 - \eta')^d \cdot \text{Tr}((\mathcal{O}^{>d})^2 \mathbb{E}[\rho]) \leq (1 - \eta')^d.$$

Comparing Eq. (2) and (3), it suffices to show that

$$(1 - \eta')^{|S|} \widehat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \widehat{\mathcal{O}}_S \geq \widehat{\mathcal{O}}_S^\dagger M'{}^{\otimes |S|} \widehat{\mathcal{O}}_S$$

for all vectors $\mathcal{O}_S$. Crucially, since the Pauli coefficients $\mathcal{O}_S$ must be *real*, it suffices to prove that

$$M'{}^{\otimes |S|} \leq (1 - \eta')^{|S|} \cdot \text{Re}(M^{\otimes |S|}). \quad (4)$$

Let $M_{2 \times 2}$ be the top left 2×2 block in M , which is exactly the matrix in [Lemma 4.4](#). From [Lemma 4.4](#), we have $\lambda_{\min}(\text{Re}(M_{2 \times 2}^{\otimes |S|})) \geq (1 - \mu^2)^{|S|/2}$. By the block structure of M , it follows that $\text{Re}(M^{\otimes |S|}) \geq \text{diag}(\sqrt{1 - \mu^2}, \sqrt{1 - \mu^2}, 1)^{\otimes |S|}$. Thus, we can establish Eq. (4) by proving that

$$M' \leq (1 - \eta') \cdot \text{diag}(\sqrt{1 - \mu^2}, \sqrt{1 - \mu^2}, 1). \quad (5)$$

Now, note that M' can be written as $\tilde{\Lambda}\Sigma\tilde{\Lambda}$, where Σ is the covariance matrix of D and $\tilde{\Lambda} = \text{diag}(1, 1, (1-\mu^2)^{-1/2})$. Letting $\Delta = \text{diag}((1-\mu^2)^{-1/4}, (1-\mu^2)^{-1/4}, (1-\mu^2)^{-1/2})$ (which is the diagonal matrix defined in Lemma 4.5), we can see that Eq. (5) is equivalent to

$$\|\Delta M' \Delta\|_{\text{op}} \leq 1 - \eta',$$

By Lemma 4.5, this is true by our assumption that $\|\mathcal{S}\|_{\text{op}} \leq 1 - \eta$. This establishes Eq. (5) and thus Eq. (4), and from Eq. (2) and (3), we have

$$\begin{aligned} 1 &\geq \mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^2 \rho)] \geq \sum_{|S| > d} \hat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \hat{\mathcal{O}}_S \\ &\geq (1 - \eta')^{-d} \sum_{|S| > d} \hat{\mathcal{O}}_S^\dagger M'^{\otimes |S|} \hat{\mathcal{O}}_S \\ &= (1 - \eta')^{-d} \mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^{>d} \rho)^2]. \end{aligned}$$

Therefore,

$$\mathbb{E}_{\rho \sim D^{\otimes n}} [(\text{Tr}(\mathcal{O} \rho) - \text{Tr}(\mathcal{O}^{\leq d} \rho))^2] \leq (1 - \eta')^d.$$

Hence, we have obtained the claimed result. $\square$

Remark 4.7. In the proof of Lemma 4.1, we relate the quantity $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^{>d} \rho)^2]$ that we wish to bound to the related quantity $\mathbb{E}_{\rho \sim D^{\otimes n}} [\text{Tr}(\mathcal{O}^2 \rho)]$, which is at most 1 since $\|\mathcal{O}\|_{\text{op}} \leq 1$. Due to the choice of our biased Pauli basis $B = \{I, X, Y, \frac{Z - \mu I}{\sqrt{1 - \mu^2}}\}^{\otimes n}$, we may write both quantities as sums of $\hat{\mathcal{O}}_S^\dagger M'^{\otimes |S|} \hat{\mathcal{O}}_S$ and $\hat{\mathcal{O}}_S^\dagger M^{\otimes |S|} \hat{\mathcal{O}}_S$ (see Eq. (2) and (3)).

Suppose ρ is a product of different distributions where each qubit has mean Bloch vector $\vec{\mu}_i = (0, 0, \mu_i)$ (after rotation; see Remark 4.3) and second moment bounded by $1 - \eta$. We will instead use the basis $B = \otimes_{i=1}^n \{I, X, Y, \frac{Z - \mu_i I}{\sqrt{1 - \mu_i^2}}\}$. One can easily see that the above quantities are essentially the same, with $M^{\otimes |S|}$ replaced by $\otimes_{i \in S} M_i$ and similarly for M' . Then, the next steps in the proof (establishing Eq. (4) and (5)) are exactly the same.

5 Lower bounds

In this section, we show lower bounds in the classical case, which automatically imply hardness in the quantum case. In Section 5.1, we prove Theorem 1.5, which shows hardness of learning without the product distribution assumption in Theorem 1.3. In Section 5.2, we show that truncating in the unbiased basis fails when the distribution is not mean zero.

5.1 Lower bounds for learning non-product distributions

In this section, we prove Theorem 1.5. We show that if $\mathcal{D}$ is an arbitrary distribution, then even if $\mathcal{D}$ is supported on $[-(1 - \eta), (1 - \eta)]^n$ (i.e., in the interior of the hypercube), there is no learning algorithm.

Let C be a code over $\{0, 1\}^n$ of size $2^{\Theta(n)}$ and distance $n/4$. The following fact is standard.

Fact 5.1. For any constant $\varepsilon > 0$, any learning algorithm that can learn an arbitrary function $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ over C to ε error requires $2^{\Omega(n)}$ queries.

Proof of Theorem 1.5. Let $\eta = 0.1$. We set the distribution $\mathcal{D}$ to be the uniform distribution over $(1-\eta) \cdot C$, which is supported in $[-(1-\eta), (1-\eta)]^n$. Let $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ be an arbitrary function. Since the code C has distance $n/4$, we can without loss of generality assume that $f(x') = f(x)$ whenever $d(x, x') \leq n/8$.

Suppose for contradiction that there is an algorithm that, with only $2^{o(n)}$ queries to f , outputs a function $g : [-1, 1]^n \rightarrow \mathbb{R}$ such that $\mathbb{E}_{x \sim \mathcal{D}} [(g(x) - f(x))^2] \leq \varepsilon$. Let $p := 1 - \eta/2$, and let $\text{Ber}(p)$ be the distribution where we have $+1$ with probability p and -1 otherwise. Then, for any $x \in \{\pm 1\}^n$, we have $f((1-\eta)x) = \mathbb{E}_{z \sim \text{Ber}(p)^{\otimes n}} [f(x \circ z)]$, where $\circ$ denotes entry-wise product.

We claim that for any $x \in C$, $|\mathbb{E}_{z \sim \text{Ber}(p)^{\otimes n}} [f(x \circ z)] - f(x)| \leq o_n(1)$. The number of -1 coordinates in z is distributed as a $\text{Bin}(n, \eta/2)$. By the Chernoff bound, $\Pr[\text{Bin}(n, \eta/2) \geq (1+\delta)\frac{1}{2}\eta n] \leq e^{-O(\delta^2 \eta n)} = o_n(1)$. In particular, for $\eta = 0.1$, we have that $\Pr[\text{Bin}(n, \eta/2) \geq n/8] \leq o_n(1)$. Thus, with probability $1 - o_n(1)$, $d(x \circ z, x) < n/8$, which means that $f(x \circ z) = f(x)$. This proves that $|f((1-\eta)x) - f(x)| \leq o_n(1)$ for all $x \in C$.

Then, let $h(x) = g((1-\eta)x)$.

$$\begin{aligned} \mathbb{E}_{x \in C} [(h(x) - f(x))^2] &= \mathbb{E}_{x \in C} [(g((1-\eta)x) - f(x))^2] \\ &\leq \mathbb{E}_{x \in C} [(g((1-\eta)x) - f((1-\eta)x))^2] + o_n(1) \\ &= \mathbb{E}_{x \sim \mathcal{D}} [(g(x) - f(x))^2] + o_n(1) \\ &\leq \varepsilon + o_n(1). \end{aligned}$$

This means that h is an ε -approximation of f over C . This contradicts [Fact 5.1](#) thus completing the proof. $\square$

5.2 Lower bounds for unbiased degree truncation

In [Theorem 1.3](#), if the product distribution $\mathcal{D}$ over $[-(1-\eta), 1-\eta]^n$ has mean zero, then directly truncating f with respect to the standard monomial basis at degree $O(\log(1/\varepsilon)/\log(1/(1-\eta)))$ (independent of n) suffices. However, in this section, we will show that without the mean zero assumption, even truncating at $\Omega(n)$ degree w.r.t. the monomial basis does not give a small approximation error. Our counter-example implies that truncation in any distribution-oblivious basis will fail on some product distribution. In the quantum setting, this implies that low-degree truncation in the standard Pauli basis fails on some product distribution as well.

The counter-example is quite simple: f is the multilinear extension of the Boolean majority function, and $\mathcal{D}$ is supported on a single nonzero point (which is in fact a product distribution).

Fact 5.2 ([\[O'D14\]](#)). *Let f be the multilinear extension of the Boolean majority function. Then the ℓ_2 Fourier weight on terms of degree k is $\Theta(k^{-3/2})$.*

The following is a well-known fact in approximation theory.

Fact 5.3 (Chebyshev extremal polynomial inequality [\[Riv90\]](#)). *Let p be a degree- d univariate polynomial with leading coefficient 1. Then, $\max_{x \in [-1, 1]} |p(x)| \geq 2^{-d+1}$.*

Lemma 5.4. *Let $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ be the majority function extended to the domain $[-1, 1]^n$. Let $0 < a < b < 1$ be fixed constants. Then, there exist $\delta := \delta(a, b) \in (0, 1)$ and $t^* \in [a, b]$ such that the degree- δn truncation $f^{\leq \delta n}$ has $|f^{\leq \delta n}(t^* \cdot \vec{1})| \geq \omega_n(1)$.*

Proof. Let $g(t) := f^{\leq d}(t \cdot \vec{1})$, which is a univariate polynomial of degree d , and consider the shifted polynomial $\tilde{g}(t) = g(\frac{b-a}{2}t + \frac{a+b}{2})$. The leading coefficient of $\tilde{g}$, denoted $\tilde{c}_d$, is $c_d(\frac{b-a}{2})^d$, where c_d is the leading coefficient of g . Since $g(t) = f^{\leq d}(t \cdot \vec{1})$, we have $c_d = \sum_{S: |S|=d} \hat{f}(S)$. For the majority function, **Fact 5.2** implies that $\binom{n}{d} \hat{f}(S)^2 = \Theta(d^{-3/2})$ for all S of size d , thus

$$|\tilde{c}_d| = \left(\frac{b-a}{2}\right)^d \binom{n}{d}^{1/2} \Theta(d^{-3/4}).$$

Then, by **Fact 5.3**, there must be a $s^* \in [-1, 1]$ such that

$$|\tilde{g}(s^*)| \geq 2^{-d+1} \cdot \left(\frac{b-a}{2}\right)^d \binom{n}{d}^{1/2} \Theta(d^{-3/4}).$$

If $d = \delta n$, then $\binom{n}{d} \geq (\frac{1}{\delta})^d = e^{d \log(1/\delta)}$. Thus, given $0 < a < b < 1$, there exists a $\delta \in (0, 1)$ such that the above is $\exp(\Omega(d))$. Thus, there exists a $t^* \in [a, b]$ such that $|f^{\leq d}(t^* \cdot \vec{1})| \geq \omega_n(1)$. $\square$

Acknowledgments

SC and HH would like to thank Ryan O'Donnell for a helpful discussion at an early stage of this project. Much of this work was completed while JD and JL were interns at Microsoft Research.

References

- [ADGP24] Srinivasan Arunachalam, Arkopal Dutt, Francisco Escudero Gutiérrez, and Carlos Palazuelos. Learning low-degree quantum objects. *arXiv preprint arXiv:2405.10933*, 2024.
- [BY23] Zongbo Bao and Penghui Yao. Nearly optimal algorithms for testing and learning quantum junta channels. *arXiv preprint arXiv:2305.12097*, 2023.
- [CHI⁺20] Cristina Cirstoiu, Zoe Holmes, Joseph Iosue, Lukasz Cincio, Patrick J Coles, and Andrew Sornborger. Variational fast forwarding for quantum simulation beyond the coherence time. *npj Quantum Information*, 6(1):82, 2020.
- [CNY23] Thomas Chen, Shivam Nadimpalli, and Henry Yuen. Testing and learning quantum juntas nearly optimally. In *Proceedings of the 2023 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1163–1185. SIAM, 2023.
- [EI22] Alexandros Eskenazis and Paata Ivanisvili. Learning low-degree functions from a logarithmic number of random queries. In *Proceedings of the 54th Annual ACM SIGACT Symposium on Theory of Computing*, pages 203–207, 2022.
- [FJS91] Merrick L Furst, Jeffrey C Jackson, and Sean W Smith. Improved learning of ac 0 functions. In *COLT*, volume 91, pages 317–325, 1991.
- [FO21] Steven T Flammia and Ryan O'Donnell. Pauli error estimation via population recovery. *Quantum*, 5:549, 2021.

- [GHC⁺24] Joe Gibbs, Zoe Holmes, Matthias C Caro, Nicholas Ezzell, Hsin-Yuan Huang, Lukasz Cincio, Andrew T Sornborger, and Patrick J Coles. Dynamical simulation via quantum machine learning with provable generalization. *Physical Review Research*, 6(1):013241, 2024.
- [HBC⁺22] Hsin-Yuan Huang, Michael Broughton, Jordan Cotler, Sitan Chen, Jerry Li, Masoud Mohseni, Hartmut Neven, Ryan Babbush, Richard Kueng, John Preskill, et al. Quantum advantage in learning from experiments. *Science*, 376(6598):1182–1186, 2022.
- [HCP23] Hsin-Yuan Huang, Sitan Chen, and John Preskill. Learning to predict arbitrary quantum processes. *PRX Quantum*, 4(4):040337, 2023.
- [HFW20] Robin Harper, Steven T Flammia, and Joel J Wallman. Efficient learning of quantum noise. *Nature Physics*, 16(12):1184–1188, 2020.
- [HKP20] Hsin-Yuan Huang, Richard Kueng, and John Preskill. Predicting many properties of a quantum system from very few measurements. *Nature Physics*, 16(10):1050–1057, 2020.
- [HKP21] Hsin-Yuan Huang, Richard Kueng, and John Preskill. Information-theoretic bounds on quantum advantage in machine learning. *Physical Review Letters*, 126(19):190505, 2021.
- [HLB⁺24] Hsin-Yuan Huang, Yunchao Liu, Michael Broughton, Isaac Kim, Anurag Anshu, Zeph Landau, and Jarrod R McClean. Learning shallow quantum circuits. *arXiv preprint arXiv:2401.10095*, 2024.
- [HP07] Patrick Hayden and John Preskill. Black holes as mirrors: quantum information in random subsystems. *Journal of high energy physics*, 2007(09):120, 2007.
- [HP19] Patrick Hayden and Geoffrey Penington. Learning the alpha-bits of black holes. *Journal of High Energy Physics*, 2019(12):1–55, 2019.
- [HYF21] Robin Harper, Wenjun Yu, and Steven T Flammia. Fast estimation of sparse quantum noise. *PRX Quantum*, 2(1):010322, 2021.
- [KKMS08] Adam Tauman Kalai, Adam R Klivans, Yishay Mansour, and Rocco A Servedio. Agnostically learning halfspaces. *SIAM Journal on Computing*, 37(6):1777–1805, 2008.
- [KM91] Eyal Kushilevitz and Yishay Mansour. Learning decision trees using the fourier spectrum. In *Proceedings of the twenty-third annual ACM symposium on Theory of computing*, pages 455–464, 1991.
- [KS01] Adam R Klivans and Rocco Servedio. Learning DNF in time $2^{\tilde{O}(n^{1/3})}$. In *Proceedings of the thirty-third annual ACM symposium on Theory of computing*, pages 258–265, 2001.
- [KSVZ23] Ohad Klein, Joseph Slote, Alexander Volberg, and Haonan Zhang. Quantum and classical low-degree learning via a dimension-free remez inequality. *arXiv preprint arXiv:2301.01438*, 2023.

- [LMN93] Nathan Linial, Yishay Mansour, and Noam Nisan. Constant depth circuits, fourier transform, and learnability. *Journal of the ACM (JACM)*, 40(3):607–620, 1993.
- [Mar74] Grigorii Aleksandrovich Margulis. Probabilistic characteristics of graphs with large connectivity. *Problemy peredachi informatsii*, 10(2):101–108, 1974.
- [MO08] Ashley Montanaro and Tobias J Osborne. Quantum boolean functions. *arXiv preprint arXiv:0810.2435*, 2008.
- [MRL08] Masoud Mohseni, Ali T Rezakhani, and Daniel A Lidar. Quantum-process tomography: Resource analysis of different strategies. *Physical Review A*, 77(3):032322, 2008.
- [MRT18] Mehryar Mohri, Afshin Rostamizadeh, and Ameet Talwalkar. *Foundations of machine learning*. 2018.
- [NPVY23] Shivam Nadimpalli, Natalie Parham, Francisca Vasconcelos, and Henry Yuen. On the pauli spectrum of qac0. *arXiv preprint arXiv:2311.09631*, 2023.
- [O'D14] Ryan O'Donnell. *Analysis of Boolean Functions*. Cambridge University Press, 2014.
- [PSSY22] Geoff Penington, Stephen H Shenker, Douglas Stanford, and Zhenbin Yang. Replica wormholes and the black hole interior. *Journal of High Energy Physics*, 2022(3):1–87, 2022.
- [Riv90] Theodore J Rivlin. *Chebyshev polynomials*. Courier Dover Publications, 1990.
- [Rus81] Lucio Russo. On the critical percolation probabilities. *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete*, 56(2):229–237, 1981.
- [Rus82] Lucio Russo. An approximate zero-one law. *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete*, 61(1):129–139, 1982.
- [RWZ22] Cambyse Rouzé, Melchior Wirth, and Haonan Zhang. Quantum talagrand, kkl and friedgut's theorems and the learnability of quantum boolean functions. *arXiv preprint arXiv:2209.07279*, 2022.
- [VDBMKT23] Ewout Van Den Berg, Zlatko K Mineev, Abhinav Kandala, and Kristan Temme. Probabilistic error cancellation with sparse pauli–lindblad models on noisy quantum processors. *Nature Physics*, 19(8):1116–1121, 2023.
- [VZ23] Alexander Volberg and Haonan Zhang. Noncommutative bohnenblust–hille inequalities. *Mathematische Annalen*, pages 1–20, 2023.
- [YE23] Lisa Yang and Netta Engelhardt. The complexity of learning (pseudo) random dynamics of black holes and other chaotic systems. *arXiv preprint arXiv:2302.11013*, 2023.