

Provable learning separation for predicting time-evolution of quantum many-body systems

Rahul Bandyopadhyay,^{1,2} Riccardo Molteni,^{1,2} Jens Eisert,^{3,4,5} Vedran Dunjko,^{1,2} and Sofiene Jerbi^{3,5}

¹ $\langle aQa^L \rangle$, Applied Quantum Algorithms, Leiden University, The Netherlands

²Leiden Institute of Advanced Computer Science (LIACS),
Leiden University, Einsteinweg 55, 2333 CC, Leiden, The Netherlands

³Dahlem Center for Complex Quantum Systems, Freie Universität Berlin, 14195 Berlin, Germany

⁴Fraunhofer Heinrich Hertz Institute, 10587 Berlin, Germany

⁵Helmholtz-Zentrum Berlin für Materialien und Energie, 14109 Berlin, Germany

Given that quantum computers are naturally suited to simulate the behavior of quantum many-body systems, an immediate question arises: can one formulate *physically motivated* quantum machine learning (QML) tasks that exhibit learning separations? We address this problem by studying the learnability of quantum many-body dynamics from the perspective of probably approximately correct (PAC)-learning. Concretely, we devise a supervised learning problem where the training set consists of specifications of randomized stabilizer probe states, evolution times sampled uniformly at random from a polynomially large time interval $[0, T]$, coupled with expectation values of certain observables evaluated on the resulting time-evolved state under an unknown Hamiltonian. For this learning task, we provide an efficient quantum procedure whose training phase learns the underlying Hamiltonian from short-time training samples, and whose deployment phase combines Hamiltonian simulation with the classical shadows protocol to perform inference on a newly given data point. By contrast, the existence of $\mathcal{O}(\text{poly}(n))$ -time instances ensures classical hardness: by embedding a BQP-complete computation into the polynomially long time-dynamics of a low-intersection variant of the Feynman-Kitaev clock Hamiltonian construction, we show that, for a certain family of input distributions, no randomized classical polynomial-time algorithm can fulfill our learning condition, unless $\text{BQP} \subseteq \text{P/poly}$. Furthermore, we show that the classically hard instance maintains quantum learnability. We also give an interpretation of our results in learning-assisted certified quantum simulation. Taken together, our results demonstrate a rigorous learning separation for a natural ML task based on Hamiltonian evolution, while building connections between quantum learning theory, quantum simulation, and QML.

CONTENTS

I. Introduction	2
II. Related works	4
A. Supervised learning of quantum many-body physics via classical methods	4
B. Learning separations	5
C. Hamiltonian learning	6
III. Preliminaries	7
A. PAC-learning	7
B. Complexity theory	7
C. Hamiltonian learning	8
1. Background on the methods of Haah et al.	9
IV. Learning problem	9
V. Quantum learnability	11
VI. Quantum advantage	13
A. Classical hardness	14
B. Quantum learnability for a classically hard case	15
VII. Discussion	16
A. Distinction between the present work and earlier work	16

B. Summary and future directions	17
VIII. Acknowledgments	18
References	19
A. Classical shadows protocol	22
B. Quantum advantage	23
1. A low-intersection BQP-complete Hamiltonian with constant-precision parameters	23
2. Time-evolution under the Feynman-Kitaev clock Hamiltonian	24
a. Set-up	25
b. Time-evolution of history state	25
3. Classical hardness of learning	33
C. Quantum learnability	36
1. Additional definitions	36
2. Tail bound	37
3. Hamiltonian learning precision required for the continuous-parameter case	39
4. Sample and time complexity analysis for quantum learning algorithm	41
5. Quantum learnability of the classically hard instance	47

I. INTRODUCTION

Notions of ML are rapidly reshaping the world. Large language models and other applications of ML are now ubiquitous, substantially altering how we communicate, manage organizational tasks, and generate products in industrial settings. At the same time, although the concept of quantum computing is not new [1], recent years have seen the construction of quantum devices that suggest we may soon achieve quantum computers and simulators with practical utility [2, 3]. Indeed, a number of quantum algorithms are known to exhibit super-polynomial advantages over their classical counterparts, some of which are directly relevant to practical problems in optimization [4–6] and ML [7–11], as well as those in quantum simulation [12–19]. In light of these developments, a central challenge remains the identification of real-world ML tasks exhibiting meaningful learning separations. Initial results in this direction have already emerged, including rigorous, in-principle QML separations [7–9], as well as early evidence of robust quantum advantages for natural classical data [11, 20–25].

Nevertheless, examples of provable learning separations remain relatively scarce, with many of such results relying on cryptographic assumptions. Thus, it is fair to argue that one of the central challenges in the field remains the identification of further well-motivated separations, particularly those grounded in physical considerations. Indeed, it is widely believed that quantum computers are capable of simulating certain many-body systems more efficiently than classical computers [12, 26]. Guided by the intuition that similar advantages may hold for *learning* tasks as well, an immediate question arises:

Can one identify QML separations for physically motivated learning tasks?

An example of such a learning problem is that of predicting ground state properties of quantum systems, a quintessentially quantum learning task for which one may expect quantum learners to yield an advantage over classical ML. Within the natural regime of gapped phases of matter, however, the aforementioned learning problem has been found to be solvable using classical ML algorithms [27–30], and the question of whether quantum advantage is attainable in learning properties of ground states thus remains dauntingly open. Therefore, in this work, we instead consider a closely related problem that is equally well motivated from a physical perspective, namely, predicting Hamiltonian time-evolution. Any physical system that is sufficiently isolated from its environment undergoes quantum evolution generated by a Hamiltonian. In analog quantum simulation [31–34], one studies the time-dependent evolution generated by some precisely controlled Hamiltonian. Contrasting this, digital quantum simulation involves approximating the dynamics of some target quantum system by a discrete sequence of elementary quantum gates [12, 13].

More generally, quantum simulation, whether that concerns real-time dynamics or preparation of ground or thermal states, may be regarded as a *forward computational* problem. In turn, if we consider the problem of estimating

the parameters of a many-body quantum system from measurement data thereof: in Ref. [35], this task has been conceptualized as an *inverse, Hamiltonian learning* problem. Within an ML context, this dichotomy bears a natural analogy to *training* and *inference*, or, in the language of PAC-learning, to *identification* and *evaluation*. At a high level, the present work builds on this conceptual viewpoint, using tools from Hamiltonian learning and notions from complexity theory to formulate ML tasks based on Hamiltonian evolution that are classically hard yet quantumly tractable. In other words, we focus on showing a learning separation of the type considered in Refs. [21–25], while leveraging ideas from quantum learning theory to inform the design of our learning procedure.

A closely related work in this regard is Ref. [24], which has considered the problem of learning time-dynamics generated by an unknown Hamiltonian from classical measurement data. For this task, the authors of that work have shown a learning separation, premised on the widely believed complexity-theoretic conjecture that $\text{BQP} \not\subseteq \text{P/poly}$. However, their learning problem is formulated in terms of a relatively restrictive family of Hamiltonians, namely, ones parametrized by $\mathcal{O}(\log n)$ -dimensional coefficient vectors. Consequently, this leads to an ML task that, while being physically motivated, does not capture the learning of physically plausible local Hamiltonians with a number of terms scaling, say, linearly or polynomially in the system size. By contrast, in this work, by assuming access to a stronger form of training data, we address the problem of predicting the dynamics of a more general family of Hamiltonians. Concretely, our contributions are as follows:

1. We formulate a supervised learning task where one is required to predict expectation values of quantum states that have been time-evolved by some unknown, *low-intersection* Hamiltonian, containing a number of constituent terms scaling as $\mathcal{O}(n)$. More precisely, our setting involves input data specifying probe states given by randomized stabilizer product states, as well as evolution times sampled uniformly at random, labeled by a list of expectation values evaluated on certain observables, which may be thought of as being morally equivalent to a classical description of the probe states after time-evolution by the unknown Hamiltonian.
2. We establish an exponential learning separation for the aforementioned learning problem, premised on the well-founded complexity-theoretic assumption $\text{BQP} \not\subseteq \text{P/poly}$. In particular, we provide an efficient quantum learning algorithm whose training phase is based on a simplified variant of the Hamiltonian learning protocol of Haah, Kothari, and Tang [36], and the deployment stage of which is a combination of Hamiltonian simulation and the classical shadow formalism. To complete the argument and prove classical hardness, we rely on the construction of Oliveira and Terhal [37], who have provided a low-intersection version of the Feynman-Kitaev clock Hamiltonian whose time-evolution realizes a BQP-complete computation for polynomially large evolution times.

Intuitively, the set-up is as follows: our training dataset corresponds to the expectation values of time-evolved states for evolution times that are uniformly sampled over a time interval $[0, T]$, where $T \in \mathcal{O}(\text{poly}(n))$, so that the samples may be regarded as being divided into “short”- and “long”-time regimes. The short-time window constitutes an inverse-polynomial fraction of the time interval, and its corresponding samples are used during training to learn the data-generating Hamiltonian in a classical fashion. On the other hand, we require the presence of samples corresponding to polynomially long time-evolutions to ensure classical hardness. That is, given a time $t \sim \text{Unif}([0, T])$ during the inference stage, the learner may be required to predict expectation values even for polynomially long evolution times, which is precisely the regime where the underlying dynamics can encode BQP-complete computations.

We emphasize two aspects of our work. In our setting, we assume access to training data that is sufficiently rich to enable Hamiltonian learning, unlike that of Ref. [24], where the labels are scalars given by *single* expectation values. In other words, while the formulation of their learning problem is substantially more elegant than ours, from a methodological perspective, it precludes the use of approaches based on Hamiltonian learning which requires a set of sufficient statistics to fully reconstruct the Hamiltonian. More fundamentally, as we explain further in Subsection VII A, it also imposes complexity-theoretic limitations on the kinds of Hamiltonian families to which their approach can be applied [38].

We circumvent these difficulties by assuming stronger and more informative training data, naturally motivating the use of a simple Hamiltonian learning algorithm that is tailored to our randomized-time setting. At the same time, we draw attention to the fact that this leads to a training procedure that is somewhat trivial and algorithmically elementary, involving what amounts to classical processing of the training samples and a *proper* PAC identification of the concept generating the training data. The source of our quantum learning advantage thus lies in the evaluation of the concept and not the identification thereof, leaving open the possibility of further identification-based separations in physically motivated learning tasks.

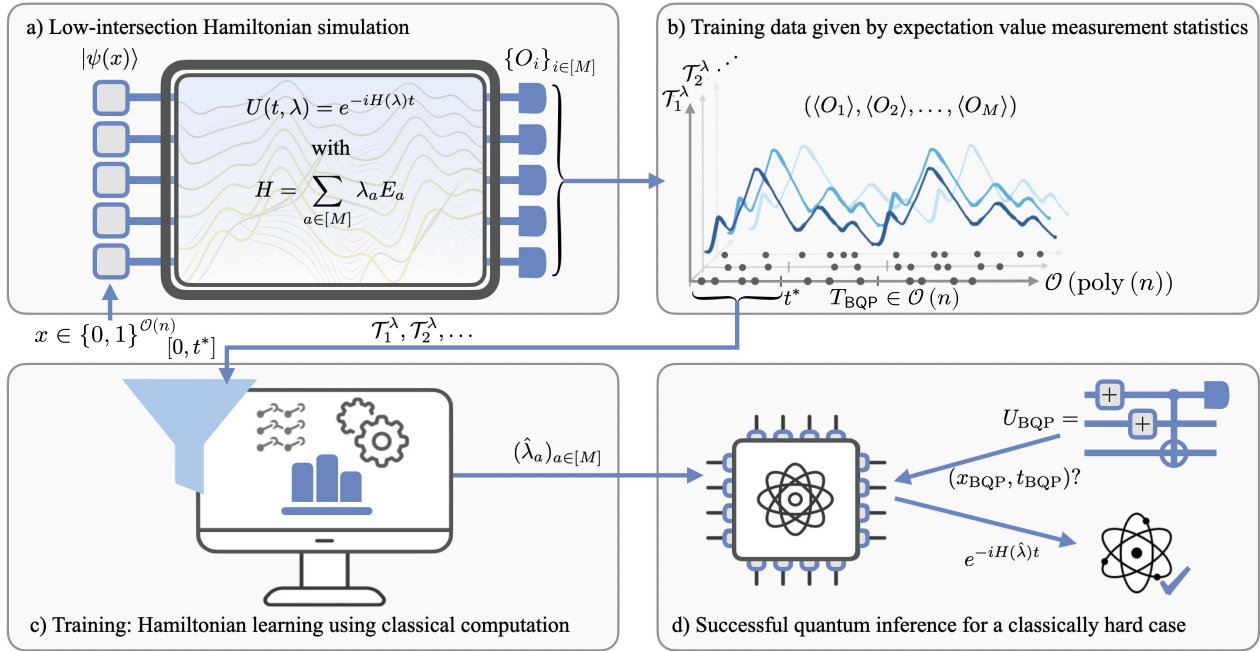


FIG. 1: A visualization of the key ideas of the present work. (a) Our training samples are generated by the time-evolution of a low-intersection Hamiltonian $H(\lambda)$ for randomly chosen times $t \sim \text{Unif}([0, T])$ acting on some initial probe state vector $|\psi(x)\rangle$. (b) We consider as our labels a vector of expectation values evaluated on a set of certain pre-determined observables. This has been schematically depicted in the above figure as a collection of measurement samples. (c) We filter out the short-time samples from our training data and use them to identify the underlying data-generating Hamiltonian via a simple Hamiltonian learning algorithm that is classical. (d) Finally, we use a quantum computer to successfully perform inference on a newly given data point. By universality of Hamiltonian evolution, this may involve predicting a BQP-complete computation, which is believed to be hard for classical learners.

Organization of this work. This work is organized as follows. In Section II, we provide an overview of the related literature on supervised learning of quantum many-body systems, learning separations, both information- and complexity-theoretic, and Hamiltonian learning. This is followed by Section III, where we provide the requisite preliminaries on computational learning theory, complexity theory, and Hamiltonian learning protocols. We then rigorously formulate our learning problem in the language of the PAC-learning formalism in Section IV. Then, in Subsection V, we provide a detailed description of our quantum learning algorithm, following which we give an overview of our proof of classical hardness in Section VI. We conclude the main text with a discussion section where, in addition to enumerating a list of open questions and future directions, we outline in detail the main differences between our work and Ref. [24].

II. RELATED WORKS

A. Supervised learning of quantum many-body physics via classical methods

In Ref. [27], the authors employ classical ML techniques to learn classical representations of ground states of gapped Hamiltonians, as well as expectation values on local observables, with rigorous guarantees on sample and computational complexities. The success of classical ML in solving the aforementioned learning tasks is essentially due to the fact that, within the gapped regime, the physical properties of quantum systems and, in particular, their ground states are well-behaved, in the sense that they are smooth functions of the Hamiltonian specifications. Thus, the authors are able to use the l_2 -Dirichlet kernel to learn such functions, using a number of training samples and computational runtime that scaling as $\mathcal{O}(n^{1/\varepsilon})$, so that, for constant precision $\varepsilon \in \mathcal{O}(1)$, their algorithm is efficient. For the same learning task, follow-up papers considered the setting where the learner is equipped with additional knowledge of certain aspects of the

quantum system being learned, such as its geometry and the knowledge of the observable being predicted. These include works by Lewis et al. [28] and Wanner et al. [29], which have improved the scaling of Ref. [27] to $\mathcal{O}(\log n \cdot 2^{\text{poly} \log(1/\varepsilon)})$ and $\mathcal{O}(2^{\text{poly} \log(1/\varepsilon)})$, respectively, while retaining the quasi-polynomial dependence on system size for non-constant ε . As extensions of the gapped scenario, subsequent works have considered important classes of states beyond ground states of gapped Hamiltonians. In Ref. [30], the authors have employed classical methods to learn properties of thermal states with exponentially decaying correlations, in addition to ground and thermal states satisfying the *generalized approximate local indistinguishability* property. The setting of Ref. [30] has further been generalized in Ref. [39] to Lindbladian phases of matter [40], which provably encompass the more conventional notions of phases. Finally, within the complementary setting of predicting dynamical properties of quantum systems, a noteworthy work is that of Huang, Chen, and Preskill [41], where they address a supervised learning problem involving quantum states evolved by quantum channels of arbitrary complexity. Remarkably, by generalizing the notion of a low-degree approximation from classical PAC-learning theory to quantum observables, they have been able to devise a quasi-polynomial time classical ML algorithm.

Relation of the present work to the above. While our work is also situated within the context of the above literature, it should be noted that the vast majority of the aforementioned papers have to do with predicting *static* properties of quantum systems, in that the learning problems they consider involve ground and thermal states. Moreover, the works mentioned above are able to successfully tackle such problems via methods that are either entirely or predominantly classical. Given the success of classical methods in this regime, in contrast, we consider the challenge of showing learning separations for the task of predicting learning quantum *dynamics*. The advantage of this setting lies in the BQP-completeness of time-evolution, a fact that allows us to prove complexity-theoretic learning separations for our task.

B. Learning separations

Within the related field of quantum learning theory, information-theoretic learning separations have been shown for tasks involving genuinely quantum data, such as predicting properties of quantum states, learning quantum processes, quantum principal component analysis, as well as certain variants of quantum property testing [42–45]. Importantly, such separations are predicated on being able to transduce quantum data onto an external quantum memory in a stable manner, followed by coherent processing of quantum information. A different strand of research has focused on the supervised learning setting, similar to that of Ref. [27], involving primarily classical data obtained from quantum experiments [9, 21, 24, 25]. In contrast to the previously discussed works on quantum learning separations, which involve unconditional information-theoretic lower bounds, the latter papers prove separations that rely on the classical hardness of simulating BQP-complete processes. In particular, the authors of Ref. [21] have laid out a systematic framework for proving separations for quantum many-body learning tasks, by leveraging the BQP-completeness of the data-generating function to prove classical heuristic hardness-of-learning in physically motivated settings under a fixed input distribution. In Ref. [9], premised on the assumption $\text{BQP} \not\subseteq \text{P/poly}$, learning separations were shown within distribution-free settings as well, marking an improvement over previously shown hardness statements based on $\text{BQP} \not\subseteq \text{HeurP/poly}$. The learning problem considered there is that of predicting expectation values of time-evolved quantum states on unknown observables, given as either linear combination of k -local Pauli operators as $O(\lambda) = \sum_i \lambda_i P_i$, or unitarily parametrized observables, given as $O(\lambda) = W(\lambda) O W(\lambda)^\dagger$, for some Hermitian operator O and parametrized unitary $\lambda \mapsto W(\lambda)$.

In Ref. [24], the learning task is that of predicting expectation values of quantum states on a *known* observable, which the authors tackle via a combination of Fourier coefficient extraction of parametrized Trotter circuits, linear regression, and kernel methods. However, in contrast to Ref. [9], the quantum states are time-evolved according to an *unknown* parameterized Hamiltonian. The work of Ref. [25] introduces the *learning under quantum privileged information* (LUQPI) framework, a quantum generalization of the classical formalism of Vapnik et al. [46, 47]. There, the authors envision a supervised learning scenario involving minimal quantum resources, whereby a quantum computer is used only in the extraction of certain features. Within this framework, the authors then construct examples of concept classes demonstrating exponential learning separations and provide numerical evidence of examples where the use of quantum privileged information is shown to enhance the performance of classical ML algorithms for learning quantum systems. To conclude, we also mention the work of Ref. [23], wherein the authors asked: can quantum advantage in QML arise from the hardness of identifying the concept generating the data, as opposed to the hardness of evaluating it? There, the authors answered this question in the affirmative, by using the notion of random generatability to provide the first proofs of identification-based learning separations.

Relation of the present work to the above. Our work is in a similar vein to the works of Refs. [9, 21, 24, 25], in that we show a complexity-theoretic learning separation for the task of predicting expectation values of certain time-evolved quantum states. In terms of research direction, we reiterate that the current work is perhaps most closely aligned with that of Ref. [24], where the task of learning quantum real-time dynamics has also been considered. However, as we document in Subsection VII A, our work differs from theirs in several key aspects. Namely, we consider Hamiltonian families satisfying properties such as low-intersection and geometric locality, leading to a learning task that is arguably more natural from a physical standpoint. Other distinctions have to do with details of how the learning problem is defined, in addition to the learning algorithm itself, which involves a combination of Hamiltonian learning and simulation during training and inference respectively. In connection with this methodological difference, we make note of the fact that the works of Refs. [9, 21, 48] were the first to envision the Hamiltonian learning problem as an identification problem. Our use of Hamiltonian learning in the training procedure may be interpreted as leveraging this insight.

C. Hamiltonian learning

Hamiltonian learning is the task of inferring the Hamiltonian corresponding to some quantum many-body system of interest, given measurement data obtained from the system. Historically, parameter estimation of specific families of Hamiltonians has been a well-studied problem in quantum control, quantum metrology, and quantum sensing [49–55]. More recently, within the field of quantum learning theory, the literature on Hamiltonian learning has expanded substantially, with works on learning many-body systems from singular eigenstates [56–59], thermal states [30, 36, 60–66], and time-evolution [35, 36, 49, 52–54, 57, 61, 66–90]. Here, we focus on the third setting in particular, since our learning problem has to do with predicting expectation values on time-evolved states. This setting is particularly relevant for experimentally learning large-scale Hamiltonians from realistic data in a robust fashion [67]. For present purposes, however, rigorous arguments about precise sample complexities in idealized situations are particularly important. A landmark result within that literature has been the work of Haah, Kothari, and Tang, wherein rigorous guarantees were provided for the first time on Hamiltonian learning from high-temperature Gibbs states, as well as query-access to time evolution [36]. Subsequent works also considered a version of Hamiltonian learning where the interaction structure of the underlying many-body system is unknown to the learner [35, 74, 90]. By this, we mean that the interaction terms constituting the Hamiltonian are not provided as inputs to the problem, so that one is required to perform *structure learning* in addition to learning the interaction strengths. A key requirement is Heisenberg-limited scaling, where one is interested in achieving a total query time of $t_{\text{time}} \in \mathcal{O}(1/\varepsilon)$, where ε is the desired precision on the Hamiltonian coefficients. Numerous works exist in this direction, such as those of Refs. [74–76, 79, 80, 88], as well as the related work of Dutkiewicz et al., where they show the fundamental necessity of a certain amount of quantum control in achieving Heisenberg-limited scaling [89].

Relation of present work to the above. As we will see later, our learning algorithm will involve a simplification of the Haah-Kothari-Tang protocol, an adjustment that will prove appropriate for our setting, since it involves training data corresponding to randomized time instances. So, while the learning problem considered in this work is a supervised one, we borrow heavily from the set-up of Ref. [36], and our analysis of quantum learnability will involve numerous mathematical objects considered in that work.

While the core aim of our work is the physically motivated learning separation as such, there are a few other implications that are of independent interest in the context of Hamiltonian learning. One interpretation of our results is in *learning-assisted certified quantum simulation*. Instead of attempting to certify a difficult long-time quantum evolution directly, we characterize an experimentally accessible short-time regime, infer the governing Hamiltonian, and use this learned model as the basis for long-time quantum prediction. In this way, Hamiltonian learning serves as a calibration stage for subsequent quantum simulation, yielding rigorous guarantees on the resulting predictions. This perspective is somewhat reminiscent of the approach taken in Ref. [77], but here comes with explicit complexity-theoretic and learning-theoretic guarantees.

III. PRELIMINARIES

A. PAC-learning

The PAC-learning framework provides a formal and mathematically precise language for describing supervised learning problems. Within this framework, a learning problem is specified by a set of functions, called the *concept class*, usually denoted by $\mathcal{C}$. The individual functions contained in the concept class are called *concepts*. Each concept $c: \mathcal{X} \rightarrow \mathcal{Y}$ maps an input/instance space $\mathcal{X}$ to an output/label space $\mathcal{Y}$. In the supervised learning setting, a learner has access to a *training dataset* of the form $\mathcal{T} = \{(x_i, c(x_i))\}_{i=1}^N$. This is a set of input-output pairs, where the inputs x_i are drawn independently from an unknown but fixed *target distribution* $\mathcal{D}$ over the input space $\mathcal{X}$, and the outputs are examples that are labeled by a concept $c \in \mathcal{C}$. Given this dataset, the learner must then output a *hypothesis* $h: \mathcal{X} \rightarrow \mathcal{Y}$ that approximately “agrees” with the labeling concept c with high probability over the target distribution $\mathcal{D}$. To formalize the aforementioned picture, we will now give a definition of *efficient PAC-learnability*, which requires that the learner outputs a hypothesis in polynomial time.

Definition 1 (Efficient PAC-learnability). *Consider a concept class $\mathcal{C} = \{c_j\}_j$, comprised of concepts $c_j: \mathcal{X} \rightarrow \mathcal{Y}$, where $\mathcal{X}$ is the instance space and $\mathcal{Y}$ is the output space. Let $\mathcal{L}: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}_{\geq 0}$ be a loss function. Then, $\mathcal{C}$ is efficiently PAC-learnable with respect to a fixed distribution $\mathcal{D}$ over $\mathcal{X}$ if there exists an algorithm $\mathcal{A}$ that takes in as input a dataset $\mathcal{T} = \{(x_i, c(x_i))\}_{i=1}^N$, accuracy and confidence parameters $\varepsilon, \delta > 0$, such that, for all concepts $c \in \mathcal{C}$, the size of the dataset $N \in \mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$, where n is the dimension of the instance space, the algorithm $\mathcal{A}$ outputs a hypothesis function h satisfying*

$$\mathbb{E}_{x \sim \mathcal{D}} [\mathcal{L}(h(x), c(x))] \leq \varepsilon, \quad (1)$$

with high probability $1 - \delta$ over $\mathcal{D}$ over $\mathcal{X}$, and runs in time $\mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$.

B. Complexity theory

The central result of our work is a learning separation premised on the widely-held complexity-theoretic assumption that $\text{BQP} \not\subseteq \text{P/poly}$. We therefore begin with the definition of the class BQP.

Definition 2 (BQP). *Let $\mathcal{L} := (\mathcal{L}_{\text{yes}}, \mathcal{L}_{\text{no}})$ be a promise problem, where $\mathcal{L}_{\text{yes}}, \mathcal{L}_{\text{no}} \subseteq \{0, 1\}^*$ and $\mathcal{L}_{\text{yes}} \cap \mathcal{L}_{\text{no}} = \emptyset$. Then, $\mathcal{L} \in \text{BQP}$ if and only if there exists a polynomial-time, uniform family of quantum circuits $\{U_n: n \in \mathbb{N}\}$, where, for all $n \in \mathbb{N}$, U_n takes in as input an n -qubit computational basis state and outputs a single bit obtained by measuring the first output qubit in the computational basis, such that the following properties are satisfied:*

1. *For all $x \in \mathcal{L}_{\text{yes}}$, the probability that $U_{|x|}$ accepts x is greater than or equal to $2/3$; that is, $\Pr[U_{|x|}(x) = 1] \geq 2/3$.*
2. *For all $x \in \mathcal{L}_{\text{no}}$, the probability that $U_{|x|}$ does not accept x is greater than or equal to $2/3$; that is, $\Pr[U_{|x|}(x) = 0] \geq 2/3$.*

An important category of complexity classes we will consider is that of advice classes, wherein a Turing machine is provided, in addition to the input, a polynomially sized *advice* string. Intuitively, the advice string should be thought of as enhancing the “power” of the computation being performed, compared to the setting where it is only supplied with the problem input. Advice classes are used to model the complexity of computational problems which involve a resource intensive pre-processing stage, followed by a post-processing stage that is computationally less expensive. This makes them suitable for analyzing the complexity of ML problems. We provide a definition of the class P/poly , as follows.

Definition 3 (P/poly , Def. 6.5, [91]). *A decision problem $\mathcal{L}: \{0, 1\}^* \rightarrow \{0, 1\}$ is contained in the complexity class P/poly if there exists a polynomial-time, classical algorithm, $\mathcal{A}$, satisfying the following property: for every $n \in \mathbb{N}$, there exists a polynomially sized bit-string, $\lambda_n \in \{0, 1\}^{\text{poly}(n)}$, such that, for all $x \in \{0, 1\}^n$, $\mathcal{A}(x, \lambda_n) = \mathcal{L}(x)$.*

C. Hamiltonian learning

Consider a quantum system defined on n qubits. The interactions characterizing the system are governed by a $2^n \times 2^n$ Hermitian operator called the *Hamiltonian*, which we will denote by H . Such an operator is expressible as a linear combination of Pauli operators, with their coefficients representing bounded interaction strengths. The task of *learning* a Hamiltonian, then, may be informally described as recovering ε -accurate estimates of the Hamiltonian coefficients with high probability, given as input either of the following:

1. Query access to the real-time evolution of the Hamiltonian, e^{-iHt} , for some $t > 0$.
2. Several, identically prepared copies of the thermal state of the Hamiltonian

$$\rho_\beta(H) = \frac{e^{-\beta H}}{\text{Tr}[e^{-\beta H}]}, \quad (2)$$

at a known inverse temperature $\beta > 0$.

The setting of the present work will involve the first access model, in that the input to our supervised learning problem will be a training dataset containing expectation values of randomized probe states that have been evolved under an unknown, low-intersection Hamiltonian, which we define below.

Definition 4 (Hamiltonian). *A Hamiltonian H is given by a set of tuples of the form $\{(\lambda_a, E_a)\}_{a \in [M]}$, where $M \in \mathbb{N}$, so that $a \in [M]$ is an index ranging over $[M] = \{1, 2, \dots, M\}$, $\lambda_a \in [-1, 1]$ for all $a \in [M]$, and $(E_a)_{a \in [M]}$ are unique, non-identity Pauli operators. The associated Hamiltonian is then given as*

$$H := \sum_{a \in [M]} \lambda_a E_a. \quad (3)$$

On occasion, we will also use $\lambda := (\lambda_1, \lambda_2, \dots, \lambda_M)$ to denote the entire vector of coefficients, and $H(\lambda)$ to express the Hamiltonian as a function of λ . We also say that H is $\mathfrak{K}$ -local if $|\text{supp}(E_a)| \leq \mathfrak{K}$ for all $a \in [M]$.

Definition 5 (Low-intersection Hamiltonian [36]). *Let $H = \sum_{a \in [M]} \lambda_a E_a$ be an n -qubit, $\mathfrak{K}$ -local Hamiltonian as in Def. 4. The dual-interaction graph $\mathfrak{G}$ of H is an undirected graph whose vertices are identified with elements of $[M]$, with an edge connecting vertices $a, b \in [M]$ if*

$$\text{supp}(E_a) \cap \text{supp}(E_b) \neq \emptyset. \quad (4)$$

Let $\mathfrak{d}$ denote the maximum degree of $\mathfrak{G}$ over all vertices. Then, we say that H is low-intersection if both $\mathfrak{K}, \mathfrak{d} \in \mathcal{O}(1)$.

Low-intersection Hamiltonians have also been referred to in the literature as *sparsely-interacting Hamiltonians* [92] and *low-interaction Hamiltonians* [75]. We note here that low-intersection $\mathfrak{K}$ -local Hamiltonians subsume the more familiar class of geometrically local $\mathfrak{K}$ -local Hamiltonians, which constitute, by far, the most physically plausible family of Hamiltonians, since the latter must satisfy the additional locality constraint that each interaction term be supported on a spatially contiguous region. We now formally state the problem of Hamiltonian learning from query access to time-evolution.

Problem 1 (Hamiltonian learning from time-evolution, [36]). *Let $H := \{(\lambda_a, E_a) : a \in [M]\}$ be an n -qubit Hamiltonian as defined in Def. 4. Let $t \leq t_c$ be a known time, where $t_c \in \mathcal{O}(1)$ is a constant and $U := e^{-iHt}$. Then, assuming knowledge of $(E_a)_{a \in [M]}$, the learning problem involves taking in as inputs precision and confidence parameters $\varepsilon, \delta > 0$, and $\mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$ queries to the time-evolution unitary U , and producing as outputs $\hat{\lambda} := (\hat{\lambda}_a)_{a \in [M]}$ such that, with probability at least $1 - \delta$, $|\lambda_a - \hat{\lambda}_a| \leq \varepsilon$ holds for all $a \in [M]$.*

1. Background on the methods of Haah et al.

We now provide some background on the results of Ref. [36]. Let $H = \{(\lambda_a, E_a) \mid a \in [M]\}$ be a low-intersection Hamiltonian as in Defs. 4 and 5. The starting point of their Hamiltonian learning algorithm is an expectation value we denote by $\mathcal{F}_{a,t}$, which, for $a \in [M]$, they define as

$$\mathcal{F}_{a,t} := \text{Tr} \left[Q_a e^{-iHt} \left(\frac{I + P_a}{D} \right) e^{iHt} \right] = \frac{1}{D} \text{Tr} [Q_a e^{-iHt} P_a e^{iHt}]. \quad (5)$$

In the above, P_a is chosen to be any single-qubit Pauli operator that anti-commutes with the Hamiltonian term E_a , $Q_a := i[P_a, E_a] = 2iP_a E_a$, $D := 2^n$, where n is the number of qubits constituting the many-body system, and $t \leq t_c \in \mathcal{O}(1/\text{poly}(\mathfrak{d}))$. It can be shown, as in the proof of Lemma 7 or Eq. (8) of Ref. [36], that the quantity $\mathcal{F}_{a,t}$ admits the expansion

$$\mathcal{F}_{a,t} = \frac{1}{D} \text{Tr} [Q_a e^{-iHt} P_a e^{iHt}] = 4\lambda_a t + \mathcal{O}(t^2), \quad (6)$$

so that the coefficient λ_a is exposed in the first-order term. This is a consequence of the Hadamard formula (see Lemma 5), orthonormality of Pauli operators with respect to the Hilbert-Schmidt inner product, as well as the specific choices of P_a and Q_a . The authors of Ref. [36] envision the Hamiltonian learning problem as one of inverting the non-linear map

$$\mathcal{F}_t: [-1, 1]^M \ni (\lambda_a)_{a \in [M]} \mapsto (\mathcal{F}_{b,t} := \text{Tr} [Q_b e^{-iHt} P_b e^{iHt}])_{b \in [M]} \in \mathbb{R}^M, \quad (7)$$

for all $a \in [M]$. In particular, their protocol is carried out via the following two steps:

1. Computing estimates $\tilde{\mathcal{F}}_{a,t}$ of $\mathcal{F}_{a,t}$ up to some precision ξ , such that $|\tilde{\mathcal{F}}_{a,t} - \mathcal{F}_{a,t}| \leq \xi$, for all $a \in [M]$. This is straightforwardly performed by measuring time-evolved randomized stabilizer states on the observable Q_a .
2. Using estimates $\tilde{\mathcal{F}}_{a,t}$, inverting the map $\mathcal{F}_t$ via the Newton-Raphson method.

Their use of the Newton-Raphson method is necessitated by the fact that the map $\mathcal{F}_t$ is non-linear, in that the forward map from the coefficients $(\lambda_a)_{a \in [M]}$ to $(\mathcal{F}_{a,t})_{a \in [M]}$ involves exponentials in the Hamiltonian. Furthermore, due to their lack of control over t , they are unable to utilize a first-order approximation of $\mathcal{F}_{a,t}$ to arbitrarily small error. However, as we explain later, our setting *will* allow us to take advantage of a linear approximation of $\mathcal{F}_{a,t}$, leading to a drastically simpler Hamiltonian learning algorithm.

IV. LEARNING PROBLEM

In this section, we formalize the central learning problem behind our quantum advantage results. At a high level, the goal of this task is to predict, for various given times and input states, a classical description of these states after they have been time evolved under an unknown Hamiltonian. The training data for this task is generated as follows: for different inputs $x \in \{0, 1\}^{\text{poly}(n)}$ and evolution times $t \in [0, T]$, an n -qubit input state vector $|\psi(x)\rangle$ is prepared and time-evolved by a low-intersection Hamiltonian $H(\lambda)$ for time t , i.e., it evolves by $U(t, \lambda) = e^{-iH(\lambda)t}$. The classical description of this time-evolved state then takes the form of expectation values of a predetermined set of certain observables that are gathered via repeated preparations and measurements. The Hamiltonian parameters $\lambda \in [-1, 1]^M$ are fixed but unknown, with $M \in \mathcal{O}(n)$, which follows from the low-intersection property. In particular, we consider Hamiltonians

$$H(\lambda) = \sum_{a \in [M]} \lambda_a E_a \quad (8)$$

of the kind described in Def. 4. The task of a learner will then be to use these classical snapshots of time-evolved states from polynomially many different (x, t) pairs to generalize to new inputs and evolution times. This learning problem can be rigorously formalized by means of the following concept class.

Definition 6 (Concept class). Let $H := \{H(\lambda) | \lambda \in [-1, 1]^M\}$ be a family of n -qubit low-intersection Hamiltonians. Let $x \in \{0, 1\}^{p(n)}$, where $p: \mathbb{N} \rightarrow \mathbb{N}$ is a function mapping n to some integer and $p(n) \in \mathcal{O}(\text{poly}(n))$. Also, consider $(Q_a)_{a \in [M]}$, the operators defined in Subsection III C 1. Then, we define the concept class

$$\mathcal{C}_{U_{\text{enc}}, T}^H := \left\{ c_\lambda: \{0, 1\}^{p(n)} \times [0, T] \ni (x, t) \mapsto \left(\langle \psi(x) | U^\dagger(t, \lambda) Q_a U(t, \lambda) | \psi(x) \rangle \right)_{a \in [M]} \in \mathbb{R}^M | \lambda \in [-1, 1]^M \right\}, \quad (9)$$

where $U(t, \lambda) := e^{-iH(\lambda)t}$, such that $|\psi(x)\rangle = U_{\text{enc}}(x) |0\rangle^{\otimes n}$, where U_{enc} is a data-encoding unitary and $T \in \mathcal{O}(\text{poly}(n))$.

Note that the concept class is specified by a fixed choice of encoding unitary U_{enc} , the Hamiltonian family $H = \{H(\lambda)\}_\lambda$ under which we perform the time evolution, and a maximal evolution time T . In this work, we are interested in an encoding unitary of a particular kind, whose action is determined by the first bit of the input bit-string x . For some input pair $(x^{(j)}, t^{(j)})$, we define it as

$$U_{\text{enc}}(x^{(j)}) |0^n\rangle := \begin{cases} |\psi(x^{(j)})\rangle = |s^{(j)}\rangle \in \{|0\rangle, |1\rangle, |+\rangle, |-\rangle, |i\rangle, |-i\rangle\}^{\otimes n}, & \text{if } x_1^{(j)} = 0, \\ |\psi(x^{(j)})\rangle = |x_2^{(j)}, \dots, x_{n+1}^{(j)}\rangle, & \text{if } x_1^{(j)} = 1, \end{cases} \quad (10)$$

where, in the case when $x_1^{(j)} = 0$, the choice of the state vector $|s^{(j)}\rangle$ is fully determined by the bit-string $x_2^{(j)}, \dots, x_{p(n)}^{(j)}$. In particular, the state vectors $|s^{(j)}\rangle$ are tensor products of single-qubit stabilizer states and can therefore be described by bitstrings of length $p(n) = 1 + \lceil n \log_2 6 \rceil$, where n is the number of qubits on which the state vector $|\psi(x^{(j)})\rangle$ is defined. In the following, we denote by $s^{(j)}$ the bit-strings $x^{(j)}$ with $x_1^{(j)} = 0$, where the remaining bits $x_2^{(j)}, \dots, x_{p(n)}^{(j)}$ specify the particular stabilizer state vector $|s^{(j)}\rangle$. The goal of the ML algorithm, then, is to learn a model $h(x, t)$ which approximates the unknown concept

$$c_\lambda(x, t) = \left(\langle \psi(x) | U^\dagger(t, \lambda) Q_a U(t, \lambda) | \psi(x) \rangle \right)_{a \in [M]}, \quad (11)$$

given training data of the form $\mathcal{T}^\lambda := \{(x^{(j)}, t^{(j)}, y^{(j)})\}_{j=1}^N$, where $y^{(j)} := c_\lambda(x^{(j)}, t^{(j)})$. On occasion, in the context of quantum learnability, i.e., when $x_1 = 0$, we will also use

$$B_a^{(j)} := \text{Tr} \left[Q_a e^{-iH(\lambda)t} \bigotimes_{i=1}^n |s_i^{(j)}\rangle \langle s_i^{(j)}| e^{iH(\lambda)t} \right] \quad (12)$$

to denote the a^{th} component of the j^{th} training sample, for notational convenience. To fully formalize the learning task, we also need to specify the family of distributions $\{\mathcal{D}_i\}_i$ from which the training data points $(x^{(j)}, t^{(j)})$ are sampled. These distributions will also govern the measure of approximation (or generalization performance) of the hypothesis $h(x, t)$ over the entire space (x, t) . We will consider in particular input distributions of the following form:

Definition 7 (Allowed input distributions). Let $p > 0$ and $\mathcal{D}_i \in \{\mathcal{D}_i\}_i$ be a family of distributions over $(x, t) \in \{0, 1\}^p \times [0, T]$ defined as:

- The time t is uniformly sampled over $[0, T]$.
- The first bit x_1 of x is randomly selected with equal probability between 0 and 1.
- If $x_1 = 0$ then the other $p - 1$ bits $x_2, x_3, \dots, x_p$ are sampled from the uniform distribution over $p - 1$ bit-strings $\text{Unif}(\{0, 1\}^{p-1})$.
- If $x_1 = 1$ then the $p - 1$ bits in $x_2, \dots, x_p$ follow any possible distribution over $\{0, 1\}^{p-1}$. For each i , the distribution over these bit-strings defines the specific overall distribution $\mathcal{D}_i$.

In our setting, we will consider $p = p(n)$.

Formally, our learning condition is defined as follows.

Definition 8 (Efficient PAC-learning condition). *The concept class $\mathcal{C}_{U_{\text{enc}}, T}^H$ is (efficiently) PAC-learnable with respect to the set of distributions $\{\mathcal{D}_i\}_i$ over $(x, t) \in \{0, 1\}^{p(n)} \times [0, T]$ if there exists an algorithm $\mathcal{A}$ such that, for all $\varepsilon, \delta > 0$, for any c_λ in $\mathcal{C}_{U_{\text{enc}}, T}^H$, and any input distribution $\mathcal{D}_i \in \{\mathcal{D}_i\}_i$, $\mathcal{A}$ receives as input a training set $\mathcal{T}^\lambda$ of size $\mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$ and outputs a hypothesis $h(\cdot) = \mathcal{A}(\mathcal{T}^\lambda, \varepsilon, \delta, \cdot)$ which satisfies*

$$\mathbb{E}_{(x,t) \sim \mathcal{D}_i} [\|c_\lambda(x, t) - h(x, t)\|_\infty] \leq \varepsilon \quad (13)$$

with probability greater than $1 - \delta$. The concept class is said to be efficiently PAC-learnable if the time complexity of $\mathcal{A}$ is $\mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$.

V. QUANTUM LEARNABILITY

Having properly formalized the learning task considered in this work, we now provide a quantum learning algorithm that is able to efficiently PAC-learn the concept class described in Def. 6 with respect to a certain family of input distributions. This entails: i) learning the Hamiltonian parameters λ using short-time-evolved states sampled from the $x_1 = 0$ subspace, and ii) guaranteeing that the learned parameters $\hat{\lambda}$ allow time-evolution of new input states up to good approximation error, even for long evolution times $t \sim T$. Recall from the previous section that our training data is of the form

$$\mathcal{T}^\lambda = \left\{ \left(x^{(j)}, t^{(j)}, c_\lambda(x^{(j)}, t^{(j)}) \right) \right\}_{j=1}^N, \quad (14)$$

$$\text{where } c_\lambda(x^{(j)}, t^{(j)}) = \left(\langle \psi(x) | U^\dagger(t, \lambda) Q_a U(t, \lambda) | \psi(x) \rangle \right)_{a \in [M]}. \quad (15)$$

Given such a training dataset as input, our learning algorithm is divided into the following two stages:

1. A training phase, where, using the expectation values provided in $\mathcal{T}^\lambda$, we identify the concept generating the data via Hamiltonian learning. This phase is specified by Algorithm 1.
2. An inference stage, involving evaluation of the learned concept, which is performed by means of Hamiltonian simulation and the classical shadows protocol. This phase is specified by Algorithm 2.

We will proceed by considering the expansion of $\mathcal{F}_{a,t}$ shown in Eq. (6). Following an alternative idea sketched in Ref. [36], we note that one could choose to determine an ε -dependent t such that higher-order terms satisfy an appropriate tail bound, in which case the coefficient becomes recoverable through simply division of an estimate of $\mathcal{F}_{a,t}$ by $4t$. For the scenario considered in this work, where the times contained in the training dataset are sampled from a uniform distribution over $[0, T]$, where $T \in \mathcal{O}(\text{poly}(n))$, we consider the map

$$\tilde{\mathcal{F}}_{t^*} : [-1, 1]^M \ni (\lambda_a)_{a \in [M]} \mapsto \left(\tilde{\mathcal{F}}_{b,t^*} := \frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{b,t} dt \right)_{b \in [M]} \in \mathbb{R}^M \quad (16)$$

(instead of the one defined in Eq. (7)), where t^* is chosen such that

$$\frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{b,t} dt = 2\lambda_a t^* + \underbrace{\dots}_{\leq \varepsilon_{\text{tail}}} \quad (17)$$

can be shown to hold, where $\varepsilon_{\text{tail}}$ will be some allocated error budget. Algorithmically, we will compute the time-averaged integral via Monte-Carlo estimation, which we perform simply by taking the empirical average over training samples contained in the short-time window $[0, t^*] \subset [0, T]$. Dividing the empirical average by $2t^*$ gives us an estimate

Algorithm 1. Training via Hamiltonian learning from short-time randomized probes

Input:

1. Final precision and confidence parameters $\varepsilon, \delta > 0$.
2. A training dataset of the form $\mathcal{T}^\lambda$ of size N as defined in Eq. (14).
3. Knowledge of the maximum degree $\mathfrak{d}$ of the dual-interaction graph $\mathfrak{G}$ of H , as well as the Hamiltonian terms $(E_a)_{a \in [M]}$.

Output: Estimates $\hat{\lambda} := (\hat{\lambda}_a)_{a \in [M]}$ of $\lambda := (\lambda_a)_{a \in [M]}$ such that, for all $a \in [M]$, $|\lambda_a - \hat{\lambda}_a| \leq \varepsilon/2MT$ with probability at least $1 - \delta/2$.

```

1:  $t^* \leftarrow \mathcal{O}\left(\varepsilon/MT \|O\|_\infty (\mathfrak{d} + 1)^2\right)$ , where  $O$  is the observable being predicted
2: for  $a = 1$  to  $M$  do
3:    $S_{a,t^*} \leftarrow 0$  ▷ Initialize sum variable.
4:    $\text{count}_{a,t^*} \leftarrow 0$  ▷ Initialize counter variable.
5:   for  $j = 1$  to  $N$  do ▷ Loop over all data points in training set.
6:     if  $x_1^{(j)} = 0$  then ▷ Check whether sample can be used for learning.
7:       if  $t^{(j)} \in [0, t^*]$  then ▷ Check if data point is contained in short-time window
8:         if  $P_a |s_i^{(j)}\rangle = + |s_i^{(j)}\rangle$  then ▷ Check if  $|s_i^{(j)}\rangle$  is a +1 eigenstate of  $P_a$ .
9:            $S_{a,t^*} \leftarrow S_{a,t^*} + 6B_a^{(j)}$ 
10:        end if
11:        $\text{count}_{a,t^*} \leftarrow \text{count}_{a,t^*} + 1$  ▷ Count number of points in short-time window.
12:     end if
13:   end for
14:   end for
15:    $\hat{\mathcal{F}}_{a,t^*} \leftarrow S_{a,t^*} / \text{count}_{a,t^*}$  ▷ Compute empirical average.
16:    $\hat{\lambda}_a \leftarrow \hat{\mathcal{F}}_{a,t^*} / 2t^*$  ▷ Invert the map  $\hat{\lambda}_a \mapsto \hat{\mathcal{F}}_{a,t^*}$  via division by  $2t^*$ .
17: end for
18: return  $(\hat{\lambda}_a)_{a \in [M]}$  ▷ Return estimates of Hamiltonian coefficients.

```

$\hat{\lambda}_a$ of λ_a , up to some precision controlled by $\varepsilon_{\text{tail}}$ and the accuracy of Monte-Carlo estimation. We then calculate the sample complexity of short-time samples, which we denote by N_{t^*} . We do this by bounding the difference between the Monte-Carlo estimate and the time-averaged integral, which we show to be the expectation of the random variable from which the training samples are drawn, conditioned on $t \in [0, t^*]$, thereby allowing us to use Hoeffding's inequality.

Finally, we guarantee that enough training samples are contained in the short-time window, as well as in the $x_1^{(j)} = 0$ subspace, so as to ensure efficient execution of the training algorithm. To address these requirements, we use the Chernoff bound to compute the total sample complexity N , and find that we incur only a polynomial blow-up with respect to N_{t^*} .

The sample complexity of the inference stage is determined by the cost of implementing the classical shadows protocol, which is logarithmic in the number of observables being predicted and the inverse of the failure probability, and inverse-polynomial in the desired precision. We also find that the overall runtime of our learning procedure, a detailed analysis of which we defer to Appendix C, is polynomial in the relevant parameters. We now give an informal statement of our main quantum learnability result.

Theorem 1 (Quantum learnability, informal). *There exists a family of distributions, namely, the one defined in Def. 7, such that, with respect to all distributions in that family, for precision and confidence parameters $\varepsilon, \delta > 0$, the concept class in Def. 6 is efficiently PAC-learnable using a quantum learning algorithm whose sample- and time-complexity scales as $\mathcal{O}(\text{poly}(n, 1/\varepsilon, 1/\delta))$.*

Proof-sketch. The above theorem is formally restated as Theorem 5, which may be found in Subsection C 4 of Appendix C. In order to show that we can efficiently PAC-learn the concept class using Algorithms 1 and 2, we need to prove that the overall procedure for solving the learning task is both sample- and time-efficient. Our sample complexity analysis proceeds by analyzing the error decomposition

$$\left| \hat{\lambda}_a - \lambda_a \right| = \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| = \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} + \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} \right| + \left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right|, \quad (19)$$

where $(\hat{\lambda}_a)_{a \in [M]}$ are the coefficients retrieved in the training phase, $(\lambda_a)_{a \in [M]}$ are the actual coefficients corresponding to the ground-truth Hamiltonian, and $\hat{\mathcal{F}}_{a,t^*}$ is the Monte-Carlo estimator used to approximate the time-averaged integral

Algorithm 2. Inference via Hamiltonian simulation and classical shadows

Input:

1. Final precision and confidence parameters $\varepsilon, \delta > 0$.
2. Learned coefficients $(\hat{\lambda}_a)_{a \in [M]}$ via Algorithm 1.
3. Inference point $(x, t) \sim \text{Unif}(\{0, 1\}^{p(n)} \times [0, T])$.
4. The Hamiltonian terms $(E_a)_{a \in [M]}$.
5. The set of observables $(Q_a)_{a \in [M]}$.
6. Sample-access to the uniform distribution over the ensemble of unitaries $\mathcal{U} := \{\mathbb{1}, H, HS^\dagger\}^{\otimes n}$.

Output: Estimates $(\langle \hat{Q}_a \rangle_{\rho(x,t)})_{a \in [M]}$ such that, for all $a \in [M]$, $|\langle \hat{Q}_a \rangle_{\rho(s,t)} - \langle Q_a \rangle_{\rho(s,t)}| \leq \varepsilon$ with probability at least $1 - \delta/2$.

- 1: $\hat{H} \leftarrow \sum_{a \in [M]} \hat{\lambda}_a E_a$
 - 2: Set $G \leftarrow \lceil 2 \log(4M/\delta) \rceil$ ▷ Number of classical shadow segments.
 - 3: Set $\tilde{N}$ to be the following: ▷ Number of copies of time-evolved states required for classical shadows protocol.
- $$\tilde{N} \leftarrow \frac{136G}{\varepsilon^2} \max_{a \in [M]} \underbrace{\left\| Q_a - \frac{\overbrace{\text{Tr}[Q_a]}^{=0}}{2^n} \mathbb{1} \right\|_{\text{shadow}}^2}_{\leq 4 \cdot 3^{|\text{supp}(E_a)|} \text{ (for Pauli measurements)}} \quad (18)$$
- 4: **for** $\ell = 1$ to $\tilde{N}$ **do** ▷ Loop over all copies for data acquisition phase.
 - 5: $\rho_{\text{init}}^{(\ell)}(x) \leftarrow U_{\text{enc}}(x) |0^n\rangle\langle 0^n| U_{\text{enc}}(x)^\dagger$ ▷ Encoded input state preparation.
 - 6: $\hat{\rho}^{(\ell)}(x, t) \leftarrow e^{-i\hat{H}t} \rho_{\text{init}}^{(\ell)}(x) e^{i\hat{H}t}$ ▷ Hamiltonian simulation.
 - 7: **end for** ▷ Obtain $\tilde{N}$ copies of the time-evolved probe state.
 - 8: Run Subroutine 1 on the states $(\hat{\rho}^{(\ell)}(x, t))_{\ell=1}^{\tilde{N}}$ ▷ Data acquisition phase of classical shadow protocol, outputs $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$.
 - 9: Run Subroutine 2 using $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$, $(Q_a)_{a \in [M]}$, and G as inputs ▷ Median-of-means estimation of expectation values on $(Q_a)_{a \in [M]}$.
 - 10: **return** $(\langle \hat{Q}_a \rangle_{\rho(s,t)})_{a \in [M]}$ ▷ Predicted expectation values.
-

$\tilde{\mathcal{F}}_{a,t^*}$. As we have already discussed, we bound the difference between $\hat{\mathcal{F}}_{a,t^*}/2t^*$ and $\tilde{\mathcal{F}}_{a,t^*}/2t^*$ using Hoeffding's inequality, and that between $\tilde{\mathcal{F}}_{a,t^*}/2t^*$ and λ_a via tail bounds on $\mathcal{F}_{a,t^*}$. This is followed by an application of the Chernoff bound to compute the total sample complexity, which, for precision and failure probability $\varepsilon, \delta > 0$ respectively, we find to be polynomial in n and $1/\varepsilon$, and logarithmic in $1/\delta$. The sample complexity of the inference stage is dominated by the cost of running the classical shadows formalism, which is inverse-polynomial in ε and logarithmic in $1/\delta$ and the number of observables being predicted. The time complexity we show to be polynomial in the relevant parameters as well. The training stage involves elementary constant-time arithmetic operations for each data sample and observable Q_a , whereas the runtime of the deployment phase is dictated by the cost of Hamiltonian simulation, which, for polynomial evolution time, is also polynomial, and that of the classical shadows protocol, which involves preparation of logarithmically many copies of time-evolved states, followed by classical polynomial-time post-processing. $\square$

VI. QUANTUM ADVANTAGE

In this section, we instantiate the learning problem described by the concept class in Def. 6 in a way that shows a quantum learning advantage. Specifically, we define a family of input distributions $\mathcal{D}_i \in \{\mathcal{D}_i\}_i$ and Hamiltonians $H(\lambda)$ for which no classical algorithm can satisfy the learning condition in Def. 8. Together with the fact that the quantum algorithm described in Section V can still learn the particular instantiation described here, this establishes a quantum learning advantage. Our construction relies on mainly two tools: i) a Hamiltonian family $\{H_{\text{BQP}}(\lambda)\}_\lambda$ due to Oliveira and Terhal [37] that satisfies the low intersection property and whose time-evolution for time $T_{\text{BQP}} \in \mathcal{O}(n)$ is

BQP-complete with high probability, ii) and the seminal boosting result of Schapire [93] that allows us to promote any efficient classical PAC-learner under our choice of input distributions to polynomial-size classical circuits for deciding BQP languages (i.e., that would imply $\text{BQP} \subseteq \text{P/poly}$).

A. Classical hardness

We start by stating our classical hardness result, which is based on the widely-believed conjecture that $\text{BQP} \not\subseteq \text{P/poly}$, and provide a proof-sketch. A detailed proof is given in Appendix B.

Theorem 2 (Classical hardness, informal). *For any (Promise) BQP-complete language, there exists a Hamiltonian family H_{BQP} specifying a concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ as per Def. 6 and a family of input distributions following the specification of Def. 7 such that no randomized polynomial-time classical algorithm $\mathcal{A}_C$ satisfies the PAC-learning condition of Def. 8 for this concept class, even when restricted to a constant learning error $\varepsilon < 1/48$, unless $\text{BQP} \subseteq \text{P/poly}$.*

Proof-sketch. Here, we summarize the main argument behind the proof. We begin by defining a Hamiltonian family $H_{\text{BQP}} = \{H_{\text{BQP}}(\lambda)\}_{\lambda}$ such that BQP computations can be embedded into the Hamiltonian evolution of $H_{\text{BQP}}(\lambda)$ by properly choosing λ . Let $\mathcal{L}$ be a (Promise) BQP-complete language over input bit-strings $z \in \{0, 1\}^m$. Since $\mathcal{L} \in \text{BQP}$, there exists a polynomial-depth quantum circuit U_{BQP} deciding $\mathcal{L}$, in the sense that

- if $z \in \mathcal{L}$, then applying U_{BQP} to $|z\rangle$ outputs 1 with probability at least $2/3$;
- if $z \notin \mathcal{L}$, then applying U_{BQP} to $|z\rangle$ outputs 0 with probability at least $2/3$.

Using the construction of Ref. [37], one obtains a low-intersection local Hamiltonian H_{BQP} acting on $n \in \mathcal{O}(\text{poly}(m))$ qubits whose real-time evolution for polynomially or linearly long times $T_{\text{BQP}} \in \mathcal{O}(n)$ implements the circuit U_{BQP} with high probability (see Appendices B 1 and B 2). The key to making this Hamiltonian low-intersection compared to the standard Feynman/Kitaev construction is to expand the simulated circuit U_{BQP} over $\mathcal{O}(\text{poly}(m))$ many qubits (or, more precisely, mK qubits for a depth K circuit) such that any qubit involved in this circuit gets only acted upon by a constant number of gates. In the Feynman-Kitaev construction, this translates into each qubit being involved in only a constant number of local Hamiltonian terms, i.e., the Hamiltonian satisfying the low-intersection property. Moreover, the construction can be easily modified so that H_{BQP} simulates U_{BQP} with arbitrarily high probability for a randomly sampled evolution time in $[0, T_{\text{BQP}}]$. Expanding each local term of H_{BQP} in the Pauli basis gives a Hamiltonian of the form in Eq. (8). The Hamiltonian family $H_{\text{BQP}} = \{H_{\text{BQP}}(\lambda)\}_{\lambda}$ is then defined by considering all Pauli expansions resulting from this construction, for arbitrary circuits U_{BQP} . In this way, the Pauli coefficients $\lambda \in [-1, 1]^M$ specify the particular Hamiltonian H_{BQP} that simulates a given circuit U_{BQP} .

We now relate this Hamiltonian evolution to the decision problem for $\mathcal{L}$. For inputs x with $x_1 = 1$, we identify the bits $z = x_2 \cdots x_{m+1}$ with an input to the BQP computation. Measuring the observable $O = Z \otimes I \otimes \cdots \otimes I$ on the state obtained after the evolution generated by H_{BQP} distinguishes whether $z \in \mathcal{L}$ or $z \notin \mathcal{L}$. More precisely, the concept associated with λ contains a component of the form

$$(c_{\lambda}(x, t))_{i^*} = \text{Tr} \left[O e^{-iH_{\text{BQP}}(\lambda)t} \rho(x) e^{iH_{\text{BQP}}(\lambda)t} \right], \quad (20)$$

where i^* is the index associated with component corresponding to O . Averaging this component over $t \in [0, T_{\text{BQP}}]$, and restricting our attention to the subset of inputs with $x_1 = 1$, we find that the resulting function

$$\bar{q}(z) = \frac{1}{T} \int_{t=0}^T (c_{\lambda}((\mathbf{1}, z, \mathbf{0}), t))_{i^*} \quad (21)$$

approximates the BQP computation deciding $\mathcal{L}$ up to good margin for all $z \in \{0, 1\}^m$. Armed with this observation, the only difficulty from here onward is that a classical learner achieving non-trivial *average* performance with respect to some input distributions $\{\mathcal{D}_i\}_i$ does not immediately guarantee that it can solve the BQP problem. We need additional machinery to leverage this hypothetical efficient classical learner to construct an algorithm in P/poly that can decide the BQP language.

Suppose now, toward a contradiction, that there were a randomized polynomial-time classical algorithm satisfying the learning condition for the concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ with small overall error $\varepsilon \in \mathcal{O}(1)$. Applying this learner to the concept

specified by λ would produce, with high probability, a hypothesis that approximates the component $(c_\lambda(x, t))_{i^*}$ with average error $\mathcal{O}(\varepsilon)$ over the subspace $x_1 = 1$ and $t \in [0, T_{\text{BQP}}]$. By further averaging the hypothesis over the time window $t \in [0, T_{\text{BQP}}]$ and suitably setting the threshold error ε to $1/48$, one obtains a classical hypothesis that correctly decides $\mathcal{L}$ with probability strictly more than $1/2$ with respect to any distribution μ over $z \in \{0, 1\}^m$.

The final step uses the result of Ref. [93], which states that if a function is PAC-learnable under arbitrary distributions, then it can be computed by a polynomial-size classical circuit on every input. We recall this result and explain its consequences in Appendix B. Applying it here would imply the existence of a polynomial-size classical circuit deciding the (Promise) BQP-complete language $\mathcal{L}$. Therefore $\text{BQP} \subseteq \text{P/poly}$, contradicting the assumed hardness. Hence no randomized polynomial-time classical learner can satisfy the learning condition for $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$, unless $\text{BQP} \subseteq \text{P/poly}$. $\square$

B. Quantum learnability for a classically hard case

Recall that the objective of our quantum learner is to output an estimate $\hat{\lambda}$ such that the hypothesis function h , defined by $h(x, t) = c_{\hat{\lambda}}(x, t)$, satisfies the learning condition in Eq. (13). For this, the learning algorithm devised in the last section (Algorithm 1) is directly applicable, as per Theorem 1. An interesting observation however is that, in the case of our Hamiltonian family H_{BQP} , it is sufficient to recover an $\hat{\lambda}$ such that $\|\hat{\lambda} - \lambda\|_\infty \leq \varepsilon \in \mathcal{O}(1)$ in order to recover $H_{\text{BQP}}(\lambda)$ exactly. In turn, this means that we can learn the concept class perfectly, i.e., with *zero* generalization error, and that too using less samples than stated in Theorem 1.

Lemma 1 (Quantum learnability of $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$). *Consider the learning problem corresponding to the concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$ that is defined by the Hamiltonian family H_{BQP} presented in Theorem 2. Then, Algorithm 1 is able to meet the learning condition in Def. 8 with zero error, i.e., $h(\cdot) = \mathcal{A}(\mathcal{T}^\lambda, \varepsilon, \delta, \cdot)$ is such that with probability $1 - \delta$,*

$$h(x, t) = c_\lambda(x, t) \quad \forall (x, t) \in \{0, 1\}^{p(n)} \times [0, T], \quad (22)$$

for any concept c_λ in $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$. Moreover, the number of samples required by $\mathcal{A}$ is $N_{\text{BQP}} \in \mathcal{O}(T \log(M/\delta))$, and the time-complexity of $\mathcal{A}$ is given by $\mathcal{O}(MN_{\text{BQP}})$.

Proof. The proof follows from Lemmas 2 and 13 in the Appendix. In particular, Lemma 2 guarantees that the Hamiltonian associated with the classically hard instance can be written as a 5-local, low-intersection Hamiltonian $H_{\text{BQP}}(\lambda) = \sum_{a \in [M]} \lambda_a E_a$, where $\{E_a\}_{a \in [M]}$ are local Pauli strings and the coefficients λ_a belong to a fixed finite set $\Lambda \subset \mathbb{R}$ that is independent of the system size. Hence, to reconstruct $H_{\text{BQP}}(\lambda)$ exactly, it suffices to estimate each coefficient with precision smaller than half the minimum separation between any two distinct admissible values. In the Pauli expansion provided in Lemma 2, two identical Pauli terms can appear in H_{BQP} only in the case we consider the Z -dependent part of the gates

$$T = \left(\frac{1 + e^{i\pi/4}}{2} \right) I + \left(\frac{1 - e^{i\pi/4}}{2} \right) Z \quad (23)$$

and $H = (X + Z)/\sqrt{2}$. In this case, the difference of the coefficients is

$$\varepsilon_P = \left| \frac{1 - e^{i\pi/4}}{2} - \frac{1}{\sqrt{2}} \right| \approx 0.66. \quad (24)$$

After accounting for the normalization factors introduced by the clock-register terms and the two SWAP operations, the minimum separation between two distinct admissible coefficients is lower-bounded by $\varepsilon'_P = \varepsilon_P/32$. We may therefore run $\mathcal{A}$ with coefficient-estimation precision strictly smaller than $\varepsilon'_P/2$. With probability at least $1 - \delta$, each estimated coefficient $\hat{\lambda}_a$ can then be uniquely rounded to the corresponding exact value λ_a . Consequently, $H_{\text{BQP}}(\hat{\lambda}) = H_{\text{BQP}}(\lambda)$, and the hypothesis returned by the learner coincides with the target concept

$$h(x, t) = c_\lambda(x, t) \quad \forall (x, t) \in \{0, 1\}^{p(n)} \times [0, T] \quad (25)$$

for every input and every evolution time. Moreover, the ability to learn the Hamiltonian exactly also has an impact on the sample complexity of our learning algorithm. Indeed, unlike the general case considered in Theorem 1, the error propagation in Hamiltonian simulation (see Lemma 8) is trivially accounted for, which significantly reduces the sample complexity of Algorithm 1. As detailed in Lemma 13, we can get away with a sample complexity $N_{\text{BQP}} \in \mathcal{O}(T \log(M/\delta))$, as opposed to

$$N \in \mathcal{O}\left(\frac{M^5 T^6 \|O\|_\infty^3 (\mathfrak{d} + 1)^6}{\varepsilon^5} \log\left(\frac{M}{\delta}\right)\right) \quad (26)$$

in the general case of Theorem 1, which gives $N = \mathcal{O}(M^5 T^6 \log(M/\delta))$ upon plugging in the scaling $\varepsilon \in \mathcal{O}(1)$. $\square$

VII. DISCUSSION

A. Distinction between the present work and earlier work

We recall the PAC-learnable version of the concept class in Ref. [24]:

Definition 9 (PAC-learnable Hamiltonian dynamics concept class, Def. 4, [24]). *Consider a sequence of n -qubit Hamiltonians $\{H_n(x, \lambda)\}_{n \in \mathbb{N}}$, such that each Hamiltonian $H_n(x, \lambda)$ is parametrized by a bit-string $x \in \{0, 1\}^{p(n)}$ of length $p(n)$, where $p: \mathbb{N} \rightarrow \mathbb{N}$, and $\lambda \in [0, 1]^{d(n)}$, a vector of continuous, unknown parameters, with $d(n) \in \mathcal{O}(\log n)$. For $\{O_n\}_{n \in \mathbb{N}}$ a sequence of observables, consider the concept class*

$$\mathcal{C}_{n,d}^H := \left\{ c_\lambda: \{0, 1\}^{p(n)} \ni x \mapsto \langle 0 | U_n^\dagger(x, \lambda) O_n U_n(x, \lambda) | 0 \rangle \in \mathbb{R} \mid \lambda \in [0, 1]^{d(n)} \right\}, \quad (27)$$

where $U_n(x, \lambda) := e^{-iH_n(x, \lambda)\tau}$, where $\tau \in \mathbb{R}_{>0}$ is a fixed, positive real number.

Let us enumerate the key differences between the work of Barthe et al. and that of our paper. To begin with, in Ref. [24], for $n \in \mathbb{N}$, the unknown part of the Hamiltonian $H_n(x, \lambda)$, i.e., the part parametrized by λ , can contain only logarithmically many terms, since the dimensionality of λ must be $d(n) \in \mathcal{O}(\log n)$. This is in contrast with our setting, since we impose no such restriction. In particular, for our learning task, the unknown Hamiltonian generating the training data satisfies the low-intersection property described in Def. 5, which leads to the number of constituent terms being $\mathcal{O}(n)$, a scenario that encompasses many-body systems of the kind typically encountered in condensed matter physics.

Another difference between our work and Ref. [24] lies in the form of the training data considered. Specifically, each of our training examples is labeled by a list of expectation values on the observables $(Q_a)_{a \in [M]}$, whereas the authors of Ref. [24] require only a single expectation value on some arbitrary observable O per label. In this sense, our training data is richer and potentially more informative, which leads to a methodological difference between our work and Ref. [24]—the stronger access model enables us to learn the unknown Hamiltonian generating the data, which in turn allows us to predict a wider range of observables than in Ref. [24]. Put differently, our training procedure approximately identifies the concept generating the training data by learning the underlying Hamiltonian itself. Our approach is thus a form of proper PAC-learning, where the learner is required to output a hypothesis that is contained in the concept class underlying the learning problem. By contrast, in Ref. [24], the hypothesis class differs from the concept class described in Def. 9; in particular, it is given by a family of generalized trigonometric polynomials associated with the parametrized Trotter circuit implementing the time-evolution operator $U_n(x, \lambda) = e^{-iH_n(x, \lambda)t}$.

$$\mathcal{H}_{n,d}^{\text{PQC}} := \left\{ h_\lambda: \{0, 1\}^{p(n)} \ni x \mapsto \sum_{\ell \in \mathcal{L}} b_\ell(x) e^{i\langle \lambda, \ell \rangle} \in \mathbb{R} \mid \lambda \in [0, 1]^{d(n)} \right\}, \quad (28)$$

where $\mathcal{L}$ is the frequency spectrum of the parametrized quantum circuit (PQC) and $|\mathcal{L}| = d(n) \in \mathcal{O}(\log n)$.

It has been observed in Ref. [24] that the aforementioned logarithmic scaling of $d(n)$ is a fundamental limitation imposed by complexity theory. Specifically, this follows from the results of Arunachalam et al. [38], which showed conditional quantum hardness for learning certain classes of shallow classical circuits. More precisely, they define the following concept class:

$$\mathcal{C}_n \subseteq \{c: \{0, 1\}^n \rightarrow \{0, 1\}\}, \quad (29)$$

where $\mathcal{C}_n$ corresponds to the circuit classes TC^0 , AC^0 , and TC_2^0 . Then, working within the framework of Ref. [94], they consider a *quantum* PAC-learner $\mathcal{A}$, which is given examples of the form

$$|\psi_{c,\mathcal{D}}\rangle := \sum_{x \in \{0,1\}^n} \sqrt{\mathcal{D}(x)} |x, c(x)\rangle, \quad (30)$$

where $\mathcal{D}: \{0,1\}^n \rightarrow [0,1]$ is an unknown distribution, taken to be fixed in Ref. [38] as opposed to arbitrary, and the learning criteria to be satisfied by $\mathcal{A}$ are analogous to those in Def. 1. Additionally provided in their setting is access to a quantum membership oracle $O_c: |x, b\rangle \mapsto |x, b \oplus c(x)\rangle$, where $b \in \{0,1\}$. Note that the quantum PAC-learning framework assumes stronger access to training data than its classical counterpart, since examples are provided in the form of the coherent superposition $\sum_x \sqrt{\mathcal{D}(x)} |x, c(x)\rangle$, so that one exactly recovers the ordinary classical example $(x, c(x))$, with $x \sim \mathcal{D}$, via measurement in the computational basis. Consequently, one can simulate classical PAC-learning within this framework, and it follows that hardness results for the quantum PAC-learning setting also apply to classical PAC-learners.

We briefly recall the hardness results of Ref. [38]. There, assuming that *ring-learning with errors* (RLWE) cannot be solved in quasi-polynomial-time using quantum computers [95], the authors showed that efficient (weak) quantum PAC-learners do not exist for polynomial-size TC^0 circuits. Furthermore, under stronger sub-exponential RLWE assumptions, polynomially sized AC^0 circuits were also shown to be unlearnable using $n^{\mathcal{O}(\log^\nu n)}$ -time quantum PAC-learners, with ν being some constant. Additionally, under LWE hardness, depth-2 circuits contained in TC_2^0 were found to not be weakly quantum PAC-learnable in polynomial-time. A sufficiently general PQC, i.e., one with polynomially many unknown parameters, or a Hamiltonian containing $\mathcal{O}(n)$ -many terms whose time-evolution can be compiled into such a circuit, is capable of encoding these shallow circuit classes through a single expectation value obtained by measuring Z_1 on the output qubit. Hence, an efficient quantum learner capable of solving the learning task of Ref. [24] in full generality would violate the hardness results of Ref. [38]. Our work avoids this obstruction by considering a different form of training data, in which the labels are given by a richer vector of expectation values chosen to enable Hamiltonian learning.

Finally, in Def. 9, it should be noted that the unitary $U_n(x, \lambda) := e^{-iH_n(x, \lambda)\tau}$ corresponds to time-evolution by a fixed time τ , so that the training data considered in Ref. [24] are expectation values of states that have been time-evolved by τ . This is unlike the setting considered in our work, since our training samples correspond to a uniform distribution over the interval $[0, T]$, with $T \in \mathcal{O}(\text{poly}(n))$.

B. Summary and future directions

To summarize, in this work, we have formulated a physically motivated learning problem in which the goal is to predict expectation values of quantum states that have been time-evolved by an unknown, low-intersection Hamiltonian. Assuming $\text{BQP} \not\subseteq \text{P/poly}$ to hold, we prove that the aforementioned task demonstrates an exponential quantum-classical learning separation with respect to a certain family of input distributions. We show the existence of an efficient quantum learner, whose training phase involves a simplified version of the Hamiltonian learning protocol of Ref. [36], while its inference stage relies on Hamiltonian simulation and the classical shadows formalism. On the other hand, to prove classical hardness, we use a low-intersection variant of the Feynman-Kitaev clock Hamiltonian, whose polynomial-time dynamics encode BQP-complete computations. In this sense, our work shows a proper quantum-classical separation for a meaningful learning task that is naturally motivated by a plausible physical application. We finally conclude by listing a few compelling future directions and questions that have been left open in our work.

Applicability of long-time Hamiltonian learning. Currently, our approach in this work relies on short-time samples to perform Hamiltonian learning in the training phase. One could also consider whether the unknown Hamiltonian can be learned within a setting involving long-time samples only. In this regard, we highlight the recent works of Refs. [88, 96]. The former shows that Heisenberg-limited Hamiltonian learning can be performed even without access to short-time control; however, their method interleaves queries to long-time Hamiltonian evolution with certain learned auxiliary operations, which does not correspond exactly to our setting.

The latter work of Ref. [96] studies the recoverability of local Hamiltonians from arbitrarily long times, and provides *average-case identifiability* guarantees over certain local, random ensembles of Hamiltonians. However, whether their methods apply to the PAC-learning setting considered here, where the concept class may contain instances corresponding to worst-case Hamiltonian instances such as those in H_{BQP} , remains unclear. Nonetheless, we view the question of whether long-time Hamiltonian learning algorithms akin to Refs. [88, 96] can be applied to our supervised learning problem as an interesting future direction and worthy of further study.

Extensions to other physical settings of interest. Obvious extensions of our work would be to show learning separations for predicting ground and thermal state properties. In particular, for the analogous scenario involving Gibbs states, it is foreseeable that a similar approach could hold for some distribution over inverse temperatures $\beta \sim [\beta_{\text{low}}, \beta_{\text{high}}]$, with high (or constant) temperatures corresponding to samples enabling Hamiltonian learning [36, 65, 97] and low temperatures ensuring classical hardness [98–100]. Other interesting settings could possibly include states evolved by processes more general than unitary time-evolution, such as Lindbladians, for which learning algorithms also exist [101–106].

Identification-based learning separations. We stress here again that a limitation of our construction lies in the fact that our labels are vector-valued lists of expectation values, corresponding to sufficient statistics from which the unknown Hamiltonian can be recovered via elementary classical post-processing. Thus, our learner obtains an *explicit* description of the target concept in a manner that is entirely classical. Therefore, our learning separation does not arise from the hardness of identifying the data-generating concept. Instead, the classical hardness enters at the evaluation stage: even after the unknown Hamiltonian has been identified in a proper PAC sense, evaluating the learned concept on potentially long-time inputs requires simulating time-dynamics that may encode BQP-complete computations. Thus, a pressing open direction is the formulation of physically inspired learning tasks similar to the one considered in the present work, where the learning separation stems from the hardness of identification as opposed to evaluation.

Another future direction of interest could thus be to consider a similar learning task to that considered in this work, but one involving more frugal training data and possibly involving a combination of the setting of the present work with the PQC learning-based approach of Ref. [24]. Alternatively, one could contrive constructions wherein only part of the Hamiltonian is identifiable from short-time dynamics, with additional parameters or auxiliary subsystems “switching on” only after long evolution times. Such settings could potentially exhibit identification-based learning separations.

Certified quantum simulation. One related interpretation of our results is in learning-assisted certified quantum simulation. Rather than certifying long-time quantum dynamics directly, we learn the governing Hamiltonian from experimentally accessible short-time data and use it for long-time quantum prediction, obtaining rigorous learning-theoretic and complexity-theoretic guarantees.

Connections to cryptographic notions. Another intriguing implication of our work arises in the context of quantum cryptography. While our result does not provide a construction of one-way functions, it identifies a physically motivated family of quantum-generated prediction tasks that are efficiently solvable on a quantum computer while remaining inaccessible to efficient classical learners. In this sense, our work lies near the conceptual boundary explored in the emerging area of Microcrypt [107]: quantum processes may give rise to useful asymmetric computational tasks without necessarily implying classical one-way functions.

Complexity-theoretic no-gos. Finally, an exciting direction would be to place fundamental limitations on the learnability of time-dynamics. That is, it would be of interest to explore an intermediate regime between labels that are singular expectation values, as in Ref. [24], and ones that amount to maximally informative classical descriptions of probe states. In other words, one could try to prove hardness results of the type shown in Ref. [38], while considering different generalizations of their setting, such as input states that go beyond computational basis states, like stabilizer states, Choi states, and so on, as well as measurement of observables more general than single-qubit observables.

VIII. ACKNOWLEDGMENTS

This project was funded by the European Union. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them. This work is supported by ERC grant BeMAIQuantum (grant agreement No. 101124342; project DOI: 10.3030/101124342) and the Dutch National Growth Fund (NGF) as part of the QDNL programme. This publication is part of the NWA research program Research on Routes by Consortia (ORC), financed by the Dutch Research Council (NWO), through the projects “Divide & Quantum” (Project No. 1389.20.241) and “Quantum Inspire the Dutch Quantum Computer in the Cloud” (Project No. NWA.1292.19.194). The Berlin team would like to acknowledge support from the BMFTR (MUNIQC-Atoms, QuSol, Hybrid++), BIFOLD, the DFG (CRC 183 and SPP 2514), the Quantum Flagship (Millenion, PasQuanS2), the Munich Quantum Valley, Berlin

Quantum, and the European Research Council (DebuQC) for financial support.

-
- [1] P. W. Shor, *Proc. 50th Ann. Symp. Found. Comp. Sc.*, 124 (1994).
 - [2] J. Eisert and J. Preskill, *arXiv* (2026), [arXiv:2510.19928](#).
 - [3] H.-Y. Huang, S. Choi, J. R. McClean, and J. Preskill, *arXiv* (2025), [arXiv:2508.05720](#).
 - [4] N. Pirnay, V. Uhlitzsch, F. Wilde, J. Eisert, and J.-P. Seifert, *Science Adv.* **10**, eadj5170 (2024).
 - [5] S. P. Jordan, N. Shutty, M. Wootters, A. Zalcman, A. Schmidhuber, R. King, S. V. Isakov, T. Khattar, and R. Babbush, *Nature* **646**, 831836 (2025).
 - [6] A. Abbas, A. Ambainis, B. Augustino, A. Bärttschi, H. Buhrman, C. Coffrin, G. Cortiana, V. Dunjko, D. J. Egger, B. G. Elmegreen, N. Franco, F. Fratini, B. Fuller, J. Gacon, C. Gondiulea, S. Gribling, S. Gupta, S. Hadfield, R. Heese, G. Kircher, T. Kleinert, T. Koch, G. Korpas, S. Lenk, J. Marecek, V. Markov, G. Mazzola, S. Mensa, N. Mohseni, G. Nannicini, C. O’Meara, E. P. Tapia, S. Pokutta, M. Proissl, P. Rebentrost, E. Sahin, B. C. B. Symons, S. Tornow, V. Valls, S. Woerner, M. L. Wolf-Bauwens, J. Yard, S. Yarkoni, D. Zechiel, S. Zhuk, and C. Zoufal, *Nature Rev. Phys.* **6**, 718 (2024).
 - [7] R. Sweke, J.-P. Seifert, D. Hangleiter, and J. Eisert, *Quantum* **5**, 417 (2021).
 - [8] Y. Liu, S. Arunachalam, and K. Temme, *Nature Phys.* **17**, 1013 (2021).
 - [9] R. Molteni, C. Gyurik, and V. Dunjko, *npj Quant. Inf.* **12**, 19 (2026).
 - [10] H.-Y. Huang, R. Kueng, and J. Preskill, *Phys. Rev. Lett.* **126**, 190505 (2021).
 - [11] J. Liu, M. Liu, J.-P. Liu, Z. Ye, Y. Wang, Y. Alexeev, J. Eisert, and L. Jiang, *Nature Comm.* **15**, 434 (2024).
 - [12] S. Lloyd, *Science* **273**, 1073 (1996).
 - [13] I. Georgescu, S. Ashhab, and F. Nori, *Rev. Mod. Phys.* **86**, 153185 (2014).
 - [14] A. M. Childs, D. Maslov, Y. Nam, N. J. Ross, and Y. Su, *Proc. Natl. Ac. Sc.* **115**, 94569461 (2018).
 - [15] A. J. Daley, I. Bloch, C. Kokail, S. Flannigan, N. Pearson, M. Troyer, and P. Zoller, *Nature* **607**, 667 (2022).
 - [16] S. Flannigan, N. Pearson, G. H. Low, A. Buyskikh, I. Bloch, P. Zoller, M. Troyer, and A. J. Daley, *Quant. Sc. Tech.* **7**, 045025 (2022).
 - [17] R. Trivedi, A. Franco Rubio, and J. I. Cirac, *Nature Comm.* **15**, 6507 (2024).
 - [18] V. Kashyap, G. Styliaris, S. Mouradian, J. I. Cirac, and R. Trivedi, *Phys. Rev. X* **15**, 021017 (2025).
 - [19] J. Rao, J. Eisert, and T. Guaita, *arXiv* (2025), [arXiv:2510.08467](#).
 - [20] Y. Wang, R. Jiang, Y. Fan, X. Jia, J. Eisert, J. Liu, and J.-P. Liu, *arXiv* (2025), [arXiv:2502.14252](#).
 - [21] C. Gyurik and V. Dunjko, *arXiv* (2024), [arXiv:2306.16028](#).
 - [22] R. Molteni, C. Gyurik, and V. Dunjko, *npj Quant. Inf.* **12**, 19 (2026).
 - [23] R. Molteni, S. C. Marshall, and V. Dunjko, *arXiv* (2026), [arXiv:2504.15964](#).
 - [24] A. Barthe, M. Y. Rad, M. Grossi, and V. Dunjko, *arXiv* (2025), [arXiv:2506.17089](#).
 - [25] V. Bokov, L. Kohl, S. Schmitt, and V. Dunjko, *arXiv* (2026), [arXiv:2601.22006](#).
 - [26] R. P. Feynman, *Found. Phys.* **16**, 507 (1986).
 - [27] H.-Y. Huang, R. Kueng, G. Torlai, V. V. Albert, and J. Preskill, *Science* **377**, 6613 (2022).
 - [28] L. Lewis, H.-Y. Huang, V. T. Tran, S. Lehner, R. Kueng, and J. Preskill, *Nature Comm.* **15**, 895 (2024).
 - [29] M. Wanner, L. Lewis, C. Bhattacharyya, D. Dubhashi, and A. Gheorghiu, *arXiv* (2024), [arXiv:2405.18489](#).
 - [30] C. Rouzé, D. Stilck França, E. Onorati, and J. D. Watson, *Nature Comm.* **15**, 7755 (2024).
 - [31] J. I. Cirac and P. Zoller, *Nature Phys.* **8**, 264 (2012).
 - [32] I. Bloch, J. Dalibard, and S. Nascimbene, *Nature Phys.* **8**, 267 (2012).
 - [33] I. M. Georgescu, S. Ashhab, and F. Nori, *Rev. Mod. Phys.* **86**, 153 (2014).
 - [34] S. Trotzky, Y.-A. Chen, A. Flesch, I. P. McCulloch, U. Schollwöck, J. Eisert, and I. Bloch, *Nature Phys.* **8**, 325 (2012), [arXiv:1101.2659](#).
 - [35] A. Bakshi, A. Liu, A. Moitra, and E. Tang, in *2024 IEEE 65th Annual Symposium on Foundations of Computer Science (FOCS)* (IEEE, 2024) p. 10371050.
 - [36] J. Haah, R. Kothari, and E. Tang, *Nature Phys.* **20**, 10271031 (2024).
 - [37] R. Oliveira and B. M. Terhal, *arXiv* (2005), [arXiv:quant-ph/0504050](#).
 - [38] S. Arunachalam, A. B. Grilo, and A. Sundaram, *SIAM J. Comp.* **50**, 972 (2021).
 - [39] E. Onorati, C. Rouzé, D. S. França, and J. D. Watson, *arXiv* (2023), [arXiv:2311.07506](#).
 - [40] A. Coser and D. Prez-Garca, *Quantum* **3**, 174 (2019).
 - [41] H.-Y. Huang, S. Chen, and J. Preskill, *arXiv* (2023), [arXiv:2210.14894](#).
 - [42] H.-Y. Huang, R. Kueng, and J. Preskill, *Phys. Rev. Lett.* **126**, 190505 (2021).
 - [43] S. Chen, J. Cotler, H.-Y. Huang, and J. Li, *arXiv* (2021), [arXiv:2111.05881](#).
 - [44] H.-Y. Huang, M. Broughton, J. Cotler, S. Chen, J. Li, M. Mohseni, H. Neven, R. Babbush, R. Kueng, J. Preskill, and J. R. McClean, *Science* **376**, 11821186 (2022).
 - [45] D. Aharonov, J. Cotler, and X.-L. Qi, *Nature Comm.* **13**, 887 (2022).
 - [46] V. Vapnik and A. Vashist, *Neur. Net.* **22**, 544 (2009).
 - [47] V. Vapnik and R. Izmailov, *J. Mach. Learn. Res.* **16**, 2023 (2015).
 - [48] R. Molteni, C. Gyurik, and V. Dunjko, *arXiv* (2024), [arXiv:2405.02027](#).

- [49] N. F. Ramsey, *Phys. Rev.* **78**, 695 (1950).
- [50] C. M. Caves, *Phys. Rev. D* **23**, 1693 (1981).
- [51] M. J. Holland and K. Burnett, *Phys. Rev. Lett.* **71**, 1355 (1993).
- [52] H. Lee, P. Kok, and J. P. Dowling, *J. Mod. Opt.* **49**, 23252338 (2002).
- [53] V. Giovannetti, S. Lloyd, and L. Maccone, *Science* **306**, 1330 (2004).
- [54] M. de Burgh and S. D. Bartlett, *Phys. Rev. A* **72**, 042301 (2005).
- [55] C. L. Degen, F. Reinhard, and P. Cappellaro, *Rev. Mod. Phys.* **89**, 035002 (2017).
- [56] J. R. Garrison and T. Grover, *Phys. Rev. X* **8**, 021026 (2018).
- [57] E. Bairey, I. Arad, and N. H. Lindner, *Phys. Rev. Lett.* **122**, 020504 (2019).
- [58] X.-L. Qi and D. Ranard, *Quantum* **3**, 159 (2019).
- [59] Z. Li, L. Zou, and T. H. Hsieh, *Phys. Rev. Lett.* **124**, 160502 (2020).
- [60] A. Anshu, S. Arunachalam, T. Kuwahara, and M. Soleimanifar, *Nature Phys.* **17**, 931 (2021).
- [61] S. Arunachalam, A. Dutt, and F. E. Gutierrez, *arXiv* (2025), arXiv:2411.00082.
- [62] A. Bakshi, A. Liu, A. Moitra, and E. Tang, *arXiv* (2026), arXiv:2310.02243.
- [63] H. Fawzi, O. Fawzi, and S. O. Scalet, *Nature Comm.* **15**, 7394 (2024).
- [64] L. P. Garca-Pintos, K. Bharti, J. Bringewatt, H. Dehghani, A. Ehrenberg, N. Yunger Halpern, and A. V. Gorshkov, *Phys. Rev. Lett.* **133**, 040802 (2024).
- [65] S. Narayanan, *arXiv* (2024), arXiv:2407.04540.
- [66] S. Chen, J. Cotler, and H.-Y. Huang, *arXiv* (2025), arXiv:2510.08499.
- [67] D. Hangleiter, I. Roth, J. Fuxs, J. Eisert, and P. Roushan, *Nature Comm.* **15**, 9595 (2024).
- [68] F. Wilde, A. Kshetrimayum, I. Roth, D. Hangleiter, R. Sweke, and J. Eisert, *Mach. Learn. Sc. Tech.* **11**, 035002 (2026).
- [69] M. C. Caro, *ACM Trans. Quant. Comp.* **5**, 153 (2024).
- [70] M. P. da Silva, O. Landon-Cardinal, and D. Poulin, *Phys. Rev. Lett.* **107**, 210404 (2011).
- [71] B. Flynn, A. A. Gentile, N. Wiebe, R. Santagati, and A. Laing, *New J. Phys.* **24**, 053034 (2022).
- [72] D. S. França, T. Möbus, C. Rouzé, and A. H. Werner, *arXiv* (2025), arXiv:2510.08500.
- [73] A. A. Gentile, B. Flynn, S. Knauer, N. Wiebe, S. Paesani, C. E. Granade, J. G. Rarity, R. Santagati, and A. Laing, *Nature Phys.* **17**, 837 (2021).
- [74] H.-Y. Hu, M. Ma, W. Gong, Q. Ye, Y. Tong, S. T. Flammia, and S. F. Yelin, *PRX Quantum* **6**, 040315 (2025).
- [75] H.-Y. Huang, Y. Tong, D. Fang, and Y. Su, *Phys. Rev. Lett.* **130**, 200403 (2023).
- [76] H. Li, Y. Tong, H. Ni, T. Gefen, and L. Ying, *arXiv* (2023), arXiv:2307.04690.
- [77] T. Kraft, M. K. Joshi, W. Lam, T. Olsacher, F. Kranzl, J. Franke, L. K. Joshi, R. Blatt, A. Smerzi, D. S. Frana, B. Vermersch, B. Kraus, C. F. Roos, and P. Zoller, *arXiv* (2025), arXiv:2511.23392.
- [78] M. Ma, S. T. Flammia, J. Preskill, and Y. Tong, *arXiv* (2024), arXiv:2410.18928.
- [79] A. Mirani and P. Hayden, *Phys. Rev. A* **110**, 062421 (2024).
- [80] H. Ni, H. Li, and L. Ying, *arXiv* (2024), arXiv:2312.17390.
- [81] T. Odake, H. Kristjansson, A. Soeda, and M. Murao, *Phys. Rev. Res.* **6**, 1012063 (2024).
- [82] A. Shabani, M. Mohseni, S. Lloyd, R. L. Kosut, and H. Rabitz, *Phys. Rev. A* **84**, 012107 (2011).
- [83] S. D. Sinha and Y. Tong, *arXiv* (2025), arXiv:2509.07937.
- [84] N. Wiebe, C. Granade, C. Ferrie, and D. Cory, *Phys. Rev. A* **89**, 042314 (2014).
- [85] N. Wiebe, C. Granade, C. Ferrie, and D. G. Cory, *Phys. Rev. Lett.* **112**, 190501 (2014).
- [86] A. Zhao, in *Proceedings of the 57th Annual ACM Symposium on Theory of Computing*, STOC 25 (ACM, 2025) p. 12011211.
- [87] A. Zubida, E. Yitzhaki, N. H. Lindner, and E. Bairey, *arXiv* (2021), arXiv:2108.08824.
- [88] M. Shin, J. Lee, and C. Oh, *arXiv* (2026), arXiv:2604.27838.
- [89] A. Dutkiewicz, T. E. O'Brien, and T. Schuster, *Quantum* **8**, 1537 (2024).
- [90] A. Abbas, N. Cerrato, F. E. Gutierrez, D. Grinko, F. A. Mele, and P. Sinha, *arXiv* (2025), arXiv:2509.09813.
- [91] S. Arora and B. Barak, *Computational complexity: a modern approach* (Cambridge University Press, 2009).
- [92] A. Gu, L. Cincio, and P. J. Coles, *Nature Comm.* **15**, 312 (2024).
- [93] R. E. Schapire, *Mach. Learn.* **5**, 197 (1990).
- [94] N. H. Bshouty and J. C. Jackson, in *Proceedings of the eighth annual conference on Computational learning theory* (1995) pp. 118–127.
- [95] V. Lyubashevsky, C. Peikert, and O. Regev, in *Annual international conference on the theory and applications of cryptographic techniques* (Springer, 2010) pp. 1–23.
- [96] C. C. V. de Pradenne, J. Cotler, and H.-Y. Huang, *arXiv* (2026), arXiv:2606.05690.
- [97] A. Bakshi, A. Liu, A. Moitra, and E. Tang, *arXiv* (2023), arXiv:2310.02243.
- [98] C. Rouzé, D. Stilck França, and A. M. Alhambra, *Nature Phys.* **10.1038/s41567-026-03246-y** (2026).
- [99] S. Bravyi, A. Chowdhury, D. Gosset, and P. Wocjan, *Nature Phys.* **18**, 13671370 (2022).
- [100] S. Bravyi, A. Chowdhury, D. Gosset, V. Havlicek, and G. Zhu, *arXiv* (2024), arXiv:2302.11454 [quant-ph].
- [101] Y. Liu, J. R. Seddon, T. Kohler, E. Onorati, and T. S. Cubitt, *arXiv* (2025), arXiv:2507.07912 [quant-ph].
- [102] E. Onorati, T. Kohler, and T. S. Cubitt, *Quantum* **7**, 1197 (2023).
- [103] E. Bairey, C. Guo, D. Poletti, N. H. Lindner, and I. Arad, *New J. Phys.* **22**, 032001 (2020).
- [104] D. S. França, L. A. Markovich, V. V. Dobrovitski, A. H. Werner, and J. Borregaard, *arXiv* (2022), arXiv:2205.09567.
- [105] P. Ivashkov, N. Romanov, W. Gong, A. Gu, H.-Y. Hu, and S. F. Yelin, *arXiv* (2026), arXiv:2603.05492.

- [106] Y. Cai, [arXiv \(2026\)](#), [arXiv:2603.17736](#).
- [107] M. Barhoush, R. Nishimaki, and T. Yamakawa, [arXiv \(2025\)](#), [2505.14461](#).
- [108] H.-Y. Huang, R. Kueng, and J. Preskill, [Nature Phys.](#) **16**, 1050 (2020).
- [109] D. Nagaj, [J. Math. Phys.](#) **51**, 062201 (2010).
- [110] J. Feldman, [Duhamel's Formula](#) (2007), [Accessed 19-03-2026].

Appendix A: Classical shadows protocol

In this appendix, we provide details on the formalism of classical shadows, since we use it in Algorithm 2. In essence, one can think of the classical shadows protocol as consisting of a data acquisition phase, followed by a median-of-means procedure that estimates certain properties of interest. Below, we provide the subroutine used to construct the classical shadow of an unknown quantum state, given numerous copies thereof.

Subroutine 1 Data acquisition phase of classical shadows protocol [108]

Input:

1. $\tilde{N} \in \mathbb{N}$ copies of an n -qubit quantum state ρ .
2. Sample-access to the uniform distribution over the ensemble of unitaries $\mathcal{U} := \{\mathbb{I}, H, HS^\dagger\}^{\otimes n}$.

Output: A classical shadow given by a collection of unbiased snapshots, which we denote as $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$.

- 1: **for** $\ell = 1$ to $\tilde{N}$ **do**
- 2: Sample $U^{(\ell)} := \bigotimes_{j=1}^n U_j^{(\ell)}$ i.i.d. $\text{Unif}(\mathcal{U})$
- 3: Prepare $U^{(\ell)} \rho U^{(\ell)\dagger}$
- 4: Measure in the computational basis to obtain the pair $(U^{(\ell)}, b^{(\ell)}) \in \mathcal{U} \times \{0, 1\}^n$ ▷ Randomized measurement.
- 5: Construct the unbiased classical snapshot: ▷ Inversion of randomized measurement channel.

$$\hat{\sigma}^{(\ell)} \leftarrow \bigotimes_{j=1}^n \left(3U_j^{(\ell)\dagger} |b_j^{(\ell)}\rangle \langle b_j^{(\ell)}| U_j^{(\ell)} - \mathbb{I} \right) \quad (\text{A1})$$

- 6: **end for** ▷ Obtain a collection of classical snapshots: $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$.
 - 7: **return** $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$
-

Having outlined the method used to construct the classical shadow, let us take a look at how it can be used to estimate expectation values on some observables of interest.

Subroutine 2 Median-of-means estimation of expectation values [108]

Input:

1. A classical shadow $(\hat{\sigma}^{(\ell)})_{\ell=1}^{\tilde{N}}$ of some n -qubit quantum state ρ .
2. A set of observables $\mathcal{M} := \{O_1, O_2, \dots, O_P\}$.
3. Number of segments $G \in \mathbb{N}$.

Output: Estimates $(\langle \hat{O}_j \rangle_\rho)_{j=1}^P$ of the expectation values of the observables in $\mathcal{M}$.

- 1: Set $m \leftarrow \lfloor \tilde{N}/G \rfloor$ ▷ Dividing classical snapshot collection into G segments of size m .
 - 2: **for** $j = 1$ to P **do** ▷ Loop over all observables.
 - 3: **for** $g = 1$ to G **do** ▷ Loop over elements of each equally sized segment of classical snapshot collection.
 - 4: $S_g^{(j)} \leftarrow 0$ ▷ Initialize sum variable to zero.
 - 5: **for** $r = 1$ to m **do**
 - 6: $\ell \leftarrow (g-1)m + r$
 - 7: $S_g^{(j)} \leftarrow S_g^{(j)} + \text{Tr} [O_j \hat{\sigma}^{(\ell)}]$
 - 8: **end for**
 - 9: $\mu_g^{(j)} \leftarrow S_g^{(j)} / m$ ▷ Sample mean for each segment.
 - 10: **end for**
 - 11: $\langle \hat{O}_j \rangle_\rho \leftarrow \text{median} (\mu_1^{(j)}, \mu_2^{(j)}, \dots, \mu_G^{(j)})$ ▷ Median-of-means estimation.
 - 12: **end for**
 - 13: **return** $(\langle \hat{O}_j \rangle_\rho)_{j \in [P]}$ ▷ Predicted expectation values.
-

For theoretical details of the aforementioned subroutines, we refer the reader to the original work of Huang et al. [108].

Appendix B: Quantum advantage

1. A low-intersection BQP-complete Hamiltonian with constant-precision parameters

In this section, we construct a Hamiltonian family that satisfies the low-intersection property required by our learning algorithm while allowing us to simulate arbitrary BQP computations over long-time evolution. This Hamiltonian family also has the interesting property that the parameters λ specifying its Pauli expansion only take values from a discrete set, up to some constant precision. This means that learning λ to constant error (in the $\|\cdot\|_{\ell_\infty}$ -norm) is sufficient to recover the Hamiltonian exactly, and therefore simulate it for arbitrarily long times without any error. We prove that our quantum algorithm solves the learning problem perfectly in the classically hard case, where, in the definition of the concept class in Def. 6, $U(t, \lambda)$ is given by the time evolution of H_{BQP} as discussed in Section VI.

Lemma 2. *Let $\mathcal{L}$ be a promise problem in (Promise) BQP. There exists a local, low-intersection Hamiltonian*

$$H_{\text{BQP}}(\lambda) = \sum_{a \in [M]} \lambda_a E_a, \quad (\text{B1})$$

where the operators E_a are local Pauli strings, such that:

- For evolution times linear in system size, the real-time evolution generated by H_{BQP} simulates a quantum circuit U_{BQP} that decides $\mathcal{L}$ with bounded error.
- There exists a finite set $\Lambda \subset \mathbb{R}$, independent of the system size, such that $\lambda_a \in \Lambda$ for every $a \in [M]$.

This Hamiltonian belongs to a Hamiltonian family $\{H_{\text{BQP}}(\lambda)\}_\lambda$ that is capable of simulating any circuit U_{BQP} of the same depth, for a suitable choice of the parameters λ .

Proof. Let $\mathcal{L}$ be a (Promise) BQP problem and U_{BQP} the corresponding circuit that correctly decides it for input bit-strings in $\{0, 1\}^m$. By the Feynman-Kitaev clock Hamiltonian construction [26, 109], there exists a local Hamiltonian H_{BQP} such that evolving under it for a linear time T_{BQP} implements U_{BQP} with high probability. We will now show that the coefficients of the Pauli strings in H_{BQP} can be enforced to only take values in a fixed discrete set so that constant-precision estimates are sufficient to determine λ_{BQP} exactly and thus reconstruct H_{BQP} .

In particular, given the circuit $U_{\text{BQP}} = U_K U_{K-1} \dots U_2 U_1$ pre-compiled in a brickwork architecture using the universal gateset $\{\text{CNOT}, H, T\}$, the associated Feynman Hamiltonian is defined as

$$H_F = \sum_{j=1}^K H_j \quad (\text{B2})$$

where the terms

$$H_j = U_j \otimes |j\rangle\langle j-1| + U_j^\dagger \otimes |j-1\rangle\langle j| \quad (\text{B3})$$

act between a work register and a clock register. Later, in Subsection B 2, we show that, on average, the target concept implements U_{BQP} with high probability for a randomly sampled evolution time $t \in [0, T_{\text{BQP}}]$, where T_{BQP} is linear in the depth K of the circuit $U_{\text{BQP}} = U_K U_{K-1} \dots U_2 U_1$. Specifically, in Lemma 4, we show that this success probability of correctly implementing U_{BQP} can be brought to $1 - \delta$, that is, arbitrarily close to 1 by padding the circuit with $\mathcal{O}(K/\delta)$ identity gates and taking $T_{\text{BQP}} = \tilde{\mathcal{O}}(K/\delta^2)$.

Since the Feynman Hamiltonian is constructed with two-local gates $U_j \in \{\text{CNOT}, H, T\}$, the standard unary encoding of the clock register makes H_F 5-local. However, in order to have quantum learnability, we must ensure that H_{BQP} satisfies the low-intersection property as well. For this, we leverage the construction in Ref. [37]. This involves transforming the circuit U_{BQP} into an equivalent circuit $U'_{\text{BQP}} = U'_{K'} \dots U'_2 U'_1$ acting on $K' = mK$ qubits, such that each qubit participates in at most three gates (one U_j and two SWAPs). We refer to these additional qubits as the auxiliary work register. The Feynman Hamiltonian corresponding to this modified circuit is our final H_{BQP} . It takes the same form as in Eq. (B2), but uses unitaries $U'_j \in \{\text{CNOT}, H, T, \text{SWAP}\}$. Given that we choose a unary clock $|j\rangle = |1^j 0^{K'-j}\rangle$, the H_j terms can be rewritten as

$$H_j = U'_j \otimes I \otimes |1, 1, 0\rangle\langle 1, 0, 0|_{j-1, j, j+1} \otimes I + (U'_j)^\dagger \otimes I \otimes |1, 0, 0\rangle\langle 1, 1, 0|_{j-1, j, j+1} \otimes I \quad (\text{B4})$$

These terms can be rewritten in the form of Eq. (B1) by expanding each U'_j and clock transition operators in the Pauli basis. In particular, using the relations

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

and

$$|1\rangle\langle 1| = \frac{I - Z}{2}, \quad |0\rangle\langle 0| = \frac{I + Z}{2},$$

$$|1\rangle\langle 0| = \frac{X - iY}{2}, \quad |0\rangle\langle 1| = \frac{X + iY}{2},$$

we obtain

$$|1, 1, 0\rangle\langle 1, 0, 0| = \left(\frac{I - Z}{2}\right) \otimes \left(\frac{X - iY}{2}\right) \otimes \left(\frac{I + Z}{2}\right),$$

$$|1, 0, 0\rangle\langle 1, 1, 0| = \left(\frac{I - Z}{2}\right) \otimes \left(\frac{X + iY}{2}\right) \otimes \left(\frac{I + Z}{2}\right),$$

$$\text{SWAP} = \frac{1}{2}(I \otimes I + X \otimes X + Y \otimes Y + Z \otimes Z),$$

$$\text{CNOT} = |0\rangle\langle 0| \otimes I + |1\rangle\langle 1| \otimes X = \frac{I + Z}{2} \otimes I + \frac{I - Z}{2} \otimes X,$$

$$H = \frac{1}{\sqrt{2}}(X + Z),$$

$$T = \begin{pmatrix} 1 & 0 \\ 0 & e^{i\pi/4} \end{pmatrix} = \frac{1 + e^{i\pi/4}}{2} I + \frac{1 - e^{i\pi/4}}{2} Z.$$

Therefore, each gate U'_j , as well as each clock transition operator, admits a Pauli expansion with coefficients drawn from a fixed finite set Λ_0 . Since any given Pauli string can appear in at most a constant number of terms H'_j , each coefficient in the final expansion $H_{\text{BQP}} = \sum_{a \in [M]} \lambda_a E_a$ is a sum of at most a constant number of elements of Λ_0 . Consequently, the coefficients λ_a belong to a finite set Λ that is independent of the system size. Our final Hamiltonian family $\{H_{\text{BQP}}(\lambda)\}_\lambda$ is obtained by considering all possible Pauli decompositions obtainable by following the construction above for arbitrary choices of circuits $U_{\text{BQP}} = U_K U_{K-1} \dots U_2 U_1$. $\square$

2. Time-evolution under the Feynman-Kitaev clock Hamiltonian

Here, we provide a rigorous analysis of the dynamics generated by the Feynman-Kitaev Hamiltonian. In particular, we show that, on average over some polynomially large interval $[0, T]$ and upon padding with sufficiently many identity gates, long-time-evolution with respect to the clock Hamiltonian is able to correctly simulate BQP-complete computations with success probability arbitrarily close to 1.

a. Set-up

Our basic set-up will be as follows. Let $L, K \in \mathbb{N}$, with $L > K$. Consider $U_{\text{BQP}} = U_L U_{L-1} \dots U_K U_{K-1} \dots U_2 U_1$ with $\{U_L, U_{L-1} \dots U_{K+1}\}$ all identities and $L - K = R$. Let U_{BQP} decide membership of $z \in \{0, 1\}^m$ in some language $\mathcal{L}$ belonging to the class BQP, and let

$$H_F = \sum_{j=1}^L \left(U_j \otimes |j\rangle\langle j-1| + U_j^\dagger \otimes |j-1\rangle\langle j| \right) \quad (\text{B5})$$

be the corresponding Feynman-Kitaev clock Hamiltonian. For each $j \in \{1, \dots, L\}$, define:

$$|\psi_j(z)\rangle := U_j U_{j-1} \dots U_2 U_1 (|z\rangle \otimes |0\rangle^{\otimes a}) \in \mathcal{H}_{\text{work}}, \quad (\text{B6})$$

$$\text{and } |\eta_j(z)\rangle := |\psi_j(z)\rangle \otimes |j\rangle \in \mathcal{H}_{\text{work}} \otimes \mathcal{H}_{\text{clock}}, \quad (\text{B7})$$

with $|\psi_0(z)\rangle := |z\rangle \otimes |0\rangle^{\otimes a}$ for $j = 0$ and $|\eta_0(z)\rangle := (|z\rangle \otimes |0\rangle^{\otimes a}) \otimes |0\rangle$, to be the work and history states respectively. We will also consider the following “path” Hamiltonian, which acts on the clock subspace $\mathcal{H}_{\text{clock}}$:

$$H_{\text{path}} := \sum_{j=1}^L (|j\rangle\langle j-1| + |j-1\rangle\langle j|). \quad (\text{B8})$$

b. Time-evolution of history state

Proposition 1 (Unitary relation between H_F and H_{path}). *Let H_F be the Feynman-Kitaev clock Hamiltonian as defined in Eq. (B5) and H_{path} be the path Hamiltonian defined in Eq. (B8). Then, there exists a block-diagonal unitary*

$$W := \sum_{j=0}^L V_j \otimes |j\rangle\langle j|, \quad (\text{B9})$$

such that $H_F = W (I_{\text{work}} \otimes H_{\text{path}}) W^\dagger$.

Proof. Consider the following:

$$W (I_{\text{work}} \otimes H_{\text{path}}) W^\dagger \quad (\text{B10})$$

$$= W \left(I_{\text{work}} \otimes \sum_{j=1}^L |j\rangle\langle j-1| + |j-1\rangle\langle j| \right) W^\dagger \quad (\text{B11})$$

$$= \sum_{j=1}^L W (I_{\text{work}} \otimes (|j\rangle\langle j-1| + |j-1\rangle\langle j|)) W^\dagger \quad (\text{B12})$$

$$= \sum_{j=1}^L \left[\left(\sum_{k=0}^L V_k \otimes |k\rangle\langle k| \right) (I_{\text{work}} \otimes (|j\rangle\langle j-1| + |j-1\rangle\langle j|)) \left(\sum_{\ell=0}^L V_\ell \otimes |\ell\rangle\langle \ell| \right)^\dagger \right] \quad (\text{B13})$$

$$= \sum_{j=1}^L \left[\sum_{k=0}^L \sum_{\ell=0}^L V_k V_\ell^\dagger \otimes (|k\rangle\langle k|j\rangle\langle j-1|\ell\rangle\langle \ell| + |k\rangle\langle k|j-1\rangle\langle j|\ell\rangle\langle \ell|) \right] \quad (\text{B14})$$

$$= \sum_{j=1}^L \left[V_j V_{j-1}^\dagger \otimes |j\rangle\langle j-1| + V_{j-1} V_j^\dagger \otimes |j-1\rangle\langle j| \right] \quad (\text{B15})$$

$$= \sum_{j=1}^L \left[U_j U_{j-1} \dots U_1 U_1^\dagger \dots U_{j-1}^\dagger \otimes |j\rangle\langle j-1| + U_{j-1} \dots U_1 U_1^\dagger \dots U_{j-1}^\dagger U_j^\dagger \otimes |j-1\rangle\langle j| \right] \quad (\text{B16})$$

$$= \sum_{j=1}^L \left(U_j \otimes |j\rangle\langle j-1| + U_j^\dagger \otimes |j-1\rangle\langle j| \right) = H_F. \quad (\text{B17})$$

□

Lemma 3 (Time-evolved history state). *Let $|\eta_0(z)\rangle := |\psi_0(z)\rangle \otimes |0\rangle$ be the initial history state, and let $|\Psi_z(t)\rangle := e^{-iH_F t} |\eta_0(z)\rangle$ be the history state after time-evolution under the Feynman-Kitaev clock Hamiltonian H_F for some time $t > 0$. The state $|\Psi_z(t)\rangle$ is then given as follows:*

$$|\Psi_z(t)\rangle = \sum_{j=0}^L \langle j| e^{-iH_{\text{path}} t} |0\rangle |\eta_j(z)\rangle. \quad (\text{B18})$$

Proof. We compute $|\Psi_z(t)\rangle$ as follows:

$$|\Psi_z(t)\rangle := e^{-iH_F t} |\eta_0(z)\rangle \quad (\text{B19})$$

$$= e^{-iH_F t} (|\psi_0(z)\rangle \otimes |0\rangle) \quad (\text{B20})$$

$$= e^{-iW(I_{\text{work}} \otimes H_{\text{path}}) W^\dagger t} (|\psi_0(z)\rangle \otimes |0\rangle) \quad (\text{B21})$$

$$= W(I_{\text{work}} \otimes e^{-iH_{\text{path}} t}) W^\dagger (|\psi_0(z)\rangle \otimes |0\rangle) \quad (\text{B22})$$

$$= W(I_{\text{work}} \otimes e^{-iH_{\text{path}} t}) \left(\sum_{j=0}^L V_j^\dagger \otimes |j\rangle\langle j| \right) (|\psi_0(z)\rangle \otimes |0\rangle) \quad (\text{B23})$$

$$= W(I_{\text{work}} \otimes e^{-iH_{\text{path}} t}) (V_0^\dagger |\psi_0(z)\rangle \otimes |0\rangle\langle 0|) \quad (\text{B24})$$

$$= W(I_{\text{work}} \otimes e^{-iH_{\text{path}} t}) (|\psi_0(z)\rangle \otimes |0\rangle) \quad (\text{B25})$$

$$= W(|\psi_0(z)\rangle \otimes e^{-iH_{\text{path}} t} |0\rangle) \quad (\text{B26})$$

$$= \left(\sum_{j=0}^L V_j \otimes |j\rangle\langle j| \right) (|\psi_0(z)\rangle \otimes e^{-iH_{\text{path}} t} |0\rangle) \quad (\text{B27})$$

$$= \sum_{j=0}^L V_j |\psi_0(z)\rangle \otimes \langle j| e^{-iH_{\text{path}} t} |0\rangle |j\rangle \quad (\text{B28})$$

$$= \sum_{j=0}^L \langle j| e^{-iH_{\text{path}} t} |0\rangle |\psi_j(z)\rangle |j\rangle = \sum_{j=0}^L \langle j| e^{-iH_{\text{path}} t} |0\rangle |\eta_j(z)\rangle. \quad (\text{B29})$$

This concludes the proof. □

Lemma 4 (Average probability mass in the completed subspace). *Define $\alpha_j(t) := \langle j| e^{-iH_{\text{path}} t} |0\rangle$ to be the j^{th} probability amplitude of the time-evolved history state $|\Psi_z(t)\rangle$. Let*

$$\bar{p}_T := \frac{1}{T} \int_0^T \sum_{j=K}^L |\alpha_j(t)|^2 dt = \frac{1}{T} \int_0^T \langle \Psi_z(t) | I_{\text{work}} \otimes \Pi_{\text{done}} | \Psi_z(t) \rangle dt, \text{ where } \Pi_{\text{done}} := \sum_{j=K}^L |j\rangle\langle j|, \quad (\text{B30})$$

denote the time-averaged probability mass in the completed subspace, namely, the average probability that the clock register lies in the completed portion of the clock, $\{K, K+1, \dots, L\}$. Then, for some $\delta > 0$, in order for $\bar{p}_T \geq 1 - \delta$ to hold, it is sufficient to take a number of padded gates R and a time-averaging window $[0, T]$ such that

$$R \in \mathcal{O}\left(\frac{K}{\delta}\right) \quad (\text{B31})$$

$$\text{and } T \in \mathcal{O}\left(\frac{K}{\delta^2} \log\left(\frac{K}{\delta}\right)\right) \quad (\text{B32})$$

respectively.

Proof. From Equations (A2) and (A3) of Ref. [109], we know that the eigenvalues and eigenvectors of H_{path} are given as follows:

$$|v_r\rangle := \sqrt{\frac{2}{L+2}} \sum_{\ell=0}^L \sin\left(\frac{\pi r(\ell+1)}{L+2}\right) |\ell\rangle, \quad (\text{B33})$$

$$E_r = 2 \cos\left(\frac{\pi r}{L+2}\right), \quad (\text{B34})$$

for $r = 1, \dots, L+1$. Let us now compute

$$\sum_{j=K}^L |\alpha_j(t)|^2 = \sum_{j=K}^L |\langle j | e^{-iH_{\text{path}}t} | 0 \rangle|^2 \quad (\text{B35})$$

$$= \sum_{j=K}^L \langle 0 | e^{iH_{\text{path}}t} | j \rangle \langle j | e^{-iH_{\text{path}}t} | 0 \rangle \quad (\text{B36})$$

$$= \sum_{j=K}^L \langle 0 | \left(\sum_{s=1}^{L+1} e^{iE_s t} | v_s \rangle \langle v_s | \right) | j \rangle \langle j | \left(\sum_{r=1}^{L+1} e^{-iE_r t} | v_r \rangle \langle v_r | \right) | 0 \rangle \quad (\text{B37})$$

$$= \sum_{j=K}^L \sum_{r=1}^{L+1} \sum_{s=1}^{L+1} e^{-i(E_r - E_s)t} \langle 0 | v_s \rangle \langle v_s | j \rangle \langle j | v_r \rangle \langle v_r | 0 \rangle \quad (\text{B38})$$

$$= \frac{4}{(L+2)^2} \sum_{r,s=1}^{L+1} \left\{ \sin\left(\frac{\pi r}{L+2}\right) \sin\left(\frac{\pi s}{L+2}\right) \sum_{j=K}^L \sin\left(\frac{\pi r(j+1)}{L+2}\right) \sin\left(\frac{\pi s(j+1)}{L+2}\right) e^{-i(E_r - E_s)t} \right\}. \quad (\text{B39})$$

Let us define

$$A_{r,s} := \frac{4}{(L+2)^2} \sin\left(\frac{\pi r}{L+2}\right) \sin\left(\frac{\pi s}{L+2}\right) \sum_{j=K}^L \sin\left(\frac{\pi r(j+1)}{L+2}\right) \sin\left(\frac{\pi s(j+1)}{L+2}\right) \quad (\text{B40})$$

for notational convenience. We may rewrite the time-averaged probability mass as the following:

$$\bar{p}_T := \frac{1}{T} \int_0^T \sum_{j=K}^L |\alpha_j(t)|^2 dt \quad (\text{B41})$$

$$= \frac{1}{T} \int_0^T \sum_{r,s=1}^{L+1} A_{r,s} e^{-i(E_r - E_s)t} dt \quad (\text{B42})$$

$$= \underbrace{\sum_{r=1}^{L+1} A_{r,r}}_{=:\bar{p}_{\text{diag}}} + \underbrace{\sum_{\substack{r,s=1, \\ r \neq s}}^{L+1} \left(A_{r,s} \cdot \frac{1}{T} \int_0^T e^{-i(E_r - E_s)t} dt \right)}_{=:\bar{p}_{\text{off}}}, \quad (\text{B43})$$

where we have split the time-average into diagonal and off-diagonal parts. Let us focus on the diagonal part first. We have the following:

$$\bar{p}_{\text{diag}} = \sum_{r=1}^{L+1} A_{r,r} \quad (\text{B44})$$

$$= \frac{4}{(L+2)^2} \sum_{r=1}^{L+1} \sin^2 \left(\frac{\pi r}{L+2} \right) \sum_{j=K}^L \sin^2 \left(\frac{\pi r(j+1)}{L+2} \right) \quad (\text{B45})$$

$$= \frac{4}{(L+2)^2} \sum_{j=K}^L \sum_{r=1}^{L+1} \sin^2 \left(\frac{\pi r}{L+2} \right) \sin^2 \left(\frac{\pi r(j+1)}{L+2} \right) \quad (\text{B46})$$

$$= \frac{4}{(L+2)^2} \sum_{j=K}^L \sum_{r=1}^{L+1} \frac{1}{2} \left\{ 1 - \cos \left(\frac{2\pi r}{L+2} \right) \right\} \cdot \frac{1}{2} \left\{ 1 - \cos \left(\frac{2\pi r(j+1)}{L+2} \right) \right\} \quad (\text{B47})$$

$$= \frac{1}{(L+2)^2} \sum_{j=K}^L \sum_{r=1}^{L+1} \left\{ 1 - \cos \left(\frac{2\pi r}{L+2} \right) - \cos \left(\frac{2\pi r(j+1)}{L+2} \right) + \cos \left(\frac{2\pi r}{L+2} \right) \cos \left(\frac{2\pi r(j+1)}{L+2} \right) \right\} \quad (\text{B48})$$

$$= \frac{1}{(L+2)^2} \sum_{j=K}^L \sum_{r=1}^{L+1} \left\{ 1 - \cos \left(\frac{2\pi r}{L+2} \right) - \cos \left(\frac{2\pi r(j+1)}{L+2} \right) + \frac{1}{2} \left(\cos \left(\frac{2\pi r(j+1)}{L+2} - \frac{2\pi r}{L+2} \right) \right. \right. \quad (\text{B49})$$

$$\left. + \cos \left(\frac{2\pi r(j+1)}{L+2} + \frac{2\pi r}{L+2} \right) \right) \right\} \quad (\text{B50})$$

$$= \frac{1}{(L+2)^2} \sum_{j=K}^L \left\{ \sum_{r=1}^{L+1} 1 - \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r}{L+2} \right) - \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r(j+1)}{L+2} \right) + \frac{1}{2} \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r j}{L+2} \right) \right. \quad (\text{B51})$$

$$\left. + \frac{1}{2} \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r(j+2)}{L+2} \right) \right\}. \quad (\text{B52})$$

At this point, we will make use of the elementary roots-of-unity identity, which is given as the following, for some $q \in \mathbb{N}$:

$$\sum_{r=0}^{L+1} \exp \left(i \frac{2\pi r q}{L+2} \right) = \begin{cases} L+2, & q \equiv 0 \pmod{L+2}, \\ 0, & q \not\equiv 0 \pmod{L+2}. \end{cases} \quad (\text{B53})$$

Taking the real part of the above:

$$\sum_{r=1}^{L+1} \cos \left(\frac{2\pi r q}{L+2} \right) = \sum_{r=0}^{L+1} \cos \left(\frac{2\pi r q}{L+2} \right) - 1 \quad (\text{B54})$$

$$= \Re \left[\sum_{r=0}^{L+1} \exp \left(i \frac{2\pi r q}{L+2} \right) \right] - 1 = \begin{cases} L+1, & q \equiv 0 \pmod{L+2}, \\ -1, & q \not\equiv 0 \pmod{L+2}. \end{cases} \quad (\text{B55})$$

As we can observe from Eqs. (B51) and (B52), in our case, $q \in \{1, j, j+1, j+2\}$. Also, in our case, $j \in \{K, K+1, \dots, L\}$. Note that, for $K \leq j < L$, we have that for all $q \in \{1, j, j+1, j+2\}$, $q \not\equiv 0 \pmod{L+2}$. Thus, picking up again from Eq. (B52):

$$\sum_{r=1}^{L+1} A_{r,r} = \frac{1}{(L+2)^2} \sum_{j=K}^L \left\{ \sum_{r=1}^{L+1} 1 - \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r}{L+2} \right) - \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r(j+1)}{L+2} \right) + \frac{1}{2} \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r j}{L+2} \right) \right. \quad (\text{B56})$$

$$\left. + \frac{1}{2} \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r(j+2)}{L+2} \right) \right\} \quad (\text{B57})$$

$$= \frac{1}{(L+2)^2} \left(\sum_{j=K}^{L-1} \left\{ (L+1) - (-1) - (-1) + \frac{1}{2}(-1) + \frac{1}{2}(-1) \right\} \right) + \frac{3}{2(L+2)} \quad (\text{B58})$$

$$= \frac{1}{(L+2)^2} \sum_{j=K}^{L-1} (L+2) + \frac{3}{2(L+2)} = \frac{R}{L+2} + \frac{3}{2(L+2)} = \frac{2R+3}{2(L+2)}. \quad (\text{B59})$$

In the above, the term $\frac{3}{2(L+2)}$ is the result of computing the final summand. The derivation proceeds as follows:

$$\frac{4}{(L+2)^2} \sum_{r=1}^{L+1} \sin^2 \left(\frac{\pi r}{L+2} \right) \sin^2 \left(\frac{\pi r (L+1)}{L+2} \right) \quad (\text{B60})$$

$$= \frac{4}{(L+2)^2} \sum_{r=1}^{L+1} \sin^4 \left(\frac{\pi r}{L+2} \right) \quad (\text{B61})$$

$$= \frac{4}{(L+2)^2} \sum_{r=1}^{L+1} \left(\frac{3}{8} - \frac{1}{2} \cos \left(\frac{2\pi r}{L+2} \right) + \frac{1}{8} \cos \left(\frac{4\pi r}{L+2} \right) \right) \quad (\text{B62})$$

$$= \frac{4}{(L+2)^2} \cdot \left\{ \sum_{r=1}^{L+1} \frac{3}{8} - \frac{1}{2} \sum_{r=1}^{L+1} \cos \left(\frac{2\pi r}{L+2} \right) + \frac{1}{8} \sum_{r=1}^{L+1} \cos \left(\frac{4\pi r}{L+2} \right) \right\} \quad (\text{B63})$$

$$= \frac{4}{(L+2)^2} \left\{ \frac{3}{8} (L+1) - \frac{1}{2} (-1) + \frac{1}{8} (-1) \right\} \quad (\text{B64})$$

$$= \frac{4}{(L+2)^2} \left\{ \frac{3}{8} (L+1) + \frac{3}{8} \right\} = \frac{3}{2(L+2)}, \quad (\text{B65})$$

where, in the first line, we have used the fact that:

$$\sin^2 \left(\frac{\pi r}{L+2} \right) = \sin^2 \left(\frac{\pi r ((L+2) - (L+1))}{L+2} \right) \quad (\text{B66})$$

$$= \sin^2 \left(\pi r - \frac{\pi r (L+1)}{L+2} \right) \quad (\text{B67})$$

$$= \left\{ \underbrace{\sin(\pi r)}_{=0} \cos \left(\frac{\pi r (L+1)}{L+2} \right) - \sin \left(\frac{\pi r (L+1)}{L+2} \right) \underbrace{\cos(\pi r)}_{=(-1)^r} \right\}^2 \quad (\text{B68})$$

$$= \sin^2 \left(\frac{\pi r (L+1)}{L+2} \right). \quad (\text{B69})$$

Let us now turn our attention to the off-diagonal part of Eq. (B43). Recall that this is given as the following:

$$\bar{p}_{\text{off}} = \sum_{\substack{r,s=1, \\ r \neq s}}^{L+1} A_{r,s} \cdot \underbrace{\frac{1}{T} \int_0^T e^{-i(E_r - E_s)t} dt}_{=: I_{r,s}(T)} = \sum_{\substack{r,s=1, \\ r \neq s}}^{L+1} A_{r,s} I_{r,s}(T). \quad (\text{B70})$$

Consider now that

$$I_{r,s}(T) = \frac{1}{T} \int_0^T e^{-i(E_r - E_s)t} dt = \frac{1}{T} \int_0^T e^{i(E_s - E_r)t} dt = I_{s,r}^*(T). \quad (\text{B71})$$

The above will allow us to rewrite $\bar{p}_{\text{off}}$ in the following manner, grouping summands corresponding to the indices (r, s) and (s, r) together:

$$\bar{p}_{\text{off}} = \sum_{\substack{r,s=1, \\ r \neq s}}^{L+1} A_{r,s} I_{r,s}(T) \quad (\text{B72})$$

$$= \sum_{1 \leq r < s \leq L+1} A_{r,s} I_{r,s}(T) + A_{s,r} I_{s,r}(T) \quad (\text{B73})$$

$$= \sum_{1 \leq r < s \leq L+1} A_{r,s} (I_{r,s}(T) + I_{r,s}^*(T)) \quad (\text{B74})$$

$$= \sum_{1 \leq r < s \leq L+1} A_{r,s} \cdot 2\Re [I_{r,s}(T)] \quad (\text{B75})$$

$$= \sum_{1 \leq r < s \leq L+1} A_{r,s} \cdot 2\Re \left[\frac{1}{T} \int_0^T \left(e^{-i(E_r - E_s)t} \right) dt \right] \quad (\text{B76})$$

$$= \sum_{1 \leq r < s \leq L+1} 2A_{r,s} \cdot \frac{1}{T} \int_0^T \Re \left[e^{-i(E_r - E_s)t} \right] dt \quad (\text{B77})$$

$$= \sum_{1 \leq r < s \leq L+1} 2A_{r,s} \cdot \frac{1}{T} \int_0^T \cos((E_r - E_s)t) dt \quad (\text{B78})$$

$$= \sum_{1 \leq r < s \leq L+1} 2A_{r,s} \frac{\sin((E_r - E_s)T)}{(E_r - E_s)T}. \quad (\text{B79})$$

We therefore have

$$\bar{p}_T = \frac{2R+3}{2(L+2)} + 2 \sum_{1 \leq r < s \leq L+1} A_{r,s} \frac{\sin((E_r - E_s)T)}{(E_r - E_s)T}. \quad (\text{B80})$$

We now proceed to compute an upper-bound on $|\bar{p}_{\text{off}}|$.

$$|\bar{p}_{\text{off}}| = \left| 2 \sum_{1 \leq r < s \leq L+1} A_{r,s} \frac{\sin((E_r - E_s)T)}{(E_r - E_s)T} \right| \leq \frac{2}{T} \sum_{1 \leq r < s \leq L+1} \frac{|A_{r,s}|}{|E_r - E_s|}. \quad (\text{B81})$$

Now, let us introduce

$$\theta_r := \frac{\pi r}{L+2}, \quad \theta_s := \frac{\pi s}{L+2}. \quad (\text{B82})$$

for notational simplicity. Note that we have the following:

$$|E_r - E_s| = |2(\cos \theta_r - \cos \theta_s)| \quad (\text{B83})$$

$$= 4 \left| \sin \left(\frac{\theta_r + \theta_s}{2} \right) \sin \left(\frac{\theta_s - \theta_r}{2} \right) \right|, \quad (\text{B84})$$

as well as

$$|A_{r,s}| = \left| \frac{4}{(L+2)^2} \sin \theta_r \sin \theta_s \sum_{j=K}^L \sin(\theta_r(j+1)) \sin(\theta_s(j+1)) \right| \quad (\text{B85})$$

$$\leq \frac{4}{(L+2)^2} |\sin \theta_r \sin \theta_s| \sum_{j=K}^L \underbrace{|\sin(\theta_r(j+1)) \sin(\theta_s(j+1))|}_{\leq 1} \quad (\text{B86})$$

$$\leq \frac{4(R+1)}{(L+2)^2} |\sin \theta_r \sin \theta_s|. \quad (\text{B87})$$

Therefore, we have:

$$\frac{|A_{r,s}|}{|E_r - E_s|} \leq \frac{\frac{4(R+1)}{(L+2)^2} |\sin \theta_r \sin \theta_s|}{4 \left| \sin \left(\frac{\theta_r + \theta_s}{2} \right) \sin \left(\frac{\theta_s - \theta_r}{2} \right) \right|} \quad (\text{B88})$$

$$= \frac{(R+1)}{(L+2)^2} \cdot \frac{|\sin((\frac{\theta_r+\theta_s}{2}) - (\frac{\theta_s-\theta_r}{2})) \sin((\frac{\theta_r+\theta_s}{2}) + (\frac{\theta_s-\theta_r}{2}))|}{|\sin(\frac{\theta_r+\theta_s}{2}) \sin(\frac{\theta_s-\theta_r}{2})|} \quad (\text{B89})$$

$$= \frac{(R+1)}{(L+2)^2} \cdot \frac{|\sin^2(\frac{\theta_r+\theta_s}{2}) - \sin^2(\frac{\theta_s-\theta_r}{2})|}{|\sin(\frac{\theta_r+\theta_s}{2}) \sin(\frac{\theta_s-\theta_r}{2})|} \quad (\text{B90})$$

$$\leq \frac{(R+1)}{(L+2)^2} \cdot \frac{|\sin^2(\frac{\theta_r+\theta_s}{2})|}{|\sin(\frac{\theta_r+\theta_s}{2}) \sin(\frac{\theta_s-\theta_r}{2})|} \quad (\text{B91})$$

$$= \frac{(R+1)}{(L+2)^2} \cdot \frac{|\sin(\frac{\theta_r+\theta_s}{2})|}{|\sin(\frac{\theta_s-\theta_r}{2})|} \leq \frac{(R+1)}{(L+2)^2} \cdot \frac{1}{|\sin(\frac{\theta_s-\theta_r}{2})|} = \frac{(R+1)}{(L+2)^2} \cdot \frac{1}{\left|\sin\left(\frac{\pi(s-r)}{2(L+2)}\right)\right|}. \quad (\text{B92})$$

Since $1 \leq s-r \leq L$, we have that:

$$0 < \frac{\pi}{2(L+2)} \leq \frac{\pi(s-r)}{2(L+2)} \leq \frac{L}{(L+2)} \cdot \frac{\pi}{2} < \frac{\pi}{2}, \quad (\text{B93})$$

allowing us to use Jordan's inequality, as follows:

$$\frac{|A_{r,s}|}{|E_r - E_s|} \leq \frac{(R+1)}{(L+2)^2} \cdot \frac{1}{\left|\sin\left(\frac{\pi(s-r)}{2(L+2)}\right)\right|} \leq \frac{R+1}{(L+2)(s-r)}. \quad (\text{B94})$$

Going back to $\bar{p}_{\text{off}}$, we have:

$$|\bar{p}_{\text{off}}| \leq \frac{2}{T} \sum_{1 \leq r < s \leq L+1} \frac{R+1}{(L+2)(s-r)} = \frac{2(R+1)}{T(L+2)} \sum_{1 \leq r < s \leq L+1} \underbrace{\frac{1}{s-r}}_{=:k} \quad (\text{B95})$$

$$= \frac{2(R+1)}{T(L+2)} \sum_{k=1}^L \frac{L+1-k}{k} = \frac{2(R+1)}{T(L+2)} \left(\sum_{k=1}^L \frac{L+1}{k} - \sum_{k=1}^L 1 \right) \quad (\text{B96})$$

$$= \frac{2(R+1)}{T(L+2)} ((L+1)H_L - L) \leq \frac{2(R+1)}{T(L+2)} ((L+1)(1+\log L) - L) = \frac{2(R+1)}{T(L+2)} (1 + (L+1)\log L), \quad (\text{B97})$$

where, in the fourth equality, H_L denotes the L^{th} harmonic number, and in the inequality that follows, we have used the elementary bound on harmonic numbers: $H_L \leq 1 + \log L$. Now, recall that $\bar{p}_T \geq \bar{p}_{\text{diag}} - |\bar{p}_{\text{off}}|$, from which we get the following:

$$\bar{p}_T = \bar{p}_{\text{diag}} + \bar{p}_{\text{off}} \quad (\text{B98})$$

$$\geq \bar{p}_{\text{diag}} - |\bar{p}_{\text{off}}| \quad (\text{B99})$$

$$\geq \frac{2R+3}{2(L+2)} - \frac{2(R+1)}{T(L+2)} (1 + (L+1)\log L) \quad (\text{B100})$$

$$= 1 - \frac{2(L+2) - 2R - 3}{2(L+2)} - \frac{2(R+1)}{T(L+2)} (1 + (L+1)\log L) \quad (\text{B101})$$

$$= 1 - \frac{2K+1}{2(L+2)} - \frac{2(R+1)}{T(L+2)} (1 + (L+1)\log L). \quad (\text{B102})$$

We would like $\bar{p}_T \geq 1 - \delta$, for some failure probability $\delta > 0$. Let us now require:

$$\frac{2K+1}{2(L+2)} \leq \frac{\delta}{2} \implies L+2 \geq \frac{2K+1}{\delta} \implies R \geq \left\lceil \frac{2K+1}{\delta} - K - 2 \right\rceil. \quad (\text{B103})$$

Finally,

$$\frac{2(R+1)}{T(L+2)}(1+(L+1)\log L) \leq \frac{\delta}{2} \implies T \geq \frac{4(R+1)}{\delta(L+2)}(1+(L+1)\log L). \quad (\text{B104})$$

That is, in order to have $\bar{p}_T \geq 1 - \delta$, it is sufficient for the number of padded identity gates and length of time-averaging to scale as:

$$R \in \mathcal{O}\left(\frac{K}{\delta}\right), \quad (\text{B105})$$

$$T \in \mathcal{O}\left(\frac{K}{\delta^2} \log\left(\frac{K}{\delta}\right)\right). \quad (\text{B106})$$

This completes the proof. $\square$

Throughout this subsection, we have written the “done” projector in the following form:

$$\Pi_{\text{done}} = \sum_{j=K}^L |j\rangle\langle j|. \quad (\text{B107})$$

However, note that we use a unary encoding to represent the clock states; in particular, we write $|j\rangle_{\text{clock}} = |1^j 0^{L-j}\rangle$, where $j = 0, 1, \dots, L$. This leads to the K^{th} qubit in the clock region to be equal to 1 precisely when $j \geq K$. Hence, with this convention in mind, the done projector may be written as the following:

$$\Pi_{\text{done}} = \sum_{j=K}^L |j\rangle\langle j| \equiv |1\rangle\langle 1|_K, \quad (\text{B108})$$

and $I - \Pi_{\text{done}} \equiv |0\rangle\langle 0|_K$. Intuitively, we may thus think of the completed-clock information as being accessible by the following single-qubit Pauli Z operator acting on the K^{th} clock qubit:

$$-Z_K^{\text{clock}} = -|0\rangle\langle 0|_K + |1\rangle\langle 1|_K, \quad (\text{B109})$$

such that the $+1$ -eigenstates of the above operator are given by clock states in the region corresponding to $j \geq K$, and the -1 -eigenstates by those lying in the region where $j < K$. Our choice of observable for which we will show classical hardness will thus be given by the following two-body operator:

$$Q := Z_1^{\text{work}} \otimes (-Z_K^{\text{clock}}) = Z_1^{\text{work}} \otimes (-|0\rangle\langle 0|_K) + Z_1^{\text{work}} \otimes |1\rangle\langle 1|_K \quad (\text{B110})$$

$$= Z_1^{\text{work}} \otimes \Pi_{\text{done}} - Z_1^{\text{done}} \otimes (I - \Pi_{\text{done}}), \quad (\text{B111})$$

where the operator $Z_1^{\text{work}} = Z \otimes I \otimes \dots \otimes I$ is acting on the first qubit of the fully simulated circuit (i.e., right before padding with identities). When convenient, we will drop the superscripts denoting which subspace the tensored operators belong to.

Remark. Note that the observable Q above is indeed in the set of operators $(Q_a)_{a \in [M]}$ used for Hamiltonian learning (see Subsection III C 1). Indeed, for this to be the case, we need Q to be of the form $Q_a = i[P_a, E_a] = 2iP_a E_a$ for P_a a single-qubit Pauli operator that anti-commutes with a Pauli E_a appearing in the decomposition of our Hamiltonian. For that, we make use of a Pauli from the transition term

$$|1, 1, 0\rangle\langle 1, 0, 0|_{K', K'+1, K'+2}^{\text{clock}} \otimes U_{K'+1}^{\text{work}}, \quad (\text{B112})$$

which corresponds to the first transition in the idling phase. $U_{K'+1}^{\text{work}}$ is assumed to act on the first qubit of the fully simulated circuit and to be freely defined in our Hamiltonian family (but effectively set to identity). This allows us to use $E_a = IXI_{mK, mK+1, mK+2}^{\text{clock}} \otimes Z_1^{\text{work}}$ and simply choose the associated single-qubit $P_a = IYI_{mD, mK+1, mK+2}^{\text{clock}} \otimes I^{\text{work}}$ such that $O = Q_a = 2iP_a E_a = 2IZI_{mK, mK+1, mK+2}^{\text{clock}} \otimes Z_1^{\text{work}}$. The factor -2 is irrelevant.

3. Classical hardness of learning

Theorem 3 (Classical hardness of learning, formal). *For any (Promise) BQP-complete language, there exists a Hamiltonian family H_{BQP} specifying a concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ as per Def. 6 and a family of input distributions following the specification of Def. 7 such that no randomized polynomial-time classical algorithm $\mathcal{A}_c$ satisfies the PAC-learning condition of Def. 8 for this concept class, even when restricted to a constant learning error $\varepsilon < 1/48$, unless $\text{BQP} \subseteq \text{P/poly}$.*

Proof. We consider Q as defined in Eq. (B110), and begin by computing the expectation value of this observable evaluated on the time-evolved history state $|\Psi_z(t)\rangle$.

$$q(z, t) := \text{Tr}[Q|\Psi_z(t)\rangle\langle\Psi_z(t)|] \quad (\text{B113})$$

$$= \text{Tr}[(Z_1 \otimes (-Z_K))|\Psi_z(t)\rangle\langle\Psi_z(t)|] \quad (\text{B114})$$

$$= \text{Tr}\left[(Z_1 \otimes (-Z_K)) \sum_{j=0}^L \sum_{k=0}^L \alpha_j(t) \alpha_k^*(t) |\eta_j(z)\rangle\langle\eta_k(z)|\right] \quad (\text{B115})$$

$$= \sum_{j=0}^L \sum_{k=0}^L \alpha_j(t) \alpha_k^*(t) \text{Tr}[(Z_1 \otimes (-Z_K))(|\psi_j(z)\rangle\langle\psi_k(z)| \otimes |j\rangle\langle k|)] \quad (\text{B116})$$

$$= \sum_{j=0}^L \sum_{k=0}^L \alpha_j(t) \alpha_k^*(t) \text{Tr}[Z_1|\psi_j(z)\rangle\langle\psi_k(z)|] \text{Tr}[(-Z_K)|j\rangle\langle k|] \quad (\text{B117})$$

$$= \sum_{j=0}^L \sum_{k=0}^L \alpha_j(t) \alpha_k^*(t) \langle\psi_k(z)|Z_1|\psi_j(z)\rangle \underbrace{\langle k|(-Z_K)|j\rangle}_{s_j \delta_{j,k}}, \quad (\text{B118})$$

where we set

$$s_j := \begin{cases} +1, & \text{if } j \geq K, \\ -1, & \text{if } j < K. \end{cases} \quad (\text{B119})$$

We resume our computation of $q(z, t)$:

$$q(z, t) = \sum_{j=0}^L s_j |\alpha_j(t)|^2 \langle\psi_j(z)|Z_1|\psi_j(z)\rangle \quad (\text{B120})$$

$$= \sum_{j=K}^L |\alpha_j(t)|^2 \langle\psi_j(z)|Z_1|\psi_j(z)\rangle - \sum_{j=0}^{K-1} |\alpha_j(t)|^2 \langle\psi_j(z)|Z_1|\psi_j(z)\rangle \quad (\text{B121})$$

$$= \underbrace{\langle\psi_K(z)|Z_1|\psi_K(z)\rangle}_{=:m(z)} \underbrace{\sum_{j=K}^L |\alpha_j(t)|^2}_{=:p_{\text{done}}(t)} - \underbrace{\sum_{j=0}^{K-1} |\alpha_j(t)|^2 \langle\psi_j(z)|Z_1|\psi_j(z)\rangle}_{=:r(z, t)} \quad (\text{B122})$$

$$= m(z) p_{\text{done}}(t) - r(z, t), \quad (\text{B123})$$

where we denote by $r(z, t)$ a “garbage” term, corresponding to stages of the computation where $j < K$ and therefore not necessarily encoding the answer of the BQP computation. We note that:

$$|r(z, t)| = \left| \sum_{j=0}^{K-1} |\alpha_j(t)|^2 \langle\psi_j(z)|Z_1|\psi_j(z)\rangle \right| \quad (\text{B124})$$

$$\leq \sum_{j=0}^{K-1} |\alpha_j(t)|^2 |\langle\psi_j(z)|Z_1|\psi_j(z)\rangle| \quad (\text{B125})$$

$$\leq \sum_{j=0}^{K-1} |\alpha_j(t)|^2 = 1 - p_{\text{done}}(t). \quad (\text{B126})$$

Now, consider:

$$\bar{q}(z) := \frac{1}{T} \int_0^T q(z, t) dt = \frac{1}{T} \int_0^T (m(z) p_{\text{done}}(t) - r(z, t)) dt \quad (\text{B127})$$

$$= m(z) \cdot \underbrace{\frac{1}{T} \int_0^T p_{\text{done}}(t) dt}_{\bar{p}_T} - \underbrace{\frac{1}{T} \int_0^T r(z, t) dt}_{=: \bar{r}(z)} = m(z) \bar{p}_T - \bar{r}(z). \quad (\text{B128})$$

Let us note that:

$$|\bar{r}(z)| = \left| \frac{1}{T} \int_0^T r(z, t) dt \right| \leq \frac{1}{T} \int_0^T |r(z, t)| dt \leq \frac{1}{T} \int_0^T (1 - p_{\text{done}}(t)) dt = 1 - \bar{p}_T. \quad (\text{B129})$$

Thus, for some $\delta' > 0$, if by Lemma 4 we have that $\bar{p}_T \geq 1 - \delta'$, we also have that $|\bar{r}(z)| \leq \delta'$. Now, consider the following:

$$|\bar{q}(z) - m(z)| = |m(z) \bar{p}_T - \bar{r}(z) - m(z)| \quad (\text{B130})$$

$$= |m(z) (\bar{p}_T - 1) - \bar{r}(z)| \quad (\text{B131})$$

$$\leq |m(z)| |1 - \bar{p}_T| + |\bar{r}(z)| \leq \delta' + \delta' = 2\delta'. \quad (\text{B132})$$

Now, for $\gamma > 0$, assume that the underlying circuit U_{BQP} decides a language $\mathcal{L}$ such that, for $z \in \{0, 1\}^m$, we have

$$z \in \mathcal{L} \implies m(z) \leq -\gamma \implies \bar{q}(z) \leq -(\gamma - 2\delta'), \quad (\text{B133})$$

$$\text{and } z \notin \mathcal{L} \implies m(z) \geq \gamma \implies \bar{q}(z) \geq \gamma - 2\delta'. \quad (\text{B134})$$

We now show that an efficient learner for the concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$ would imply an efficient learner for q under an *arbitrary* distribution on z . Fix any distribution μ over $\{0, 1\}^m$. By the definition of the family of input distributions $\{\mathcal{D}_i\}_i$ (Def. 7), we can consider a distribution $\mathcal{D}_\mu$ over $x = (x_1, z, x_{[m+2, p(n)]}) \in \{0, 1\}^{p(n)}$ and $t \in [0, T]$ such that:

- t is uniformly distributed over $[0, T]$;
- x_1 is uniform in $\{0, 1\}$;
- conditioned on $x_1 = 1$, the bits $z = x_2 \cdots x_{m+1}$ are distributed according to μ , and the remaining bits $x_{m+2}, \dots, x_{p(n)}$ are deterministically set to 0;
- conditioned on $x_1 = 0$, the remaining bits are distributed uniformly in $\{0, 1\}^{p(n)-1}$.

Let $c_{\lambda_{\text{BQP}}}(x, t)$ be the target concept and, in particular,

$$(c_{\lambda_{\text{BQP}}}(\mathbf{1}, z, \mathbf{0}), t)_{i^*} = q(z, t) \quad (\text{B135})$$

be the distinguished component whose output equals the expectation value of the time-evolved history state on the observable Q . If our PAC-learning condition is met, there exists a learner that succeeds for the target concept $c_{\lambda_{\text{BQP}}}$ under any such distribution $\mathcal{D}_\mu$. Assume, towards contradiction, that $\mathcal{C}_{U_{\text{enc}}, T}^{\text{HBQP}}$ is efficiently learnable for an accuracy parameter ε to be defined later, and an arbitrary confidence parameter $\delta > 0$. Then, there is a randomized polynomial-time learner which, given training data generated from $c_{\lambda_{\text{BQP}}}$ under $\mathcal{D}_\mu$, outputs, with probability at least $1 - \delta$, a hypothesis $h(x, t)$ satisfying

$$\mathbb{E}_{(x, t) \sim \mathcal{D}_\mu} [\|c_{\lambda_{\text{BQP}}}(x, t) - h(x, t)\|_\infty] \leq \varepsilon. \quad (\text{B136})$$

Now, consider that:

$$\mathbb{E}_{(x,t) \sim \mathcal{D}_\mu} [|(h(x,t))_{i^*} - (c_{\lambda_{\text{BQP}}}(x,t))_{i^*}|] \leq \mathbb{E}_{(\mathbf{x},t) \sim \mathcal{D}_\mu} [\|c_{\lambda_{\text{BQP}}}(x,t) - h(x,t)\|_\infty] \leq \varepsilon. \quad (\text{B137})$$

Restricting the above expectation to the event $x_1 = 1$, we have the following:

$$\mathbb{E}_{\substack{z \sim \mu, \\ t \sim \text{Unif}([0,T])}} [|(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} - q(z, t)|] \cdot \Pr[x_1 = 1] \leq \varepsilon \quad (\text{B138})$$

$$\implies \mathbb{E}_{\substack{z \sim \mu, \\ t \sim \text{Unif}([0,T])}} [|(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} - q(z, t)|] \leq 2\varepsilon. \quad (\text{B139})$$

Let us define the time-averaged hypothesis as

$$\bar{h}(z) := \mathbb{E}_{t \sim \text{Unif}([0,T])} [(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*}] = \frac{1}{T} \int_0^T (h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} dt. \quad (\text{B140})$$

Recall also that we had defined

$$\bar{q}(z) = \frac{1}{T} \int_0^T q(z, t) dt = \mathbb{E}_{t \sim \text{Unif}([0,T])} [q(z, t)]. \quad (\text{B141})$$

Now, consider the following:

$$\mathbb{E}_{z \sim \mu} [|\bar{h}(z) - \bar{q}(z)|] \leq \mathbb{E}_{z \sim \mu} \left[\left| \frac{1}{T} \int_0^T [(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} - q(z, t)] dt \right| \right] \quad (\text{B142})$$

$$\leq \mathbb{E}_{z \sim \mu} \left[\frac{1}{T} \int_0^T |(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} - q(z, t)| dt \right] \quad (\text{B143})$$

$$= \mathbb{E}_{\substack{z \sim \mu, \\ t \sim \text{Unif}([0,T])}} [|(h((\mathbf{1}, z, \mathbf{0}), t))_{i^*} - q(z, t)|] \leq 2\varepsilon. \quad (\text{B144})$$

Let us choose $\delta' = \gamma/4$, so that we have $\gamma - 2\delta' = \gamma/2$. Further, let $\varepsilon < \gamma/16$. This leads to:

$$\mathbb{E}_{z \sim \mu} [|\bar{h}(z) - \bar{q}(z)|] \leq 2\varepsilon \leq \frac{\gamma}{8}. \quad (\text{B145})$$

Now, for the particular choice of $\delta' = \gamma/4$, we recall that the time-averaged distinguished component $\bar{q}(z)$ has a margin of at least $\gamma/2$ around zero. That is,

$$z \in \mathcal{L} \implies \bar{q}(z) \leq -\gamma/2, \quad (\text{B146})$$

$$\text{and } z \notin \mathcal{L} \implies \bar{q}(z) \geq \gamma/2. \quad (\text{B147})$$

Therefore, if the sign of $\bar{h}(z)$ differs from the sign of $\bar{q}(z)$, then we have

$$|\bar{h}(z) - \bar{q}(z)| \leq \frac{\gamma}{2}. \quad (\text{B148})$$

From the above, we may conclude the following:

$$\Pr_{z \sim \mu} [\text{sign}(\bar{h}(z)) \neq \text{sign}(\bar{q}(z))] \leq \Pr_{z \sim \mu} [|\bar{h}(z) - \bar{q}(z)| \geq \frac{\gamma}{2}] \quad (\text{B149})$$

$$\leq \frac{\mathbb{E}_{z \sim \mu} [|\bar{h}(z) - \bar{q}(z)|]}{\gamma/2} \quad (\text{B150})$$

$$\leq \frac{\gamma/8}{\gamma/2} = \frac{1}{4} < \frac{1}{2}, \quad (\text{B151})$$

where the second inequality follows from an application of Markov's inequality. Since the sign of $\bar{q}(z)$ encodes membership of z in $\mathcal{L}$, thresholding $\bar{h}(z)$ at zero gives a Boolean classifier with error at most $1/4$ under the arbitrary distribution μ . The assumed learner thus yields a distribution-free weak learner for $\mathcal{L}$ with constant error $1/4$. Taking the usual BQP gap of $\gamma = 1/3$ and choosing $\delta' = \gamma/4 = 1/12$, we finally get $\varepsilon < 1/48$, yielding a classification error of at most $1/4$. At this point, we invoke the standard consequence of Schapire's theorem: if a Boolean concept class is weakly learnable under arbitrary distributions, then it is strongly learnable [93]. On the other hand, if it is strongly learnable in randomized polynomial time, then for each input length there exists a polynomial-size circuit computing the target function exactly, i.e., the concept class lies in P/poly . Applying this to the language $\mathcal{L}$, we conclude that

$$\mathcal{L} \in P/\text{poly}. \quad (\text{B152})$$

Since $\mathcal{L}$ has been chosen to be (Promise) BQP-complete, this implies

$$\text{BQP} \subseteq P/\text{poly}, \quad (\text{B153})$$

contradicting the assumption. Therefore, no randomized polynomial-time classical algorithm can satisfy the PAC-learning condition for the concept class $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$, unless $\text{BQP} \subseteq P/\text{poly}$. $\square$

Remark. To put this result in perspective, we note that the trivial learning performance in this task is obtained for a hypothesis $h(x, t) = (0)_{i \in [M]}$, $\forall x, t$, which achieves $\varepsilon = 1$. The proof above shows hardness for $\varepsilon < 1/48$. Pushing the constants to their limits, namely, by taking the BQP acceptance threshold to be $\gamma = 1 - 1/\text{poly}(n)$, the failure probability of Hamiltonian simulation to be $\delta' = 1 - 1/\text{poly}(n)$, and by considering a non-uniform distribution over x_1 (we only need to have probability weight $1/\text{poly}(n)$ over $x_1 = 0$ and over short times $t \in [0, t^*]$ in order to get quantum learnability) as to get $\Pr[x_1 = 1] = 1 - 1/\text{poly}(n)$, we can easily bring the hardness threshold to $\varepsilon = 1/2 - 1/\text{poly}(n)$. The only unavoidable factor of $1/2$ originates from the advantage requirement for boosting in Eq. (B151).

Appendix C: Quantum learnability

In this appendix, we supply technical details related to the proof of quantum learnability of the concept class defined in Def. 6. First, in Subsection C 1, we provide a few technical definitions relevant to the analysis of quantum learnability. Then, following the main ideas explicated in Section V, we show a tail bound on $\mathcal{O}(t^2)$ terms of $\mathcal{F}_{a,t}$, which may be found in Lemma 7 of Subsection C 2. In Lemma 8 of Subsection C 3, we then determine how errors on Hamiltonian coefficients propagate to expectation values of time-evolved states, so as to obtain the necessary training precision for ε -accurate inference accuracy. In Subsection C 4, we provide bounds on the ideal estimator $\tilde{\mathcal{F}}_{a,t^*}$ (Lemma 9), followed by a thorough analysis of the sample complexity of PAC-learning our main concept class (Lemmas 10, 11, and 12). We conclude by formally restating our quantum learnability result in Theorem 5, where we also provide the time complexity of our quantum learning algorithm.

1. Additional definitions

We begin this subsection with the definition of a commutator and the Hadamard formula.

Definition 10 (Commutator). *Given operators $A, B \in \mathbb{C}^{D \times D}$, the commutator of A and B is defined as $[A, B] := AB - BA$. The k -fold nested commutator of A and B is given by $[A, B]_k$, which is defined recursively as*

$$[A, B]_k := \begin{cases} B, & k = 0, \\ A[A, B]_{k-1} - [A, B]_{k-1}A, & k \geq 1. \end{cases} \quad (\text{C1})$$

That is, $[A, B]_k = \underbrace{[A, [A, \dots, [A, B] \dots]]}_{k \text{ times}}$.

Lemma 5 (Hadamard formula). *For $A, B \in \mathbb{C}^{D \times D}$, it holds that*

$$e^{-iA} B e^{iA} = \sum_{k=0}^{\infty} \frac{(-i)^k}{k!} [A, B]_k. \quad (\text{C2})$$

The usefulness of Def. 10 and Lemma 5 will become readily apparent, since our analysis of quantum learnability will involve manipulating and expanding time-evolved operators of the form shown on the L.H.S. of Eq. (C2). However, our calculations will involve nested commutators that are more general than $[A, B]_k$, such as $[M_1, [M_2, [\dots, [M_{\ell-1}, M_\ell] \dots]]]$, which are nested commutators constructed using more than two operators. In particular, we will be interested in scenarios when such quantities are non-zero. This leads us to the definition of a *cluster*.

Definition 11 (Cluster, Definition 2.17, [97]). *Let $M_1, M_2, \dots, M_\ell \in \mathbb{C}^{D \times D}$. We say that the ordered ℓ -tuple $(M_1, M_2, \dots, M_\ell)$ forms a cluster if, for all $a \in [\ell - 1]$, $\text{supp}(M_{a+1}) \cap (\text{supp}(M_1) \cup \text{supp}(M_2) \cup \dots \cup \text{supp}(M_a)) \neq \emptyset$.*

To see the motivation behind Def. 11, note that the quantity $[M_1, [M_2, [\dots, [M_{\ell-1}, M_\ell] \dots]]]$ is non-zero only when $\text{supp}(M_{\ell-1}) \cap \text{supp}(M_\ell) \neq \emptyset$, $\text{supp}(M_{\ell-2})$ overlaps with $\text{supp}(M_{\ell-1}) \cup \text{supp}(M_\ell)$, and so on. This is exactly the condition imposed by Def. 11. The final technical ingredient we will require is a cluster-counting lemma that has originally been shown in Ref. [97], which we restate below for completeness:

Lemma 6 (Counting clusters, Item (b) of Lemma 2.18, [97]). *Let $\mathcal{E} \subset \mathcal{P}$ be a set of Pauli operators, where $\mathcal{P}$ is the set of n -qubit Pauli strings of the form $P_1 \otimes \dots \otimes P_n$, where each $P_i \in \{I, X, Y, Z\}$. Let every $P \in \mathcal{E}$ satisfy $|\text{supp}(P)| \leq \mathfrak{K}$ and let the dual-interaction graph associated with $\mathcal{E}$ have maximum degree $\mathfrak{d}$. Let $H_1, H_2, \dots, H_\ell$ be linear combinations of elements in $\mathcal{E}$, such that $H_i := \sum_{P \in \mathcal{E}} \lambda_{i,P} P$, for all $i \in [\ell]$. Let also $A \in \mathcal{E}$. Then, we may write*

$$[H_1, [H_2, [\dots, [H_\ell, A] \dots]]] = 2^\ell \sum_{\substack{P_1, P_2, \dots, P_\ell \in \mathcal{E}, \\ (A, P_1, \dots, P_\ell) \text{ is a cluster.}}} c_{P_1, P_2, \dots, P_\ell} Q_{P_1, P_2, \dots, P_\ell} \prod_{j=1}^{\ell} \lambda_{j, P_j}, \quad (\text{C3})$$

where $c_{P_1, P_2, \dots, P_\ell} \in \{0, \pm 1, \pm i\}$ and $Q_{P_1, P_2, \dots, P_\ell}$ are certain Pauli operators. Also, the number of terms constituting the sum in Eq. (C3) is given by $\ell! (\mathfrak{d} + 1)^\ell$.

2. Tail bound

Lemma 7 (Tail bound for higher-order terms of $\mathcal{F}_{a,t}$). *Let $H = \sum_{a \in [M]} \lambda_a E_a$ be an n -qubit Hamiltonian as defined in Def. 4. Let $U := e^{-iHt}$, for $t > 0$. For each $a \in [M]$, let P_a be any single-qubit Pauli that anti-commutes with E_a and $Q_a = i[P_a, E_a] = 2iP_a E_a$. Then, define*

$$\mathcal{F}_{a,t} := \frac{1}{D} \text{Tr} [Q_a U P_a U^\dagger]. \quad (\text{C4})$$

Then, it holds true that

$$|\mathcal{F}_{a,t} - 4t\lambda_a| \leq \frac{8t^2 (\mathfrak{d} + 1)^2}{1 - 2t (\mathfrak{d} + 1)}, \quad (\text{C5})$$

where $\mathfrak{d}$ is the degree of $\mathfrak{G}$, the dual-interaction graph of H , provided that $2t (\mathfrak{d} + 1) < 1$.

Proof. Consider

$$\begin{aligned} U P_a U^\dagger &= e^{-iHt} P_a e^{iHt} \\ &= \sum_{k=0}^{\infty} \frac{(-it)^k}{k!} [H, P_a]_k \\ &= \sum_{k=0}^{\infty} \frac{(-it)^k}{k!} \sum_{\substack{a_1, a_2, \dots, a_k \in [M]: \\ \forall j \in [k], \text{supp}(E_{a_j}) \cap (\bigcup_{i=j+1}^k \text{supp}(E_{a_i}) \cup \text{supp}(A)) \neq \emptyset}} \underbrace{[E_{a_1}, [E_{a_2}, [\dots, [E_{a_k}, P_a] \dots]]]}_{=: C_{a_1, a_2, \dots, a_k}} \lambda_{a_1} \lambda_{a_2} \dots \lambda_{a_k} \end{aligned} \quad (\text{C6})$$

$$= \sum_{k=0}^{\infty} t^k q_k(\lambda_1, \lambda_2, \dots, \lambda_M),$$

where, in the final equality, according to the notation of Ref. [36], $q_k(\lambda_1, \lambda_2, \dots, \lambda_M)$ is a degree- k , homogeneous, matrix-valued polynomial in the Hamiltonian coefficients $(\lambda_a)_{a \in [M]}$. Also, note that, in the penultimate equality, the interior sum is being taken over clusters of the form $(E_{a_1}, E_{a_2}, \dots, E_{a_k}, P_a)$, which is to ensure that the summands are non-zero. This condition is exactly the same as requiring that the $(k+1)$ -tuple $(E_{a_1}, E_{a_2}, \dots, E_{a_k}, P_a)$ is a cluster as in Def. 11. Now, consider

$$\begin{aligned} \|q_k(\lambda_1, \lambda_2, \dots, \lambda_M)\|_{\infty} &= \left\| \frac{(-i)^k}{k!} \sum_{\substack{a_1, a_2, \dots, a_k \in [M]: \\ \forall j \in [k], \text{supp}(E_{a_j}) \cap (\bigcup_{i=j+1}^k \text{supp}(E_{a_i}) \cup \text{supp}(P_a)) \neq \emptyset}} C_{a_1, a_2, \dots, a_k} \lambda_{a_1} \lambda_{a_2} \dots \lambda_{a_k} \right\|_{\infty} \\ &\leq \frac{1}{k!} \sum_{\substack{a_1, a_2, \dots, a_k \in [M]: \\ \forall j \in [k], \text{supp}(E_{a_j}) \cap (\bigcup_{i=j+1}^k \text{supp}(E_{a_i}) \cup \text{supp}(P_a)) \neq \emptyset}} \underbrace{\|C_{a_1, a_2, \dots, a_k}\|_{\infty}}_{\leq 2^k} \underbrace{|\lambda_{a_1} \lambda_{a_2} \dots \lambda_{a_k}|}_{\leq 1} \\ &\leq \frac{2^k}{k!} (\mathfrak{d} + 1)^k = (2(\mathfrak{d} + 1))^k, \end{aligned}$$

where the bound on $C_{a_1, a_2, \dots, a_k}$ follows from iterative use of the trivial commutator bound

$$\begin{aligned} \|C_{a_1, a_2, \dots, a_k}\|_{\infty} &= \|[E_{a_1}, [E_{a_2}, [\dots, [E_{a_k}, P_a] \dots]]]\|_{\infty} \\ &\leq 2^k \underbrace{\left(\prod_{i=1}^k \|E_{a_i}\|_{\infty} \right)}_{\leq 1} \underbrace{\|P_a\|_{\infty}}_{\leq 1} \leq 2^k, \end{aligned} \tag{C7}$$

and the inequality in the third line follows from Lemma 6, which counts the number of monomials (or, equivalently, the number of clusters over which the sum is taken) contained in the summation form of $[H, P]_k$. Now, consider

$$\begin{aligned} \mathcal{F}_{a,t} &= \frac{1}{D} \text{Tr} [Q_a U P_a U^{\dagger}] \\ &= \frac{1}{D} \text{Tr} \left[Q_a \sum_{k=0}^{\infty} \frac{(-it)^k}{k!} [H, P_a]_k \right] \\ &= \frac{1}{D} \sum_{k=0}^{\infty} \frac{(-it)^k}{k!} \text{Tr} [Q_a [H, P_a]_k] \\ &= \frac{1}{D} \left(\underbrace{\text{Tr} [Q_a P_a]}_{=0} - it \text{Tr} [Q_a [H, P_a]] \right) + \underbrace{\frac{1}{D} \text{Tr} \left[Q_a \sum_{k=2}^{\infty} \frac{(-it)^k}{k!} [H, P_a]_k \right]}_{\mathcal{O}(t^2)} \\ &= \frac{1}{D} \left(-it \text{Tr} \left[Q_a \left[\sum_{b \in [M]} \lambda_b E_b, P_a \right] \right] \right) + \mathcal{O}(t^2) \\ &= \frac{-it}{D} \sum_{b \in [M]} \lambda_b \text{Tr} [Q_a [E_b, P_a]] + \mathcal{O}(t^2) \\ &= \frac{-it}{D} \sum_{b \in [M]} \lambda_b \text{Tr} [2i P_a E_a (E_b P_a - P_a E_b)] + \mathcal{O}(t^2) \end{aligned} \tag{C8}$$

$$\begin{aligned}
&= \frac{2t}{D} \sum_{b \in [M]} \lambda_b \operatorname{Tr} [P_a E_a (E_b P_a - P_a E_b)] + \mathcal{O}(t^2) \\
&= \frac{2t}{D} \sum_{b \in [M]} \lambda_b (\operatorname{Tr} [P_a E_a E_b P_a] - \operatorname{Tr} [P_a E_a P_a E_b]) + \mathcal{O}(t^2) \\
&= \frac{2t}{D} \sum_{b \in [M]} \lambda_b (\operatorname{Tr} [P_a^2 E_a E_b] + \operatorname{Tr} [E_a P_a^2 E_b]) + \mathcal{O}(t^2) \\
&= \frac{4t}{D} \sum_{b \in [M]} \lambda_b \operatorname{Tr} [E_a E_b] + \mathcal{O}(t^2) = \frac{4t}{D} \cdot \lambda_a \operatorname{Tr} [E_a^2] + \mathcal{O}(t^2) = 4t\lambda_a + \mathcal{O}(t^2),
\end{aligned}$$

where, in the fifth equality, we have simply plugged in the definition of H , in the seventh equality, we have used the definition of Q_a , and the tenth equality follows from the cyclicity of the trace and the fact that $P_a E_a = -E_a P_a$. From the above, we have

$$\begin{aligned}
|\mathcal{F}_a - 4t\lambda_a| &\leq \left| \frac{1}{D} \operatorname{Tr} \left[Q_a \sum_{k=2}^{\infty} \frac{(-it)^k}{k!} [H, P_a]_k \right] \right| \tag{C9} \\
&\leq \frac{1}{D} \sum_{k=2}^{\infty} t^k \left| \operatorname{Tr} \left[Q_a \underbrace{\frac{(-i)^k}{k!} [H, P_a]_k}_{q_k(\lambda_1, \lambda_2, \dots, \lambda_M)} \right] \right| \\
&= \frac{1}{D} \sum_{k=2}^{\infty} t^k |\operatorname{Tr} [Q_a q_k(\lambda_1, \lambda_2, \dots, \lambda_M)]| \\
&\leq \frac{1}{D} \sum_{k=2}^{\infty} t^k \underbrace{\|Q_a\|_{\infty}}_{=2} \|q_k(\lambda_1, \lambda_2, \dots, \lambda_M)\|_1 \\
&\leq \frac{2}{D} \sum_{k=2}^{\infty} t^k \cdot D \|q_k(\lambda_1, \lambda_2, \dots, \lambda_M)\|_{\infty} \\
&\leq 2 \sum_{k=2}^{\infty} (2t(\mathfrak{d}+1))^k \\
&= 2 \left(\frac{1}{1-2t(\mathfrak{d}+1)} - 1 - 2t(\mathfrak{d}+1) \right) \\
&= 2 \left(\frac{1 - (1-2t(\mathfrak{d}+1))(1+2t(\mathfrak{d}+1))}{1-2t(\mathfrak{d}+1)} \right) \\
&= 2 \left(\frac{1 - (1 - (2t(\mathfrak{d}+1))^2)}{1-2t(\mathfrak{d}+1)} \right) = \frac{8t^2(\mathfrak{d}+1)^2}{1-2t(\mathfrak{d}+1)}.
\end{aligned}$$

The second equality holds when $2t(\mathfrak{d}+1) < 1$. □

3. Hamiltonian learning precision required for the continuous-parameter case

Lemma 8 (Error propagation from Hamiltonian coefficients to time-evolved expectation values). *Let $H(\lambda) := \sum_{a \in [M]} \lambda_a E_a$ and $H(\lambda') := \sum_{a \in [M]} \lambda'_a E_a$ be Hamiltonians defined on n qubits, where $\lambda := (\lambda_a)_{a \in [M]}$ and $\lambda' := (\lambda'_a)_{a \in [M]}$. Also, let $\|\lambda - \lambda'\|_{\ell_{\infty}} := \max_{a \in [M]} |\lambda_a - \lambda'_a| \leq \varepsilon$, for some $\varepsilon > 0$. Define $\rho(\lambda) := e^{-iH(\lambda)t} \rho_0 e^{iH(\lambda)t}$ and $\rho(\lambda') := e^{-iH(\lambda')t} \rho_0 e^{iH(\lambda')t}$, where $\rho_0 \in \mathbb{C}^{2^n \times 2^n}$ is an arbitrary initial state. Then:*

$$|\langle O \rangle_{\rho(\lambda)} - \langle O \rangle_{\rho(\lambda')}| \leq 2 \|O\|_{\infty} M |t| \varepsilon, \tag{C10}$$

where $O \in \text{Herm}(2^n)$ is an observable.

Proof. For notational convenience, we will denote $H(\lambda)$ and $H(\lambda')$ by H and H' respectively, and $\rho(\lambda)$ and $\rho(\lambda')$ by ρ and ρ' respectively. Recall now that, for some parametrized operator $A(t)$, by Duhamel's formula for derivatives of matrix exponentials [110], we have

$$\partial_t e^{A(t)} = \int_0^1 e^{sA(t)} \partial_t A(t) e^{(1-s)A(t)} ds. \quad (\text{C11})$$

Now, for $s \in [0, 1]$, define

$$H(s) := H' + s(H - H'), \quad (\text{C12})$$

so that $H(0) = H'$ and $H(1) = H$. Now, consider

$$\begin{aligned} e^{-iHt} - e^{-iH't} &= \int_0^1 \partial_s e^{-iH(s)t} ds \\ &= -it \int_0^1 \int_0^1 e^{-is'H(s)t} \partial_s H(s) e^{(1-s')(-iH(s)t)} ds ds' \\ &= -it \int_0^1 \int_0^1 e^{-is'H(s)t} (H - H') e^{(1-s')(-iH(s)t)} ds ds', \end{aligned} \quad (\text{C13})$$

where the first equality follows from the fundamental theorem of calculus, the second equality follows from an application of Duhamel's formula, and, in the third equality, we have made use of the fact that $\partial_s H(s) = H - H'$. Taking the norm of both sides, we get:

$$\|e^{-iHt} - e^{-iH't}\|_\infty \leq |t| \int_0^1 \int_0^1 \underbrace{\|e^{-is'H(s)t}\|_\infty}_{=1} \|H - H'\|_\infty \underbrace{\|e^{(1-s')(-iH(s)t)}\|_\infty}_{=1} ds ds' \quad (\text{C14})$$

$$\leq |t| \|H - H'\|_\infty. \quad (\text{C15})$$

Let $U := e^{-iHt}$ and $U' := e^{-iH't}$. We bound the trace norm between ρ and ρ' as

$$\begin{aligned} \|\rho - \rho'\|_1 &= \|U\rho_0 U^\dagger - U'\rho_0 U'^\dagger\|_1 \\ &= \|U\rho_0 U^\dagger + U'\rho_0 U^\dagger - U'\rho_0 U^\dagger - U'\rho_0 U'^\dagger\|_1 \\ &\leq \|(U - U')\rho_0 U^\dagger\|_1 + \|U'\rho_0(U^\dagger - U'^\dagger)\|_1 \\ &\leq \|U - U'\|_\infty + \|U^\dagger - U'^\dagger\|_\infty \\ &\leq 2|t| \|H - H'\|_\infty \leq 2t\varepsilon \sum_{a \in [M]} \|E_a\|_\infty = 2tM\varepsilon. \end{aligned} \quad (\text{C16})$$

Finally, for some observable O , we find

$$|\langle O \rangle_\rho - \langle O \rangle_{\rho'}| \leq \|O\|_\infty \|\rho - \rho'\|_1 \leq 2\|O\|_\infty tM\varepsilon, \quad (\text{C17})$$

which is exactly the advertised bound. $\square$

4. Sample and time complexity analysis for quantum learning algorithm

Lemma 9 (Bounds on time-averaged integral). *Let $\mathcal{F}_{a,t}$ be as defined in Eq. (C4), Lemma 7. For $t^* > 0$, define the time-averaged integral*

$$\tilde{\mathcal{F}}_{a,t^*} := \frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{a,t} dt. \quad (\text{C18})$$

Then, for $\varepsilon > 0$, we find

$$\left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*}, \quad (\text{C19})$$

where $\mathfrak{d}$ is the maximum degree of the dual-interaction graph $\mathfrak{G}$ of H .

Proof. Consider that

$$\begin{aligned} \tilde{\mathcal{F}}_{a,t^*} &:= \frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{a,t} dt \\ &\leq \frac{1}{t^*} \int_0^{t^*} 4\lambda_a t + \frac{8t^2 (\mathfrak{d} + 1)^2}{1 - 2t(\mathfrak{d} + 1)} dt \\ &= 2\lambda_a t^* + \frac{2}{t^*} \int_0^{t^*} \frac{1 - (1 - 4t^2 (\mathfrak{d} + 1)^2)}{1 - 2t(\mathfrak{d} + 1)} dt \\ &= 2\lambda_a t^* + \frac{2}{t^*} \int_0^{t^*} \frac{1 - (1 - 2t(\mathfrak{d} + 1))(1 + 2t(\mathfrak{d} + 1))}{1 - 2t(\mathfrak{d} + 1)} dt \\ &= 2\lambda_a t^* + \frac{2}{t^*} \int_0^{t^*} \frac{1}{1 - 2t(\mathfrak{d} + 1)} dt - \frac{2}{t^*} \int_0^{t^*} 1 + 2t(\mathfrak{d} + 1) dt \\ &= 2\lambda_a t^* - \frac{2}{t^*} \cdot \frac{1}{2(\mathfrak{d} + 1)} \int_1^{1-2t^*(\mathfrak{d}+1)} \frac{1}{u} du - \frac{2}{t^*} \left(t^* + 2(\mathfrak{d} + 1) \cdot \frac{t^{*2}}{2} \right) \\ &= 2\lambda_a t^* - \frac{\ln |1 - 2t^*(\mathfrak{d} + 1)|}{t^* (\mathfrak{d} + 1)} - 2 - 2(\mathfrak{d} + 1)t^*. \end{aligned} \quad (\text{C20})$$

The above implies

$$\left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \frac{1}{2t^*} \left(-\frac{\ln |1 - 2t^*(\mathfrak{d} + 1)|}{t^* (\mathfrak{d} + 1)} - 2 - 2(\mathfrak{d} + 1)t^* \right). \quad (\text{C21})$$

For ease of symbolic manipulation, let us define $a \leftarrow (\mathfrak{d} + 1)$ and $x \leftarrow 2at^*$. Now, consider that

$$\ln(1 - x) = -\sum_{n=1}^{\infty} \frac{x^n}{n} \implies -\frac{\ln(1 - x)}{x} - 1 - \frac{x}{2} = \sum_{n=2}^{\infty} \frac{x^n}{n+1} \leq \frac{1}{2} \sum_{n=2}^{\infty} x^n = \frac{x^2}{2(1 - x)}. \quad (\text{C22})$$

From the above, we have

$$\frac{2a}{x} \left(-\frac{\ln(1 - x)}{x} - 1 - \frac{x}{2} \right) \leq \frac{ax}{(1 - x)} = \frac{(\mathfrak{d} + 1) \cdot 2at^*}{1 - 2at^*} = \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*}. \quad (\text{C23})$$

That is,

$$\left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*}. \quad (\text{C24})$$

This completes the proof. $\square$

Theorem 4 (Hoeffding's inequality). *Let $Y_1, Y_2, \dots, Y_N$ be N independent samples of some random variable Y that is bounded in the interval $[a, b]$ and with mean $\mu := \mathbb{E}[Y]$. Let $\bar{Y}_N$ be the sample mean*

$$\bar{Y}_N := \frac{1}{N} \sum_{i=1}^N Y_i. \quad (\text{C25})$$

Then, for all precision and confidence parameters $\varepsilon, \delta > 0$,

$$\Pr \left[|\bar{Y}_N - \mu| \leq \varepsilon \right] \geq 1 - \delta \quad (\text{C26})$$

holds true, provided that the number of samples $T \geq \frac{M^2}{2\varepsilon^2} \ln \left(\frac{2}{\delta} \right)$, where $M := b - a$.

Lemma 10 (Short-time sample complexity analysis). *Let*

$$B_a^{(j)} = \text{Tr} \left[Q_a e^{-iHt^{(j)}} \left(\bigotimes_{i=1}^n |s_i^{(j)}\rangle \langle s_i^{(j)}| \right) e^{iHt^{(j)}} \right] \quad (\text{C27})$$

be the a^{th} component of the j^{th} label in the training dataset. Furthermore, for each $a \in [M]$, let

$$C_a^{(j)} := \mathbb{1} \left[P_a |s_i^{(j)}\rangle = + |s_i^{(j)}\rangle \right] B_a^{(j)} \quad (\text{C28})$$

and let $\hat{\mathcal{F}}_{a,t^} := \frac{1}{N_{t^*}} \sum_{j=1}^{N_{t^*}} 6C_a^{(j)}$ be an empirical average over training samples whose times belong to the short-time interval $[0, t^*]$, where we use N_{t^*} to denote the number of such training samples. Let $\hat{\lambda} := (\hat{\lambda}_a)_{a \in [M]} := \left(\hat{\mathcal{F}}_{a,t^*} / 2t^* \right)_{a \in [m]}$ be the vector of coefficients that are obtained as the output of Algorithm 1. Then, for a desired precision $\varepsilon > 0$, some confidence parameter $\delta > 0$, and a given observable O , we have $\left| \langle O \rangle_{\rho(\lambda, t)} - \langle O \rangle_{\rho(\hat{\lambda}, t)} \right| \leq \varepsilon$ provided that*

$$t^* \in \mathcal{O} \left(\frac{\varepsilon}{MT \|O\|_{\infty} (\mathfrak{d} + 1)^2} \right), \quad (\text{C29})$$

$$\text{and } N_{t^*} \in \mathcal{O} \left(\frac{M^4 T^4 \|O\|_{\infty}^2 (\mathfrak{d} + 1)^4}{\varepsilon^4} \log \left(\frac{M}{\delta} \right) \right), \quad (\text{C30})$$

with probability at least $1 - \delta$.

Proof. Consider the error decomposition

$$\left| \hat{\lambda}_a - \lambda_a \right| = \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| = \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} + \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} \right| + \left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right|. \quad (\text{C31})$$

Note that we have already bounded the second term by means of Lemma 9. We bound the first term as

$$\begin{aligned} \left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} \right| &\leq \frac{1}{2t^*} \left| \frac{1}{N_{t^*}} \sum_{j=1}^{N_{t^*}} 6C_a^{(j)} - \frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{a,t} dt \right| \\ &= \frac{1}{2t^*} \left| \frac{1}{N_{t^*}} \sum_{j=1}^{N_{t^*}} 6C_a^{(j)} - \mathbb{E}_{s,t} [6C_a(s, t) | t \in [0, t^*]] \right|, \end{aligned} \quad (\text{C32})$$

where $6C_a^{(j)}$ are i.i.d. samples of the random variable $C_a(s, t)$, which is defined as

$$C_a(s, t) := \begin{cases} B_a(s, t), & \text{if } P_a|s_i\rangle = +|s_i\rangle, \text{supp}(P_a) = \{i\}, \\ 0, & \text{otherwise.} \end{cases} \quad (\text{C33})$$

$$\text{where} \quad B_a(s, t) = \text{Tr} \left[Q_a e^{-iHt} \left(\bigotimes_{i=1}^n |s_i\rangle\langle s_i| \right) e^{iHt} \right], \quad (\text{C34})$$

with $(s, t) \stackrel{\text{i.i.d.}}{\sim} \text{Unif}(\{0, 1, +, -, i+, i-\}^n \times [0, T])$, where $T \in \mathcal{O}(\text{poly}(n))$. We are able to replace the time-averaged integral by the expectation value of $6C_a(s, t)$ conditioned on t belonging to the short-time window $[0, t^*]$, where the expectation is taken over s and t , because of the following:

$$\mathbb{E}_{s,t} [6C_a(s, t) \mid t \in [0, t^*]] \quad (\text{C35})$$

$$= \mathbb{E}_t \left[\mathbb{E}_s [6C_a(s, t) \mid t] \mid t \in [0, t^*] \right] \quad (\text{C36})$$

$$\begin{aligned} &= \mathbb{E}_t \left[6 \mathbb{E}_s [\mathbb{1}[P_a|s_i\rangle = +|s_i\rangle] B_a(s, t) \mid t] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[6 \cdot \Pr[P_a|s_i\rangle = +|s_i\rangle \mid t] \mathbb{E}_s [B_a(s, t) \mid P_a|s_i\rangle = +|s_i\rangle, t] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[6 \cdot \frac{1}{6} \cdot \mathbb{E}_s \left[\text{Tr} \left[Q_a e^{-iHt} \left(\bigotimes_{i=1}^n |s_i\rangle\langle s_i| \right) e^{iHt} \right] \mid P_a|s_i\rangle = +|s_i\rangle, t \right] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[\text{Tr} \left[Q_a e^{-iHt} \mathbb{E}_s \left[\bigotimes_{i=1}^n |s_i\rangle\langle s_i| \mid P_a|s_i\rangle = +|s_i\rangle, t \right] e^{iHt} \right] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[\text{Tr} \left[Q_a e^{-iHt} \left\{ \frac{I}{2} \otimes \dots \otimes \underbrace{\left(\frac{I + P_a}{2} \right)}_{i^{\text{th}} \text{ qubit}} \otimes \dots \otimes \frac{I}{2} \right\} e^{iHt} \right] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[\frac{1}{2^n} \underbrace{\text{Tr} [Q_a e^{-iHt} e^{iHt}]}_{=0} + \text{Tr} \left[Q_a e^{-iHt} \left\{ \frac{I}{2} \otimes \frac{I}{2} \otimes \dots \otimes \frac{P_a}{2} \otimes \dots \otimes \frac{I}{2} \right\} e^{iHt} \right] \mid t \in [0, t^*] \right] \\ &= \mathbb{E}_t \left[\frac{1}{2^n} \text{Tr} [Q_a e^{-iHt} P_a e^{iHt}] \mid t \in [0, t^*] \right] = \mathbb{E}_t [\mathcal{F}_{a,t} \mid t \in [0, t^*]] = \frac{1}{t^*} \int_0^{t^*} \mathcal{F}_{a,t} dt, \end{aligned}$$

where the third equality follows from the elementary conditional probability identity $\mathbb{E}[\mathbb{1}_A X] = \Pr[A] \mathbb{E}[X|A]$, where $\mathbb{1}_A$ is equal to 1 if the event A occurs, and 0 otherwise. Going from the sixth to the seventh equality, we have used the traceless property of Q_a , and, per the convention of Ref. [36], replaced the n -qubit operator $I \otimes I \otimes \dots \otimes P_a \otimes \dots \otimes I$ with simply P_a to reduce notational clutter. By Theorem 4, for some failure probability $\eta > 0$, we have that

$$\Pr \left[\frac{1}{2t^*} \left| \hat{\mathcal{F}}_{a,t^*} - \tilde{\mathcal{F}}_{a,t^*} \right| \leq \frac{1}{2t^*} \sqrt{\frac{288}{N_{t^*}} \log \left(\frac{2}{\eta} \right)} \right] \geq 1 - \eta. \quad (\text{C37})$$

Now, setting $\eta \leftarrow \delta/2M$ for some desired confidence parameter $\delta > 0$ and applying the union bound over all $a \in [M]$, we find that

$$\left| \frac{\hat{\mathcal{F}}_{a,t^*}}{2t^*} - \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} \right| + \left| \frac{\tilde{\mathcal{F}}_{a,t^*}}{2t^*} - \lambda_a \right| \leq \frac{1}{2t^*} \sqrt{\frac{288}{N_{t^*}} \log \left(\frac{4M}{\delta} \right)} + \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*} \quad (\text{C38})$$

holds for all $a \in [M]$, with probability at least $1 - \delta/2$, where we have allocated the failure probability budget $\delta/2$ to the training stage. We note here in advance that the other half of the final/inference failure probability δ will be allocated

to the inference step of our algorithm, where we estimate expectation values of numerous observables via the classical shadows formalism. This will ensure that the final prediction outputs will be ε -accurate with probability $1 - \delta$. In the above, we have also plugged in the bound from Eq. (C24). Now, we allocate the error budgets

$$\frac{1}{2t^*} \sqrt{\frac{288}{N_{t^*}} \log \left(\frac{4M}{\delta} \right)} \leq \frac{\varepsilon'}{2}, \quad (\text{C39})$$

$$\text{and} \quad \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*} \leq \frac{\varepsilon'}{2}. \quad (\text{C40})$$

Now, in order to ensure that $|\langle O \rangle_{\rho(\lambda, t)} - \langle O \rangle_{\rho(\lambda', t)}| \leq \varepsilon$, in accordance with Lemma 8, we will set $\varepsilon' \leftarrow \varepsilon / 2MT \|O\|_\infty$. Then, we have the solution

$$\begin{aligned} \frac{2t^* (\mathfrak{d} + 1)^2}{1 - 2(\mathfrak{d} + 1)t^*} &\leq \frac{\varepsilon'}{2} \implies 4(\mathfrak{d} + 1)^2 t^* \leq \varepsilon' - 2(\mathfrak{d} + 1)\varepsilon' t^* \implies \left\{ 4(\mathfrak{d} + 1)^2 + 2(\mathfrak{d} + 1)\varepsilon' \right\} t^* \leq \varepsilon' \quad (\text{C41}) \\ \implies t^* &\leq \frac{\varepsilon'}{4(\mathfrak{d} + 1)^2 + 2(\mathfrak{d} + 1)\varepsilon'} \leq \frac{\varepsilon / 2MT \|O\|_\infty}{4(\mathfrak{d} + 1)^2 + 2(\mathfrak{d} + 1)(\varepsilon / 2MT \|O\|_\infty)} = \frac{\varepsilon}{8MT \|O\|_\infty (\mathfrak{d} + 1)^2 + 2(\mathfrak{d} + 1)\varepsilon} \\ \implies t^* &\in \mathcal{O} \left(\frac{\varepsilon}{MT \|O\|_\infty (\mathfrak{d} + 1)^2} \right) \end{aligned}$$

for t^* . Now, we will solve for the sample complexity N_{t^*} , to get

$$\frac{1}{2t^*} \sqrt{\frac{288}{N_{t^*}} \log \left(\frac{4M}{\delta} \right)} \leq \frac{\varepsilon'}{2} \implies \frac{1}{4t^{*2}} \cdot \frac{288}{N_{t^*}} \log \left(\frac{4M}{\delta} \right) \leq \frac{\varepsilon'^2}{4} \implies N_{t^*} \geq \frac{288}{t^{*2} \varepsilon'^2} \log \left(\frac{4M}{\delta} \right). \quad (\text{C42})$$

Plugging in the solution obtained for t^* , as well as our chosen bound on ε' , we get

$$\begin{aligned} N_{t^*} &\geq \frac{288}{\left\{ \frac{\varepsilon}{8MT \|O\|_\infty (\mathfrak{d} + 1)^2 + 2(\mathfrak{d} + 1)\varepsilon} \right\}^2 \cdot \left(\frac{\varepsilon}{2MT \|O\|_\infty} \right)^2} \log \left(\frac{4M}{\delta} \right) \quad (\text{C43}) \\ &= \frac{288 \cdot 8M^2 T^2 \|O\|_\infty^2 \cdot \left\{ 16M^2 T^2 \|O\|_\infty^2 (\mathfrak{d} + 1)^4 + 16MT \|O\|_\infty (\mathfrak{d} + 1)^3 \varepsilon + 4(\mathfrak{d} + 1)^2 \varepsilon^2 \right\}}{\varepsilon^4} \log \left(\frac{4M}{\delta} \right) \\ &= \frac{9216M^2 T^2 \left\{ 2MT \|O\|_\infty (\mathfrak{d} + 1)^2 + (\mathfrak{d} + 1)\varepsilon \right\}^2}{\varepsilon^4} \log \left(\frac{4M}{\delta} \right) \in \mathcal{O} \left(\frac{M^4 T^4 \|O\|_\infty^2 (\mathfrak{d} + 1)^4}{\varepsilon^4} \log \left(\frac{M}{\delta} \right) \right). \end{aligned}$$

□

Lemma 11 (Total sample complexity analysis). *Let N be the total sample complexity of the training dataset:*

$$\mathcal{T} := \left\{ \left(s^{(j)}, t^{(j)}, \left(B_a^{(j)} \right)_{a \in [m]} \right) \right\}_{j=1}^N, \quad (\text{C44})$$

which is given as input to Algorithm 1. Also, let N_{t^*} be the short-time sample complexity, as in Lemma 10. Then, for Lemma 10 to hold, N must scale as

$$N \in \mathcal{O} \left(\frac{M^5 T^6 \|O\|_\infty^3 (\mathfrak{d} + 1)^6}{\varepsilon^5} \log \left(\frac{M}{\delta} \right) \right). \quad (\text{C45})$$

Proof. For each $j \in [N]$, we define

$$X^{(j)} := \mathbb{1} \left[t^{(j)} \in [0, t^*] \right] \mathbb{1} \left[x_1^{(j)} = 0 \right] = \begin{cases} 1, & \text{if } t^{(j)} \in [0, t^*] \text{ and } x_1^{(j)} = 0, \\ 0, & \text{otherwise,} \end{cases} \quad (\text{C46})$$

so that $X^{(j)} \in \{0, 1\}$ is a Bernoulli random variable, since $\Pr[X^{(j)} = 1] = t^*/2T$ and $\Pr[X^{(j)} = 0] = 1 - t^*/2T$. The number of short-time samples is then given as

$$N_{t^*} := \sum_{j=1}^N X^{(j)}, \quad (\text{C47})$$

$$\text{so that } \mu := \mathbb{E}[N_{t^*}] = \sum_{j=1}^N \mathbb{E}[X^{(j)}] = \frac{Nt^*}{2T}. \quad (\text{C48})$$

so that $N_{t^*} \sim B(N, t^*/2T)$ is itself a binomial random variable, thereby allowing us to use the Chernoff bound. For $0 \leq \eta \leq 1$, it holds that

$$\Pr[N_{t^*} \leq (1 - \eta)\mu] \leq \exp\left(-\frac{\eta^2\mu}{2}\right) \implies \Pr\left[N_{t^*} \leq \frac{(1 - \eta)Nt^*}{2T}\right] \leq \exp\left(-\frac{\eta^2Nt^*}{4T}\right). \quad (\text{C49})$$

As a notational shorthand, we will set

$$N_{\text{short}} := \frac{2304M^2T^2 \left\{ 2MT \|O\|_\infty (\mathfrak{d} + 1)^2 + (\mathfrak{d} + 1)\varepsilon \right\}^2}{\varepsilon^4} \log\left(\frac{4M}{\delta}\right) \quad (\text{C50})$$

to be the required number of short-time samples calculated in Lemma 10. Additionally, we will also require that

$$\frac{(1 - \eta)Nt^*}{2T} = N_{\text{short}} \implies \eta = 1 - \frac{2N_{\text{short}}T}{Nt^*}. \quad (\text{C51})$$

Next, in the implication of Eq. (C49), we will set the R.H.S. to be less than or equal to some δ_{short} , and then solve for N . We proceed as follows.

$$\begin{aligned} \exp\left[-\frac{Nt^*}{4T} \left(1 - \frac{2N_{\text{short}}T}{Nt^*}\right)^2\right] &\leq \delta_{\text{short}} \implies \exp\left[-\frac{1}{2} \left\{ \frac{Nt^*}{2T} - 2N_{\text{short}} + \frac{2N_{\text{short}}^2T}{Nt^*} \right\}\right] \leq \delta_{\text{short}} \\ &\implies \frac{Nt^*}{2T} - 2N_{\text{short}} + \frac{2N_{\text{short}}^2T}{Nt^*} \geq 2 \log\left(\frac{1}{\delta_{\text{short}}}\right). \end{aligned} \quad (\text{C52})$$

The above may be re-written as the quadratic inequality in $Nt^*/2T$

$$\left(\frac{Nt^*}{2T}\right)^2 - \left(2N_{\text{short}} + 2 \log\left(\frac{1}{\delta_{\text{short}}}\right)\right) \frac{Nt^*}{2T} + N_{\text{short}}^2 \geq 0. \quad (\text{C53})$$

After application of the quadratic formula and rearrangement, we will find that

$$\begin{aligned} N &\geq \frac{2T}{t^*} \left[N_{\text{short}} + \log\left(\frac{1}{\delta_{\text{short}}}\right) + \sqrt{\log\left(\frac{1}{\delta_{\text{short}}}\right) \left(2N_{\text{short}} + \log\left(\frac{1}{\delta_{\text{short}}}\right)\right)} \right] \\ &\in \mathcal{O}\left(\frac{T}{t^*} N_{\text{short}}\right) \subseteq \mathcal{O}\left(\frac{M^5T^6 \|O\|_\infty^3 (\mathfrak{d} + 1)^6}{\varepsilon^5} \log\left(\frac{M}{\delta}\right)\right), \end{aligned} \quad (\text{C54})$$

holds true, where, in the final step, we have plugged in the previously obtained scalings for N_{short} and t^* in Lemma 10. $\square$

Lemma 12 (Sample complexity of inference via classical shadows). *Let $(\hat{\lambda}_a)_{a \in [M]}$ be the learned Hamiltonian coefficients obtained as outputs of Algorithm 1 such that we have, for all $a \in [M]$ and all $(Q_a)_{a \in [M]}$, $|\hat{\lambda}_a - \lambda_a| \leq \varepsilon/2MT \max_{j \in [P]} \|O_j\|_\infty$ with probability $1 - \delta/2$, where $\varepsilon, \delta > 0$ are the final precision and confidence parameters*

associated with the inference stage. Then, for a newly provided input $(s, t) \sim \text{Unif}(\{0, 1, +, -, i+, i-\}^n \times [0, T])$, in order to ensure $|\langle \hat{Q}_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\rho(s,t)}| \leq \varepsilon$, where

$$\hat{\rho}(s, t) := e^{-i\hat{H}t} \left(\bigotimes_{i=1}^n |s_i\rangle\langle s_i| \right) e^{i\hat{H}t}, \quad \hat{H} := \sum_{a \in [M]} \hat{\lambda}_a E_a, \quad (\text{C55})$$

$$\rho(s, t) := e^{-iHt} \left(\bigotimes_{i=1}^n |s_i\rangle\langle s_i| \right) e^{iHt}, \quad H = \sum_{a \in [M]} \lambda_a E_a, \quad (\text{C56})$$

$$\text{and} \quad \|\tilde{\rho}(s, t) - \hat{\rho}(s, t)\|_1 \leq \varepsilon_{\text{sim}}, \quad (\text{C57})$$

with $\varepsilon_{\text{sim}} > 0$ being some state preparation error, we require $\tilde{N} \in \mathcal{O}(\frac{1}{\varepsilon^2} \log(\frac{M}{\delta}))$ copies of $\hat{\rho}(s, t)$ to be prepared in Algorithm 2.

Proof. The proof follows straightforwardly from the classical shadows formalism. Consider the error decomposition

$$\begin{aligned} \left| \langle \hat{Q}_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\rho(s,t)} \right| &= \left| \left(\langle \hat{Q}_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\tilde{\rho}(s,t)} \right) + \left(\langle Q_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\hat{\rho}(s,t)} \right) + \left(\langle Q_a \rangle_{\hat{\rho}(s,t)} - \langle Q_a \rangle_{\rho(s,t)} \right) \right| \\ &\leq \underbrace{\left| \langle \hat{Q}_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\tilde{\rho}(s,t)} \right|}_{\leq \varepsilon_{\text{shadow}}} + \underbrace{\left| \langle Q_a \rangle_{\tilde{\rho}(s,t)} - \langle Q_a \rangle_{\hat{\rho}(s,t)} \right|}_{\leq \|Q_a\|_\infty \|\tilde{\rho}(s,t) - \hat{\rho}(s,t)\|_1 \leq 2\varepsilon_{\text{sim}}} + \underbrace{\left| \langle Q_a \rangle_{\hat{\rho}(s,t)} - \langle Q_a \rangle_{\rho(s,t)} \right|}_{\leq \varepsilon_{\text{exp}}} \leq \varepsilon, \end{aligned} \quad (\text{C58})$$

so that we have $\varepsilon_{\text{shadow}} + \varepsilon_{\text{sim}} + \varepsilon_{\text{exp}} \leq \varepsilon$. We set $\varepsilon_{\text{shadow}} \leftarrow \varepsilon/3$, $\varepsilon_{\text{sim}} \leftarrow \varepsilon/6$, and $\varepsilon_{\text{exp}} \leftarrow \varepsilon/3$, where the first and third each holds true with probability $1 - \delta/2$. Recall again that, here, we are allocating the leftover failure probability of $\delta/2$ from the training stage. Then, using Theorem 1 of Ref. [108], we require

$$m = 2 \log \left(\frac{4M}{\delta} \right), \quad (\text{C59})$$

$$\text{and} \quad G = \frac{136}{\varepsilon^2} \max_{a \in [M]} \underbrace{\left\| Q_a - \frac{\text{Tr}[Q_a]}{2^n} \mathbb{I} \right\|_{\text{shadow}}^2}_{\leq 4 \cdot 3^{|\text{supp}(E_a)|} \text{ (for Pauli measurements)}}, \quad (\text{C60})$$

$$\text{so that} \quad \tilde{N} \leq Gm \in \mathcal{O} \left(\frac{1}{\varepsilon^2} \log \left(\frac{M}{\delta} \right) \right), \quad (\text{C61})$$

where G and m are the parameters of the median-of-means estimation procedure, as in Subroutine 2. $\square$

Theorem 5 (Quantum learnability, formal). *The concept class of Def. 6 is efficiently quantumly PAC-learnable; that is, Algorithms 1 and 2 require a total sample complexity*

$$N \in \mathcal{O} \left(\frac{M^5 T^6 \|O\|_\infty^3 (\mathfrak{d} + 1)^6}{\varepsilon^5} \log \left(\frac{M}{\delta} \right) \right), \quad (\text{C62})$$

and a runtime of $\mathcal{O}(MN)$ and $\mathcal{O}(\tilde{N}(M + t_{\text{sim}}))$, where $t_{\text{sim}} \in \mathcal{O}(\text{poly}(n, 1/\varepsilon))$, in order to satisfy the learning condition of Def. 8.

Proof. The sample complexity shown in Eq. (C62) follows from Lemma 11. Let us now thoroughly analyze the time complexities of Algorithms 1 and 2. For each $a \in [M]$ and $j \in [N]$, Algorithm 1 involves checking whether $t^{(j)} \in [0, t^*]$, $x_1^{(j)} = 0$, and $P_a |s_i^{(j)}\rangle = +|s_i^{(j)}\rangle$. Each of these checks is $\mathcal{O}(1)$; the Pauli eigenstate check being constant-time follows

from the fact that P_a is a fixed single qubit Pauli and $s_i^{(j)}$ can only be one of six symbols. This is followed by a $\mathcal{O}(M)$ cost of performing elementary arithmetic operations to obtain the estimates $(\hat{\lambda}_a)_{a \in [M]}$. Therefore, the time complexity of training is dominated by $\mathcal{O}(MN)$. We move on to Algorithm 2, the first step of which requires us to prepare, for each $\ell \in [\tilde{N}]$, the state $\rho_{\text{init}}^{(\ell)}(s)$, which can be done in $\mathcal{O}(n)$ time and $\mathcal{O}(1)$ depth. Then, we perform a time-evolution step, which, abstractly and combined with the previous initial state preparation step, incurs a cost of $t_{\text{sim}} \in \mathcal{O}(\text{poly}(n, t, 1/\varepsilon_{\text{sim}})) \subseteq \mathcal{O}(\text{poly}(n, 1/\varepsilon))$ per copy. Finally, the runtime of the classical post-processing is given by $\mathcal{O}(\tilde{N}M)$. The total time complexity of inference is then given as $\mathcal{O}(\tilde{N}M + \tilde{N}t_{\text{sim}})$. $\square$

5. Quantum learnability of the classically hard instance

Lemma 13 (Sample- and time-complexity of learning $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ using Algorithm 1). *Let H_{BQP} be the Hamiltonian family as in Theorem 2, and let $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ be the corresponding concept class defined per Def. 6. Then, with probability at least $1 - \delta$, $\mathcal{C}_{U_{\text{enc}}, T}^{\text{H}_{\text{BQP}}}$ is PAC-learnable in the sense of Def. 8 via Algorithm 1 using a number of samples scaling as*

$$N_{\text{BQP}} \in \mathcal{O}\left(T \cdot \frac{(\mathfrak{d}+1)^6}{\kappa^5} \log\left(\frac{M}{\delta}\right)\right) \subseteq \mathcal{O}\left(T \log\left(\frac{M}{\delta}\right)\right), \quad (\text{C63})$$

where, as described in Lemma 1 and for $\varepsilon_P \approx 0.66$, $\kappa := \varepsilon'_P/2 = \varepsilon_P/64$ is the minimum estimation precision required to learn $H_{\text{BQP}}(\lambda)$.

Proof. Recall, from Lemma 9, that we have

$$\left| \frac{\tilde{\mathcal{F}}_{a,t}}{2t^*} - \lambda_a \right| \leq \frac{2(\mathfrak{d}+1)^2 t^*}{1 - 2(\mathfrak{d}+1)t^*}. \quad (\text{C64})$$

Let us denote by t_{BQP}^* the short-time cutoff for the classically hard case. We then require that

$$\begin{aligned} \frac{2(\mathfrak{d}+1)^2 t_{\text{BQP}}^*}{1 - 2(\mathfrak{d}+1)t^*} &\leq \frac{\kappa}{2} \implies 4(\mathfrak{d}+1)^2 t_{\text{BQP}}^* \leq \kappa - 2(\mathfrak{d}+1)\kappa t_{\text{BQP}}^* \\ &\implies \left\{ 4(\mathfrak{d}+1)^2 + 2(\mathfrak{d}+1)\kappa \right\} t_{\text{BQP}}^* \leq \kappa \\ &\implies t_{\text{BQP}}^* \leq \frac{\kappa}{4(\mathfrak{d}+1)^2 + 2(\mathfrak{d}+1)\kappa} \\ &\implies t_{\text{BQP}}^* \leq \frac{\varepsilon_P/64}{4(\mathfrak{d}+1)^2 + 2(\mathfrak{d}+1) \cdot (\varepsilon_P/64)} \\ &\implies t_{\text{BQP}}^* \leq \frac{66}{25600(\mathfrak{d}+1)^2 + 132(\mathfrak{d}+1)} \in \mathcal{O}\left(\frac{1}{(\mathfrak{d}+1)^2}\right) \in \mathcal{O}(1) \end{aligned} \quad (\text{C65})$$

holds true, with probability at least $1 - \delta/2$. On the other hand, we also require

$$\begin{aligned} \frac{1}{2t_{\text{BQP}}^*} \sqrt{\frac{288}{N_{t_{\text{BQP}}^*}^*} \log\left(\frac{4M}{\delta}\right)} &\leq \frac{\kappa}{2} \implies \frac{1}{4t_{\text{BQP}}^{*2}} \cdot \frac{288}{N_{t_{\text{BQP}}^*}^*} \log\left(\frac{4M}{\delta}\right) \leq \frac{\kappa^2}{4} \\ &\implies N_{t_{\text{BQP}}^*}^* \geq \frac{288}{t_{\text{BQP}}^{*2} \kappa^2} \log\left(\frac{4M}{\delta}\right), \end{aligned} \quad (\text{C66})$$

where $N_{t_{\text{BQP}}^*}^*$ is the number of short-time samples. Now, let us compute the total number of samples, N_{BQP} , following the logic of Lemma 11. By analogy with Eq. (C54), we will find that N_{BQP} satisfies

$$N_{\text{BQP}} \geq \frac{2T}{t_{\text{BQP}}^*} \left[N_{t_{\text{BQP}}^*}^* + \log\left(\frac{1}{\delta_{\text{BQP}}}\right) + \sqrt{\log\left(\frac{1}{\delta_{\text{BQP}}}\right) \left(2N_{t_{\text{BQP}}^*}^* + \log\left(\frac{1}{\delta_{\text{BQP}}}\right) \right)} \right] \in \mathcal{O}\left(\frac{T}{t^*} N_{t_{\text{BQP}}^*}^*\right). \quad (\text{C67})$$

Plugging in the previously acquired expressions for t^* and $N_{t_{\text{BQP}}^*}$, we finally obtain

$$N_{\text{BQP}} \in \mathcal{O} \left(T \cdot \frac{(\mathfrak{d} + 1)^6}{\kappa^5} \log \left(\frac{M}{\delta} \right) \right) \in \mathcal{O} \left(T \log \left(\frac{M}{\delta} \right) \right), \quad (\text{C68})$$

where we have used the facts that $\mathfrak{d} \in \mathcal{O}(1)$ and $\kappa = \varepsilon_P/64 \in \mathcal{O}(1)$. The time-complexity of reconstructing H_{BQP} is then simply given as $\mathcal{O}(MN_{\text{BQP}})$. $\square$