

Characterizing quantum hardware with imperfect control

Inaugural dissertation

presented to the faculty of Mathematics and Natural Sciences
of Heinrich-Heine-University Düsseldorf
for the degree of
Doctor of Natural Sciences (Dr. rer. nat.)

by
Raphael Brieger
from Penzberg, Germany

Düsseldorf, June 2024

From the Institute of Theoretical Physics III
at the Heinrich-Heine-University Düsseldorf

Published with permission of the
Faculty of Mathematics and Natural Sciences at
Heinrich-Heine-University Düsseldorf

Supervisor: Prof. Dr. Martin Kliesch
Co-supervisor: Prof. Dr. Dagmar Bruf

Date of the oral examination: 18.06.2024

Declaration

Ich versichere an Eides Statt, dass die Dissertation von mir selbstständig und ohne unzulässige fremde Hilfe unter Beachtung der „Grundsätze zur Sicherung guter wissenschaftlicher Praxis an der Heinrich-Heine-Universität Düsseldorf“ erstellt worden ist.

Düsseldorf, 14.03.2024

Raphael Brieger

Zusammenfassung

Quantencomputer versprechen eine Vielzahl an Rechenaufgaben schneller zu lösen als herkömmliche Computer, unter anderem die Primfaktorzerlegung, die Suche in einer unstrukturierten Datenbank, und die Simulation von Systemen in der Quantenchemie. Wir befinden uns momentan in einer Phase in der mit bereits verfügbaren Quantensystemen tatsächlich Aufgaben gelöst werden, welche auf einem herkömmlichen Supercomputer etliche Jahre in Anspruch nehmen würden. Allerdings sind diese Aufgaben noch nicht von praktischem Nutzen. Um diese frühen Quantensysteme zu verbessern sind zuverlässige und effiziente Methoden notwendig, die einerseits die korrekte Funktionsweise verifizieren können, und andererseits in der Lage sind noch vorhandene Fehlerquellen zu charakterisieren. Durch den laufenden Fortschritt und immer größer werdende Systeme werden auch die Anforderungen an diese Charakterisierungs- und Verifizierungsmethoden stetig erhöht. Zum einen ist eine höhere Effizienz in der Anzahl an benötigten Messungen und der Rechenzeit zur Auswertung dieser Messungen nötig, und zum anderen müssen immer kleiner werdende Fehler noch aufgelöst werden. Um letzteres zu erreichen ist es auch vonnöten Fehler im Messprozess durch fehlerhafte Kontrolloperationen mit einzubeziehen. Aus diesem Grund wurden in sich selbst konsistente Protokolle vorgeschlagen, welche nicht nur einzelne Kontrollprozesse (Gatter) charakterisieren, sondern ein Modell des gesamten Systems mit allen Kontrolloperationen gemeinsam. Diese Herangehensweise ist unter dem Namen Gatterset-Tomographie bekannt, und sie wurde mittlerweile zu einer Standardtechnik zur Charakterisierung und Verbesserung von Quantensystemen. Wie zu erwarten führt das lernen eines Modells des gesamten Systems zu hohen Kosten im Messaufwand und einer hohen Rechenzeit in der Auswertung. Aus diesem Grund wird Gatterset-Tomographie meist sparsam in einem Experiment angewandt und auch nur auf kleinen Subsystemen.

Der zentrale Forschungsbeitrag dieser Arbeit ist die Entwicklung eines Datenverarbeitungssystems für Gatterset-Tomographie, welches mit weniger Messeinstellungen und Rechenaufwand als zuvor möglich auskommt. Dies erreichen wir, indem wir ein komprimiertes Modell schätzen (lernen), welches nur die relevantesten Freiheitsgrade des Systems erfasst. Da das Lernen eines solchen Modells dem Lösen eines hoch nicht-trivialen Optimierungsproblems entspricht, entwickeln wir einen neuen Optimierungsalgorithmus auf der Riemann'schen Mannigfaltigkeit physikalischer Gatter, welches mit Hilfe von Techniken aus dem Feld des maschinellen Lernens optimale Lösungen findet. Um den Algorithmus für Experimentatoren zugänglich zu machen, entwickeln wir ein öffentlich verfügbares Python-Paket. Auf dessen Basis demonstrieren wir die genaue Rekonstruktion eines zwei Qubit Systems von experimentellen Daten aus einem Ionenfallen-Experiment.

Eine weitere Aufgabe der wir uns in dieser Arbeit widmen, ist die Schätzung von physikalischen Eigenschaften eines Quantenzustands, von dem angenommen wird, er lasse sich wiederholt in einem Experiment erzeugen. Vorangegangenen Arbeiten zeigten, dass man mit zufälligen Messeinstellungen eine Vielzahl an physikalischen Eigenschaften gleichzeitig schätzen kann, wobei die Anzahl an dafür benötigten Messungen unabhängig von der Systemgröße bleibt. Allerdings wurde dabei bisher der Einfluss fehlerhafter Kontrolloperationen weitgehend außer Acht gelassen. In dieser Arbeit analysieren wir den Effekt von Gatterabhängigen Fehlern auf Protokolle, die auf zufälligen Messeinstellungen basieren. Von zentraler Bedeutung in unserer Analyse ist die Herleitung mathematischer Schranken, welche zeigen, dass diese Protokolle in den meisten Fällen eine

inhärente Resistenz gehen Fehler aufweisen. Im Gegensatz dazu identifizieren wir auch Eigenschaften eines Quantensystems, deren Schätzung fehleranfällig sein kann. Insbesondere können Fehler in den Kontrolloperationen zu einem exponentiell verstärkten Fehler in der Schätzung führen.

In ihrer Gesamtheit tragen die Resultate in dieser Arbeit dazu bei, wichtige Charakterisierungsprotokolle für frühe Quantencomputer signifikant besser mit den experimentellen Anforderungen in Einklang zu bringen.

Abstract

Quantum computers promise to solve a variety of computational problems faster than existing classical computers, among them the factorization of prime numbers, unstructured database search and the simulation of quantum chemistry systems. We are right in the era where actual computational advantages of current early quantum computers are demonstrated, albeit not yet in computation tasks of practical interest. In order to continuously improve these systems, reliable and efficient verification of their correct operation, as well as methods to characterize remaining error sources are vital. The ongoing maturing process of quantum computers comes with tighter demands on characterization and verification tasks. Not only is there an increased need for efficiency in the number of measurements and classical post-processing time, but progressively smaller errors need to be resolved as well. In order to achieve the latter task and to estimate ('learn') properties of quantum systems in a reliable way, errors in the measurement process due to imperfect control need to be taken into account. For this reason, self-consistent protocols were proposed, which not only characterize individual quantum operations (gates), but learn a mathematical model of the system and its dynamics as a whole. This approach is known as gate set tomography, and it has become a standard technique to characterize and improve quantum experiments. Unsurprisingly, learning a model of the whole system comes at a huge cost in the measurement and classical post-processing effort. For this reason it has only been applied sparingly and to small subsystems.

The main contribution of this thesis is the development of a new data processing framework for gate set tomography, which allows us to construct a mathematical model of a subsystem from few random measurement settings with a reduced post-processing time compared to the previous state-of-the-art method. We achieve this by learning a compressed model, which reduces the number of parameters that need to be learned, while still capturing the most relevant degrees of freedom. Since the underlying optimization problem for learning a compressed gate set model is highly non-trivial to solve, we develop a new optimization algorithm on the Riemannian manifold of physical gate sets, which uses techniques from machine learning to arrive at an optimal point. We make this algorithm accessible through a publicly available Python package and use it to accurately learn a full description of a two-qubit trapped ion system.

A different task that we focus on in this thesis is the estimation of physical properties of a quantum state, which we assume can be repeatedly prepared on the system. Previous works showed that using randomized measurements, numerous properties of the quantum state can be estimated, while the number of required measurements remains independent of the system size. Thus far, errors introduced due to imperfect control in the implementation of randomized measurements were insufficiently accounted for. In this thesis we analyze the effects of general gate-dependent noise on the randomized measurement protocol. Central to our results are analytical bounds, which show that for most properties of interest, the randomized measurement protocol is resilient against general gate-dependent noise. However, we also find that certain properties are more difficult to estimate, since for those properties there are noise models under which errors are exponentially amplified in the estimation procedure.

Overall, the results of this thesis contribute to bringing important characterization protocols for early quantum computers significantly closer to current experimental requirements.

Acknowledgment

I first and foremost want to thank my supervisor Martin Kliesch for always being supportive, for teaching me many things about our field and for countless discussions that left me excited and gave me new motivation to push forward. A special thanks goes to Ingo Roth, who took the role of an unofficial co-supervisor for me and always had great ideas and a highly contagious positivity. I also want to thank Lennart Bittel for all the interesting discussions and for his friendship throughout our time in Düsseldorf. My thanks further go to Markus Heinrich for a very enjoyable and productive cooperation and for teaching me about representation theory. I am moreover very grateful to the people from Siegen, Markus Nünnerich, Patrick Huber and Pau Dietz Romero, for our fruitful cooperation.

In addition, I want to thank all the people that made my life enjoyable and helped me in difficult times during the Thesis, among them the people in our group and in the group of Dagmar Bruß, as well as my friends Kilian Ender, Yaiza Aragonés-Soria, Oran Greier, Maximilian George, Florian Huber, Gláucia Murta Guimarães and Andres Vargas-Toscano.

A special thanks for proofreading the thesis and for very helpful comments goes out to Christopher Cedzich, Kilian Ender, Markus Heinrich, Gláucia Murta Guimarães, Ingo Roth and Matthias Zipper.

Overall, I want to thank my family, especially my brother Maximilian and my parents Reinhild and Michael, for their unconditional love and support throughout my life. Finally, my biggest gratitude goes to Jihene, for her patience, her love, and for all the happiness that she brought me.

Contents

Table of contents	x
1 Introduction	1
2 Theoretical background	5
2.1 Mathematical preliminaries	5
2.1.1 Notation	5
2.1.2 Representation theory	6
2.1.3 Optimization on matrix manifolds	9
2.2 Characterization of quantum dynamics	18
2.2.1 Quantum operations and the superoperator formalism	19
2.2.2 Distance measures	25
2.2.3 Process tomography	29
2.2.4 Gate set tomography	41
2.3 Estimating quantum state properties with randomized measurements	49
3 Compressive gate set tomography	54
3.1 The compressive GST framework and algorithm	54
3.2 Characterizing quantum gates on a trapped ion system using mGST	59
3.2.1 Single qubit GST	60
3.2.2 Two qubit GST	62
3.3 Improving shadow estimation with compressive GST	63
4 Shadow estimation under gate-dependent noise	66
5 Conclusion	69
6 Appendix	84
A Paper - Compressive gate set tomography	84
B Paper - Stability of classical shadows under gate-dependent noise	118
C Full reports of compressive GST on a trapped ion system	145
D Tutorial notebook for the mGST python package	160

Chapter 1

Introduction

Since the invention of the transistor [1], we have witnessed a rapid increase in capabilities of computer systems, which are nowadays omnipresent in daily life. This technological development was made possible through the theory of quantum mechanics and its ability to explain the physics of semiconductors. Yet it was argued since the 1980s [2] that we are thus far not harnessing the full potential of a quantum mechanical system, since exclusively quantum effects like superposition and entanglement remained out of our control. It is thus not surprising that high hopes are placed on quantum computing, which is expected by many to encompass a second computer revolution with a lasting impact on society. In fact the 1980s and 1990s saw the invention of promising applications, like Shor’s algorithm [3], which can perform prime factorization exponentially faster than best known classical algorithms, Grover’s [4] algorithm that sees a quadratic advantage in database search and quantum key distribution [5, 6], which offers the possibility of secure encryption without the need for assumptions on an adversary’s computing power. Apart from standard quantum algorithms, proposed use cases for quantum computing also include the simulation of quantum mechanical systems of interest in physics and chemistry [7]. For near term applications, so-called variational quantum algorithms (VQEs) [8] are promising. VQEs are quantum-classical algorithms that aim to approximate the ground state of a Hamiltonian or to solve combinatorial optimization problems.

Quantum algorithms assume the existence of a universal quantum computer [9], i.e. a device which can implement any unitary operation to vanishing error. It was shown that a universal quantum computer can be realized via circuits composed of a small set of controllable unitary operations (gates) [10]. How such gates can be implemented was first theoretically proposed in the 1990s [11, 12] and subsequently implemented in small scale experiments [13, 14]. Even though almost 30 years have past since, we are only now seeing the demonstration of what is termed a quantum advantage—the solution of a computational problem by a quantum computer which can not be solved by any existing classical computer in a reasonable amount of time [15–18]. Quantum advantage claims still come with the caveat that a faster classical algorithm for the studied task might be devised, which has already happened for some of the proposals [19, 20].

What slows the progress of quantum computing is the difficulty of isolating a quantum system from the environment and performing precise control operations without unwanted side effects. In practice, each gate applied in a circuit comes with its own noise contribution, either due to noise induced by the realization of the gate or background noise present during the gate time. Due to the unique properties of quantum systems such as the inability to copy a quantum state (known as the no-cloning theorem) standard error correction techniques from classical computing can not be directly applied and quantum error correction [21] had to be developed in its own right. In quantum error correction, the state of a single qubit used for computation is encoded into the state of many physical qubits. It was shown that provided errors on the physical qubits are below a given threshold, arbitrarily large circuits can be error corrected [22]. However, this still requires a large experimental overhead in practice, and current devices are just reaching

the size where error correction can be demonstrated as a proof of principle [23, 24]. In the meantime, we are stuck with what is termed noisy intermediate scale quantum (NISQ) devices [25], which need to mitigate errors, rather than fully correct them. The term quantum device is often used as an umbrella term that encompasses not only universal quantum computers, but also specialized systems used for quantum cryptography, quantum sensing or other applications that use quantum mechanical properties.

For the further development in the current NISQ era, it is of central importance to assess the performance of quantum devices via benchmarks and to reliably learn which errors occur, such that they can be fixed or at least mitigated. These tasks are addressed by the sub-field of quantum characterization and benchmarking [26, 27], to which we contribute in this thesis. Learning a mathematical description of the inner workings of a quantum device (characterization) is complicated by several factors. First, quantum measurements are inherently probabilistic and in order to estimate a single outcome probability or an expectation value, many repetitions (shots) of an experiment need to be performed. Second, the number of parameters needed to describe a quantum system grows exponentially in the number of qubits, making it impossible to learn a full mathematical description of even moderate scale quantum systems. Third, characterizing a noisy device needs to be done with control operations and measurements which are plagued by noise themselves. It is thus difficult to assign the blame for deviations from the ideal measurement outcomes to a specific component of a quantum experiment, since any of them might be responsible. To solve the latter problem, *self-consistent* approaches were developed that take errors in the probing operations into account [28–30] by simultaneously characterizing the measurement mechanism, the initial state in which an experiment is prepared, as well as all gates that are applied. This approach is known as *gate set tomography* (GST) and it has become the default method for the characterization of small subsystems of 1-2 qubits. The reason for this limitation to small systems lies again in the dimensionality of the problem: A single n -qubit gate is described by 2^{4n} parameters, and an entire set of gates has to be learned simultaneously. Conventional gate sets which are universal for quantum computing typically consist of single- and two-qubit gates and learning the errors of these individual gates is already a crucial step in understanding a system’s noise profile. What complicates things in real experiments is the phenomenon of *cross-talk*: The unwanted side effects of an operation on qubits that were not part of said operation. Characterizing cross-talk requires methods that can learn the action of a gate on the qubits where it acts on, as well as on neighboring qubits. This makes it highly desirable to extend gate set tomography to larger system sizes. An additional reason is that for many platforms such as trapped ions or neutral atoms with Rydberg interactions, multiple qubits are entangled with the use of a single unitary.

An adjacent task to device characterization is benchmarking, which aims to give quality measures for single gates or entire quantum devices, with the aim to compare them between different experimental platforms. Such quality measures could be computed from a mathematical model obtained from a characterization experiment, but this is rarely done due to the demand for scalability of a benchmarking method. Thankfully, several methods have been devised that can assess the quality of much larger quantum systems by directly estimating a quality measure from a constant number of data points. The most widely used is randomized benchmarking (RB) [31, 32], which estimates a single noise parameter quantifying the average quality of a set of gates, using measurements on random gate sequences. These sequences however require large depths for current experiments of more than 50 qubits, and current noise levels still lead to prohibitively large noise accumulations on long sequences. For this reason various extensions and new methods were recently proposed, promising among them are cross entropy benchmarking (used for the first quantum advantage demonstration [15]), quantum volume [33, 34] and randomized benchmarking with mirror circuits [35].

Apart from learning about noise we would also like to learn about the physics of a device, for instance through the estimation of ground state energies, correlation functions and entanglement measures. The main challenge remains the same as for the above case of noise characterization:

We can not learn the full state of the system efficiently. In some instances observables can be measured directly, for instance the expectation value of a multi-qubit tensor product of Pauli matrices can be measured using entangling gates and additional qubits. Yet this procedure comes with its own errors, since entangling gates are still imperfect in current quantum hardware. A unifying approach to estimating observables and higher order functions of quantum states was introduced very recently under the names of *shadow estimation* [36] and *randomized measurements* [37]. Here, the observable or property in question does not have to be known before the experiment. Instead, a classical description of the quantum state (called the shadow) is learned from random measurements. This description was shown to be efficiently learnable, i.e.

the number of required measurements does not scale with the system dimension. Furthermore, many observables can be estimated to low error with high probability from the classical description in post-processing, provided the observables satisfy certain conditions. Although shadow estimation has been very successful and has led to countless applications [37–43], it is thus far still plagued by a lack of self-consistency: Errors in the measurement process could only be accounted for in special cases [44].

In this thesis we take a significant step toward bringing the frameworks of GST and shadow estimation closer to real world experimental requirements. With respect to GST, we introduce a new approach that is able to reconstruct a compressed description of the quantum device. This description reduces the number of required measurements and the computation time for classical post-processing, while at the same time covering the physically most relevant degrees of freedom. To be able to reconstruct the compressive model from few measurement settings and under physicality constraints, we develop an optimization algorithm using the tensor network formalism, optimization on Riemannian manifolds and tools from machine learning. For the optimization method, we further derive analytic expressions of the Riemannian Hessian and geodesics on the manifold of physical device descriptions. In numerical simulations, we demonstrate that this approach either outperforms or matches previous state-of-the-art methods in accuracy, while reducing the classical post-processing time. This makes it possible for us to reconstruct 3 qubit gate sets from few random sequences. Furthermore, we develop a Python package called *mGST*, which implements our algorithm and allows users to quickly analyze data via a tutorial. To showcase a use case, we perform numerical simulations of a shadow estimation protocol under realistic noise, and learn a compressive description of the noise model via mGST. This description then allows us to partly correct errors in the shadow estimation protocol, reducing them by an order of magnitude. Going beyond numerical demonstrations, we use mGST to characterize single- and two-qubit gates in a trapped ion quantum experiment. The resulting analysis allows us to identify different errors sources and to quantify their contribution. Furthermore, we construct a compressed and an uncompressed error model for the two-qubit gates and show that both are in excellent agreement, thus validating the compressive description.

Apart from the error mitigation technique for shadow estimation using mGST, we also fill a critical gap in the theoretical understanding of the effect of noise in shadow estimation. To show why noise can indeed be a large problem, we give analytical examples of noise models that lead to an exponentially large bias in the standard shadow estimation protocol [36], as well as in the most prominent noise mitigation approach [44]. We complement these examples with tight analytical upper bounds on the bias under general noise models. The analytical upper bounds paint a much more hopeful picture: For most observables of interest, noise does not accumulate, in fact it averages out. Only observables with high ‘magic’ - a property often used in the literature to quantify the hardness of classical simulation - can lead to a noise accumulation effect. We complement these results with bounds on the variance of the shadow estimator, which tell us that for the same low magic observables, noise does not significantly change the number of measurements needed for a given accuracy.

The content of this thesis is organized as follows.

- In Chapter 2 we start by giving an introduction to the mathematical techniques which are essential for the understanding of our proofs.

- Section 2.1.2 contains elementary definitions and results of representation theory, as well a derivation of the methods used in the quantum information literature to compute moments over unitary representations. These derivations form the basis of our bounds on the variance of the shadow estimator.
 - Section 2.1.3 reviews the basic definitions of matrix manifolds and summarizes key results needed for the computation of gradients, Hessians and geodesics. The focus lies on the use of these objects for optimization algorithms and on the complex Stiefel manifold as the key example.
 - Section 2.2 contains a background on the necessary concepts for the characterization of quantum dynamics, namely the quantum channel formalism, tensor networks and different quantum channel parametrizations (Section 2.2.1), as well as distance measures between states, between measurements and between channels (Section 2.2.2). We then review the main frameworks employed in quantum process tomography, highlighting the results of key works in Section 2.2.3.
 - In Section 2.2.4 we review gate set tomography, its formalism, challenges and state-of-the-art results in the literature.
 - Finally, in Section 2.3 we end the theoretical background with an explanation of shadow estimation, its central results and the concepts employed to derive them.
- Chapter 3 summarizes our results on compressive gate set tomography based largely on our work [45]. This includes our framework in terms of a tensor network model, in which we identify individual tensors as elements of the Stiefel manifold. We then explain the techniques employed in our reconstruction algorithm, and the results of numerical simulations on the compressive behavior of the model. We also touch on the improvements of compressive GST in comparison to previous state-of-the-art methods and explain the method behind our results on error mitigation for shadow estimation. Furthermore, we show unpublished recent results on the application of our algorithm [46] to a trapped ion experiment at the university of Siegen. This includes a discussion of gate set reports generated in the process (see also Appendix C).
 - In Chapter 4 we summarize our analytical results on the bias and variance of the shadow estimator under gate-dependent noise, based on our work in [47]. This includes an overview of how we derive upper bounds on the bias and variance through an intuitive observation about the structure of the shadow estimator in terms of average noise channels. We mention concrete examples that lead to a worst case bias, as well as additional results where we identify a class of noise channels whose effect can be efficiently mitigated.
 - Finally, in Chapter 5 we conclude by discussing the central results of this thesis and their implications.

Chapter 2

Theoretical background

In this Chapter we first give a short explanation of representation theory, matrix manifolds, and concepts from quantum information theory, such as quantum operations, measurements, noise models and frequently used distance measures. Representation theory is used throughout the quantum information literature and very prevalent in the subfield of characterization and benchmarking. Most notably, theoretical guarantees for randomized benchmarking [32] and shadow estimation [36] heavily rely on the calculation of moments over the unitary group. Optimization algorithms on matrix manifolds have only recently been proposed as a tool for quantum information tasks in Luchnikov *et al.* [48] around the same time our application to gate set tomography was formalized. The second part of this chapter provides the context for our work. In Section 2.2.3 and Section 2.2.4 we familiarize ourselves with a variety of approaches that have been proposed for the characterization of quantum dynamics. Section 2.3 concludes the chapter with an introduction to the theory of shadow estimation, where we explain the protocol, show how estimators are formed and summarized the sample complexity analysis.

2.1 Mathematical preliminaries

2.1.1 Notation

We denote a finite dimensional Hilbert space by $\mathcal{H} = \mathbb{C}^d$, where $d = q^n$ with the local dimension q and the number of subsystems n . We also write $[n] := \{1, 2, \dots, n\}$. The space of linear operators on $\mathcal{H}$ is written as $L(\mathcal{H})$ and the set of quantum states is defined as $\mathcal{S}(\mathcal{H}) := \{\rho \in L(\mathcal{H}) : \rho \succeq 0, \text{Tr}[\rho] = 1\}$. The space of linear maps $L(\mathcal{H}) \mapsto L(\mathcal{H})$, whose elements are commonly called *superoperators*, is denoted by $L(L(\mathcal{H}))$. We occasionally use Dirac notation with rounded brackets for elements of the vector space $L(\mathcal{H})$, whereby $|X\rangle \in L(\mathcal{H})$ and $\langle Y| \in L(\mathcal{H})^*$. For the inner product on $L(\mathcal{H})$ we use the standard Hilbert-Schmidt inner product $\langle Y|X\rangle = \text{Tr}(Y^\dagger X)$. We further write $\mathbb{R}^{n \times p}$ ($\mathbb{C}^{n \times p}$) for the space of real (complex) n by p matrices. The vectorization of a matrix is defined as the map $\text{vec} : \mathbb{C}^{n \times m} \mapsto \mathbb{C}^{nm} : X \mapsto \text{vec}(X)$, and we use the row major vectorization convention, i.e.

the rows of a matrix are stacked into a vector. The Pauli-matrices are defined as usual:

$$\sigma_X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \sigma_Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \sigma_Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (2.1)$$

and the n -qubit Pauli group is given by $\mathcal{P}_n = \langle \sigma_X, \sigma_Y, \sigma_Z \rangle^{\otimes n}$. The Pauli group up to a global phase is defined as $\mathbf{P}_n = \mathcal{P}_n / \langle i \rangle$, and we can index elements of $\mathbf{P}_n$ by $x, z \in \mathbb{F}_2^n$, such that $\sigma_{x,z} \propto \sigma_x^{z_1} \sigma_z^{z_1} \otimes \dots \otimes \sigma_x^{z_n} \sigma_z^{z_n}$. The subset of $[n]$ on which $\sigma_{x,z} = \sigma_{0,0}$ is called the support of $a = (x, z)$, and we write $\text{supp}(\sigma_a)$ and $\text{supp}(a)$ interchangeably.

The ℓ_p -norms on $\mathbb{C}^n$ are defined to be $\|x\|_{\ell_p} := (\sum_{i=1}^n |x_i|^p)^{1/p}$. We will also use the ℓ_p -norms on matrices via the identification $\|X\|_{\ell_p} := \|\text{vec}(X)\|_{\ell_p}$. Let $s(X)$ be the vector of singular values of a matrix $X \in \mathbb{C}^{n \times n}$, then the Schatten p -norms are defined as $\|X\|_p := \|s(X)\|_{\ell_p}$. We also

call $\|X\|_F = \|X\|_2 = \sqrt{\text{Tr}[X^\dagger X]}$ the Frobenius norm and $\|X\|_{\text{op}} = \|X\|_\infty = s_{\max}(X)$ the operator norm. The Schatten 1-norm $\|X\|_1$ is also called the *trace norm*.

2.1.2 Representation theory

For the theory of shadow estimation (Section 2.3) as well as for our own results on the stability of shadow estimation against noise (Chapter 4), mathematical tools from representation theory are essential in the derivation of estimators and sample complexity bounds. In the following, we give a brief background on representation theory with a focus on the use of Schur's Lemma to calculate averages over the unitary group. This section follows the standard textbook by Fulton and Harris [49], as well as the books by Simon [50] and Goodman and Wallach [51]. For an introduction to representation theory in the context of characterization and benchmarking see Kliesch and Roth [27].

Let G be a compact group and V be a complex vector space, then we call a continuous homomorphism $\omega : G \mapsto \text{GL}(V)$ a *representation* of G with respect to V . If a subspace $V_\lambda \subseteq V$ is left invariant under the action of ω , i.e.

$\omega(g)v = v$ for all $g \in G$ and $v \in V_\lambda$, then we call V_λ an invariant subspace. We write the restriction of ω on V_λ as $\sigma_\lambda := \omega|_{V_\lambda}$, where σ_λ is called a subrepresentation of ω . We further call a representation ω an *irreducible* representation (or short: irrep) if its only invariant subspaces are the trivial ones, V and $\{0\}$.

Given two representations $\omega : G \mapsto \text{GL}(V)$ and $\tilde{\omega} : G \mapsto \text{GL}(W)$, a linear map $\varphi : V \mapsto W$ that satisfies $\varphi\omega(g) = \tilde{\omega}(g)\varphi$ for all $g \in G$ is called a G -equivariant map or *intertwiner*. If in addition φ is invertible, then ω and $\tilde{\omega}$ are said to be *equivalent* and $\varphi\omega(g)\varphi^{-1} = \tilde{\omega}(g)$ for all $g \in G$. For two equivalent representations we write $\omega \simeq \tilde{\omega}$. Note that for unitary representations $\omega, \tilde{\omega} : G \mapsto U(d)$, equivalence implies that there exists a unitary U such that $U\omega U^\dagger = \tilde{\omega}$ (see [50], Theorem II.1.2). The number m_λ of equivalent subrepresentations to σ_λ is called the multiplicity of σ_λ in ω . The set of inequivalent irreducible subrepresentations of ω will be denoted by $\text{Irr}(\omega)$. We moreover define $\text{Comm}(\omega) := \{\varphi : \varphi\omega(g) = \omega(g)\varphi \ \forall g \in G\}$, the *commutant* of ω .

In the following, we will consider subgroups of $U(\mathbb{C}^n)$ and their representations on $V = \mathbb{C}^d$, such as the unitary group $U(d)$ or the Pauli group $\mathcal{P}_n$. Apart from $U(n)$ the perhaps most widely used subgroup of $\text{GL}(\mathbb{C}^n)$ in quantum information theory is the *Clifford group* Cl_n , which can be defined as $\text{Cl}_n = \langle \{H_i, S_i\}_{i \in [n]} \cup \{CNOT_{i,i+1}\}_{i \in [n-1]} \rangle$, where H_i, S_i are the usual Hadamard and phase gates on qubit i and $CNOT_{i,j}$ is the controlled-not gate with control i and target j .

For finite groups there is a clear intuitive notion of uniformly random selection, where the probability of selecting an element of a subset $S \subset G$ is given by $|S|/|G|$ and subsets of G are measured via the counting measure $\mu_c(S) = |S|$. For compact groups such as the unitary group, one can extend the notion of a uniform measure by requiring invariance under group action, i.e. $\mu(g) = \mu(gh) = \mu(hg)$ for all $g, h \in G$. It can be shown that there is a unique measure on $U(d)$, called the *Haar measure*, that is invariant under group action and satisfies basic finiteness and regularity conditions.

Of interest to us in the context of quantum information are group averages of the following form

$$\mathcal{T}(X) = \int \tilde{\omega}(g)^\dagger X \omega(g) d\mu(g). \quad (2.2)$$

For the case of $\tilde{\omega} = \omega$, these are known as *group twirls* and for $\tilde{\omega} \simeq \sigma_\lambda$ with σ_λ being an irrep of ω , they define a Fourier transform $\mathcal{T}[\lambda] = \int \sigma_\lambda^\dagger(g)(\cdot)\omega(g)d\mu(g)$ with $\mathcal{T} : \text{Irr}(G) \mapsto \mathcal{L}(\mathcal{H} \otimes \mathcal{H}^*)$. The properties of the Fourier transform are for instance used in the theory of randomized benchmarking (RB) [31, 32, 52–54]. To compute these group averages, we use the following ubiquitous result of representation theory.

Theorem 1 (Schur). *Given two irreducible representations $\sigma : G \mapsto \text{GL}(V)$, $\tilde{\sigma} : G \mapsto \text{GL}(W)$ and an intertwining map $\varphi : V \mapsto W$ such that $\varphi\sigma = \tilde{\sigma}\varphi$, then*

- (i) φ is an isomorphism if $\sigma \simeq \tilde{\sigma}$,

(ii) $\varphi = c\mathbb{1}$ for some $c \in \mathbb{C}$ if $\sigma = \tilde{\sigma}$ and

(iii) $\varphi = 0$ else.

Proof. Let us assume there exists a $v \in V$ such that $\varphi v = 0$, then from $\varphi\sigma = \tilde{\sigma}\varphi$ we get $\varphi\sigma v = 0$, meaning that the kernel of φ is an invariant subspace of σ and therefore either $\text{Ker}(\varphi) = \{0\}$ or $\text{Ker}(\varphi) = V$. We can make the analogous argument for the cokernel, that is, if there exists a $w \in W$ such that $w\varphi = 0$, then $w\tilde{\sigma}\varphi = 0$ and $\text{Coker}(\varphi) = \{0\}$ or $\text{Coker}(\varphi) = W$. If either $\text{Ker}(\varphi)$ or $\text{Coker}(\varphi)$ are nonzero, then $\varphi = 0$ (case (iii)), otherwise ϕ is a bijection and thus an isomorphism. If φ is an isomorphism, then per definition $\sigma \simeq \tilde{\sigma}$. For case (ii) we can specify the isomorphism up to a constant by noting that if $\varphi \neq 0$, it has an eigenvalue $c \neq 0$ and it further holds that $(\varphi - c\mathbb{1})\sigma = \sigma(\varphi - c\mathbb{1})$. By the previous argument we again know that $\varphi - c\mathbb{1}$ is either equal to zero or an isomorphism. We can exclude the latter case, since $\text{Ker}(\varphi - c\mathbb{1}) \neq \{0\}$ by assumption. Thus, $\varphi = c\mathbb{1}$. $\square$

Corollary 2. *If in addition σ and $\tilde{\sigma}$ are unitary representations, then φ in case (i) is itself proportional to a unitary transformation and φ is unique up to a scalar.*

Proof. For unitary representations $\varphi\sigma = \tilde{\sigma}\varphi$ implies $\sigma\varphi^\dagger = \varphi^\dagger\tilde{\sigma}$ and hence $\varphi^\dagger\varphi\sigma = \sigma\varphi^\dagger\varphi$. Now condition (ii) in Theorem 1 implies that $\varphi^\dagger\varphi = c\mathbb{1}$ and therefore $\varphi = \sqrt{c}U$ for some unitary U . If we assume there exists a second intertwiner $\tilde{\varphi}$ satisfying $\tilde{\varphi}\sigma = \tilde{\sigma}\tilde{\varphi}$, then we can combine this condition with $\sigma\varphi^\dagger = \varphi^\dagger\tilde{\sigma}$ to get $\varphi^\dagger\tilde{\varphi}\sigma = \sigma\varphi^\dagger\tilde{\varphi}$ and thus $\tilde{\varphi} \propto \varphi$. $\square$

We can now immediately recover the familiar results for unitary twirls on quantum states over $\mathbb{C}^d$ given as

$$\mathcal{T}(\rho) = \int U\rho U^\dagger d\mu(U). \quad (2.3)$$

First, note that $\mathcal{T}(\rho)$ commutes with all elements of $U(d)$, since the Haar measure is invariant under unitaries and therefore $\tilde{U} \int U\rho U^\dagger d\mu(U) \tilde{U}^\dagger = \int U\rho U^\dagger d\mu(U)$. Since the standard unitary representation is irreducible, $\mathcal{T}(\rho)$ is an intertwiner and per the second statement in Theorem 1, we get that $\mathcal{T}(\rho) = \alpha\mathbb{1}$. The constant α can be easily determined by noting that $\alpha d = \text{Tr}[\alpha\mathbb{1}] = \text{Tr}[\int U\rho U^\dagger d\mu(U)] = \text{Tr}[\rho]$. The same holds true for the Pauli group, since its standard unitary representation is also irreducible, and we find that

$$\frac{1}{|\mathcal{P}_n|} \sum_{U \in \mathcal{P}_n} U\rho U^\dagger = \frac{\text{Tr}[\rho]}{d} \mathbb{1}. \quad (2.4)$$

To compute more complicated twirls with unitary representations that are not irreducible, we turn to the following standard result.

Lemma 3. *Let $\omega : G \mapsto U(\mathbb{C}^d)$ be a unitary representation, then there exists a set S of irreducible representations $\sigma_\lambda : G \mapsto U(V_\lambda)$ on orthogonal subspaces $V_\lambda \subseteq \mathbb{C}^d$ such that ω can be decomposed as*

$$\omega \simeq \bigoplus_{\lambda \in S} \sigma_\lambda \quad \text{and} \quad \mathbb{C}^d \simeq \bigoplus_{\lambda \in S} V_\lambda. \quad (2.5)$$

Proof. This result can be easily verified by noting that if there exists an invariant subspace V of ω , then its orthogonal complement $V^\perp$ is also an invariant subspace: For all $v \in V, \tilde{v} \in V^\perp$ and $g \in G$, we have $0 = \langle \omega(g)v, \tilde{v} \rangle = \langle v, \omega^\dagger \tilde{v} \rangle$. We can thus write $\mathbb{C}^d \simeq V \oplus V^\perp$ and $\omega \simeq \omega|_V \oplus \omega|_{V^\perp}$. Since $\omega|_{V^\perp}$ is itself a unitary representation (but not necessarily an irrep), we can find again an invariant subspace and its complement in $V^\perp$ and the result follows by iterating until all subrepresentations are irreducible. $\square$

We now turn to a simplified version of the twirl defined in Eq. (2.2), where we deal with irreducible representations: $\mathcal{T}(X) = \int \sigma(g)^\dagger X \tau(g) d\mu(g)$. Since the Haar measure is again invariant under left or right multiplication, $\sigma^\dagger(\tilde{g})\mathcal{T}(X)\tau(\tilde{g}) = \mathcal{T}(X)$ for all $\tilde{g} \in G$ and $X : V_\tau \mapsto V_\sigma$,

meaning $\mathcal{T}(X)$ are intertwiners. Schur's Lemma now tells us that $\mathcal{T}(X)$ is only nonzero if σ and τ are either equal or equivalent.

Let us first consider the general case where σ and τ are equivalent. Corollary 2 then implies that $\mathcal{T}(Y) = c(Y)U$ with a unique U that determines the equivalence $\sigma = U\tau U^\dagger$. It is straightforward to see that $c(U) = 1$, since

$$1 = \int \sigma^\dagger(g)\sigma(g)d\mu(g) = \int \sigma^\dagger(g)U\tau(g)U^\dagger d\mu(g) = c(U)UU^\dagger.$$

We can then see that the map $\mathcal{T}$ is a projection, since $\mathcal{T}(\mathcal{T}(Y)) = c(Y)\mathcal{T}(U) = c(Y)c(U)U = c(Y)U = \mathcal{T}(Y)$. Let now $d = \dim(V^\omega) = \dim(V^{\tilde{\omega}})$. We further know that

$$c(Y)d = \text{Tr}[\mathcal{T}(Y)U^\dagger] = \text{Tr}\left[\int \sigma^\dagger(g)Y\tau(g)U^\dagger d\mu(g)\right] = \text{Tr}\left[Y\int \tau(g)U^\dagger\sigma(g)^\dagger d\mu(g)\right] = \text{Tr}[YU^\dagger],$$

where the last equality follows from the fact that $\int \tau(g)U^\dagger\sigma(g)^\dagger d\mu(g) = \mathcal{T}(U)^\dagger = U^\dagger$. Thus we can write the action of $\mathcal{T}$ as $\mathcal{T}(Y) = \frac{\text{Tr}[YU^\dagger]}{d}U = \frac{1}{d}|U\rangle\langle U|Y$. If σ and τ are equal, we get $\mathcal{T}(Y) = \frac{1}{d}|\mathbb{1}_d\rangle\langle\mathbb{1}_d|Y$.

With the understanding of a twirl under irreps in place, we will now generalize it to twirls with arbitrary unitary representations. To this end, we first identify the set of orthogonal invariant subspaces $\{V_\lambda^\omega\}$ of a given unitary representation ω with a set of orthogonal projectors $\Pi_\lambda^\omega : \mathbb{C}^d \mapsto \mathbb{C}^d$, such that $\Pi_\lambda^\omega V_\lambda = V_\lambda$ and $\Pi_\lambda^\omega V_{\lambda'} = 0$ for $\lambda \neq \lambda'$. Let us now assume we are given two representations $\omega = \bigoplus_{\lambda \in S} \sigma_\lambda$ and $\tilde{\omega} = \bigoplus_{\lambda' \in \tilde{S}} \tau_{\lambda'}$. Moreover, for $Y : V_{\lambda'}^{\tilde{\omega}} \mapsto V_\lambda^\omega$ we define $\mathcal{T}_{\lambda,\lambda'}(Y) = \int \sigma_\lambda(g)^\dagger Y \tau_{\lambda'}(g) d\mu(g)$. Lemma 3 allows us to write the general twirl of Eq. (2.2) as

$$\int \omega(g)^\dagger X \tilde{\omega}(g) d\mu(g) = \sum_{\lambda \in S, \lambda' \in \tilde{S}} \int \sigma_\lambda^\dagger \Pi_\lambda^\omega X \Pi_{\lambda'}^{\tilde{\omega}} \tau_{\lambda'} d\mu(g) = \sum_{\lambda \in S, \lambda' \in \tilde{S}} \mathcal{T}_{\lambda,\lambda'}(\Pi_\lambda^\omega X \Pi_{\lambda'}^{\tilde{\omega}}).$$

For each pair λ, λ' satisfying $\sigma_\lambda \simeq \tau_{\lambda'}$ we now have that $\mathcal{T}_{\lambda,\lambda'}(\Pi_\lambda^\omega X \Pi_{\lambda'}^{\tilde{\omega}})$ is an intertwiner, meaning there is a $U_{\lambda,\lambda'} \in U(d_\lambda)$ such that $\mathcal{T}_{\lambda,\lambda'}(Y) = \frac{1}{d_\lambda}|U_{\lambda,\lambda'}\rangle\langle U_{\lambda,\lambda'}|Y$. In summary, the general twirl takes the form

$$\int \omega(g)^\dagger X \tilde{\omega}(g) d\mu(g) = \sum_{\lambda \in S, \lambda' \in \tilde{S} : \sigma_\lambda \simeq \tau_{\lambda'}} \frac{1}{d_\lambda} |U_{\lambda,\lambda'}\rangle\langle U_{\lambda,\lambda'}| \Pi_\lambda^\omega X \Pi_{\lambda'}^{\tilde{\omega}}. \quad (2.6)$$

For the simpler case where $\omega = \tilde{\omega}$ and all irreps are multiplicity-free, we have that $U_{\lambda,\lambda} = \mathbb{1}_{V_\lambda} = \Pi_\lambda^\omega$ and the familiar result

$$\int \omega(g)^\dagger X \omega(g) d\mu(g) = \sum_{\lambda \in \text{Irr}(\omega)} \frac{1}{d_\lambda} |\Pi_\lambda^\omega\rangle\langle \Pi_\lambda^\omega| X \quad (2.7)$$

follows. The task of computing twirls then reduces to finding bases for all invariant subspaces of the given representations $\omega, \tilde{\omega}$. These bases then define the projectors Π_λ^ω and $\Pi_{\lambda'}^{\tilde{\omega}}$, as well as any potential basis change $U_{\lambda,\lambda'}$.

A type of twirl that is of particular importance in quantum information theory is the twirl by the product representation $\omega_k : U(\mathbb{C}^d) \mapsto (U(\mathbb{C}^d))^{\otimes k}$:

$$\mathcal{T}_k(X) = \int \omega_k^\dagger(g) X \omega_k(g) d\mu(g). \quad (2.8)$$

It is also known as the *kth moment operator* of the Haar measure, since one can write the Haar-average of any order k polynomial in g and g^* as $\int \text{Tr}[Y \omega_k^\dagger(g) X \omega_k(g)] d\mu(g)$ for suitable matrices X, Y and $g \in U(\mathbb{C}^d)$. To gain a basic understanding of the methods that are generally used to compute these twirls we will now give a brief review of them, largely based on the treatment of Roberts and Yoshida [55].

Rather than determining all irreps and invariant subspaces of ω_k by hand, one can use a duality relation between product representations and the following representation of the symmetric group: $\pi_k : \mathcal{S}_k \mapsto (U(\mathbb{C}^d))^{\otimes k}$, which is defined by

$$\pi_k(g)|\psi_1\rangle \otimes \cdots \otimes |\psi_k\rangle = |\psi_{g(1)}\rangle \otimes \cdots \otimes |\psi_{g(k)}\rangle. \quad (2.9)$$

One can immediately see that π_k and ω_k commute, and in fact the following also holds.

Lemma 4 (Schur-Weyl duality).

$$\begin{aligned} \text{Comm}(\omega_k) &= \text{Span}(\{\pi_k(g)\}_{g \in \mathcal{S}_k}), \\ \text{Comm}(\pi_k) &= \text{Span}(\{\omega_k(g)\}_{g \in U(d)}). \end{aligned} \quad (2.10)$$

For the proof see for instance Theorem 3.3.8 in the book by Goodman and Wallach [51]. Lemma 4 implies that we can write

$$\mathcal{T}_k(X) = \sum_{g \in \mathcal{S}_k} \pi_k(g) c_g(X), \quad (2.11)$$

where c_k are linear maps since $\mathcal{T}_k$ is a linear map, and therefore $c_g(X) = \text{Tr}[C_g X]$ for some matrix C_g . From the left- and right- invariance of the Haar measure it again follows that $\mathcal{T}_k(X) = \mathcal{T}_k(\omega_k^\dagger(g) X \omega_k(g))$ and thus also $\sum_{g \in \mathcal{S}_k} \pi_k c_g(X) = \sum_{g \in \mathcal{S}_k} \pi_k c_g(\omega_k(g)^\dagger X \omega_k(g))$ and $\text{Tr}[C_{g'} X] = \text{Tr}[C_{g'} \omega_k(g)^\dagger X \omega_k(g)]$ for all $X \in \mathcal{L}((\mathbb{C}^d)^{\otimes k})$ and all $g \in U(\mathbb{C}^d)$. Thus $C_{g'} \in \text{Comm}(\omega_k)$, and we can again express them in terms of $\{\pi_k(g)\}$. In summary, we have

$$\mathcal{T}_k(X) = \sum_{g, g' \in \mathcal{S}_k} \pi_k(g) W_{g, g'} \text{Tr}[\pi_k(g') X]. \quad (2.12)$$

In order to determine the coefficient matrix $W_{g, g'}$ (known as the Weingarten matrix) we first observe that $\mathcal{T}_k(\pi_k(g)) = \pi_k(g)$, and we also use the known identity $Q_{g, g'} := \text{Tr}[\pi_k(g) \pi_k(g')] = d^{\#\text{cycles}(g \cdot g')}$. We then arrive at $\pi_k(g'') = \mathcal{T}(\pi_k(g'')) = \sum_{g, g' \in \mathcal{S}_k} \pi_k(g) W_{g, g'} \text{Tr}[\pi_k(g') \pi_k(g'')]$. Multiplying from the left with $\pi_k(h)$ for $h \in G$ and taking the trace, we are left with the matrix equation $Q = QWQ$ and assuming $k \leq d$ (which ensures Q is invertible), we get $W = Q^{-1}$.

We now have a recipe for calculating the k th moment operator for the Haar measure, in terms of the Weingarten matrix and the permutation operators $\pi_k(g)$. An explicit construction for $k = 2$ and a related solution in terms of projectors onto irreducible subspaces akin to 2.7 can be found in Kliesch and Roth [27].

The k th moment operator for the Haar measure can also be realized without taking the Haar average over the full unitary group. If we are given an arbitrary distribution ν on $U(\mathbb{C}^d)$, we say ν is a k -design if its k th moment operator coincides with the k th moment operator of the Haar measure, i.e.

$$\int \omega_k^\dagger(g) X \omega_k(g) d\nu(g) = \int \omega_k^\dagger(g) X \omega_k(g) d\mu(g). \quad (2.13)$$

It can easily be seen that a k -design is also a $(k-1)$ -design, since for $X \in \mathcal{L}((\mathbb{C}^d)^{\otimes(k-1)})$, $\text{Tr}_k[\omega_k^\dagger(g) X \otimes \mathbb{1} \omega_k(g)] = \omega_{k-1}^\dagger(g) X \omega_{k-1}(g)$. An important example for our purposes is the uniform distribution over the Clifford group, for which we know that it forms a unitary 3-design but not a unitary 4-design [56, 57].

2.1.3 Optimization on matrix manifolds

Differential geometry, or the study of smooth manifolds, is a fundamental mathematical branch which has a wide range of uses in physics, most notably in general relativity and electromagnetism. Another area where manifolds, and especially matrix manifolds, play a role is constraint

optimization. In this scenario the constraint set can be cast as an embedded (Riemannian) submanifold of $\mathbb{R}^{n \times m}$ or $\mathbb{C}^{n \times m}$, resulting in optimization algorithms that produce iterates on the manifold. One refers to this task as *Riemannian optimization*, which has largely risen to prominence in the 90s and early 2000s [58–60] and is still very active [61–63]. It has also been observed that many optimization problems in quantum information theory can be formulated in the language of Riemannian optimization [48]. In this section we give a condensed introduction to Riemannian optimization on matrix manifolds. This will provide the necessary background to understand the optimization algorithm developed in Chapter 3, where we treat the set of quantum channels as a Riemannian manifold. The section mostly follows the standard textbooks by Absil *et al.* [59] and Boumal [64].

Manifolds and embedded submanifolds

The intuitive idea of a manifold $\mathcal{M}$ is that of a topological space which can be locally identified with an open subset of the familiar and easy to handle Euclidean space $\mathbb{R}^n$. For $\mathcal{M}$ to be called a manifold, restrictions are placed on the topology of $\mathcal{M}$, namely that it is Hausdorff and second-countable. The Hausdorff condition ensures that a convergent sequence on $\mathcal{M}$ has a unique limit point, a property which is very desirable for optimization problems defined on $\mathcal{M}$. Given a subset $\mathcal{U} \in \mathcal{M}$, the local correspondence between $\mathcal{U}$ and $\mathbb{R}^n$ is defined via *charts*—bijective functions $\varphi : \mathcal{U} \mapsto \mathbb{R}^n$. For a point $X \in \mathcal{U}$, the corresponding element $\varphi(X) \in \mathbb{R}^n$ is then called the *coordinate* of X . To cover the whole set $\mathcal{M}$ with charts in a meaningful way, we need what is called an *atlas*, i.e.

a set of charts and domains $\mathcal{A} := \{(\varphi_i, \mathcal{U}_i)\}$ such that $\bigcup_i \mathcal{U}_i = \mathcal{M}$. If $\mathcal{M}$ admits an atlas such that for all overlapping domains $\mathcal{U}_i, \mathcal{U}_j$ it holds that the coordinate change $\varphi_i \circ \varphi_j^{-1} : \mathbb{R}^n \mapsto \mathbb{R}^n$ is a smooth (C^∞) function, we call $\mathcal{M}$ an n -dimensional differentiable manifold. Note that in principle a manifold is given by the pair $(\mathcal{M}, \mathcal{A})$ of a set and an a corresponding atlas, but since we typically do not need to distinguish between atlases, we just omit $\mathcal{A}$ from the notation.

We will also require the notion of a product manifold $\mathcal{M}_1 \times \mathcal{M}_2$, elements of which are the pairs (X_1, X_2) with $X_1 \in \mathcal{M}_1$ and $X_2 \in \mathcal{M}_2$. An atlas for $\mathcal{M}_1 \times \mathcal{M}_2$ can then be readily obtained by combining charts of the individual atlases via $\varphi_1 \times \varphi_2$ (X_1, X_2) = $(\varphi_1(X_1), \varphi_2(X_2))$.

Functions between manifolds are analyzed via charts in the following way. Consider $f : \mathcal{M}_1 \mapsto \mathcal{M}_2$, where $\mathcal{M}_1$ and $\mathcal{M}_2$ have dimension n_1 and n_2 respectively. By taking charts at points $x \in \mathcal{M}_1$ and $f(x) \in \mathcal{M}_2$, we can treat f locally around x via its coordinate representation $\tilde{f} := \varphi_2 \circ f \circ \varphi_1^{-1} : \mathbb{R}^{n_1} \mapsto \mathbb{R}^{n_2}$. The function f is then called differentiable or smooth if $\tilde{f}$ is C^∞ at each point. Functions between manifolds that are of importance to our problem are the cost functions in optimization problems, i.e.

$f : \mathcal{M} \mapsto \mathbb{R}$, where $\mathbb{R}$ is interpreted as a 1-dimensional manifold with trivial atlas.

One can also define *complex* manifolds, where charts are defined via $\varphi : \mathcal{M} \mapsto \mathbb{C}^n$ and coordinate changes of differentiable complex manifolds $\varphi_i \circ \varphi_j^{-1}$ need to be complex differentiable (holomorphic). The problem of complex manifolds for optimization problems is that non-constant cost functions defined via their coordinate representations $\tilde{f} : \mathbb{C}^n \mapsto \mathbb{R}$ can not be holomorphic. This can be directly seen via the Cauchy-Riemann equations, which for a real image imply that the derivatives with respect to the real and imaginary parts of each coordinate vanish and therefore the function is constant. The implications for optimization problems are that derivatives are not well-defined, and, since non-holomorphic functions are not analytic as well, neither is the Taylor series around points in the domain.

The solution is to treat n -dimensional complex manifolds as $2n$ -dimensional real manifolds, where the complex coordinates are split into real and imaginary parts. The additional structure due to the imaginary unit i is then added via the linear map $I : \mathbb{R}^{2n} \mapsto \mathbb{R}^{2n}$ with $I^2 = -\mathbf{1}$. A real manifold with this additional structure is also called an almost complex manifold, where every complex manifold is also an almost complex manifold, but not vice versa.

In the following we make the treatment of complex manifolds in terms of real manifolds more

explicit by defining the notion of a submanifold and looking at the complex Stiefel manifold as an example. But first we need to provide a few definitions. The (Fréchet) *differential* of a function $g : \mathbb{R}^n \mapsto \mathbb{R}^m$ at position x is the linear operator $Dg(x) : \mathbb{R}^n \mapsto \mathbb{R}^m : h \mapsto Dg(x)[h]$ that satisfies

$$\lim_{h \rightarrow 0} \frac{\|g(x+h) - g(x) - Dg(x)[h]\|}{\|h\|} = 0. \quad (2.14)$$

For the more familiar case $m = 1$, $Dg(x)$ is the gradient of g at x and $Dg(x)[h]$ is the directional derivative in direction h . Furthermore, the *rank* of g at x is defined to be the dimension of the range of $Dg(x)$. If the rank of $Dg(x) : \mathbb{R}^n \mapsto \mathbb{R}^m$ is $n \leq m$ for all $x \in \mathbb{R}^n$, then g is called an *immersion*. Reciprocally, if the rank of $Dg(x)$ is $m \leq n$ for all x , then g is called a *submersion*. These notions are readily applied to functions $f : \mathcal{M}_1 \mapsto \mathcal{M}_2$ by considering the differential of the coordinate representation $\tilde{f}$. A simple example of an immersion is the canonical immersion $(x_1, \dots, x_n) \mapsto (x_1, \dots, x_n, 0, \dots, 0)$.

Given two manifolds $\mathcal{M}_1$ and $\mathcal{M}_2$, where $\mathcal{M}_1 \subset \mathcal{M}_2$ as sets, we call $\mathcal{M}_1$ an *embedded submanifold* of $\mathcal{M}_2$ if there exists an immersion $f : X \in \mathcal{M}_1 \mapsto X \in \mathcal{M}_2$ and the topology of $\mathcal{M}_1$ is equal to the subspace topology induced by $\mathcal{M}_2$. In this context, $\mathcal{M}_2$ is also called the ambient space of $\mathcal{M}_1$. The following result given in Absil *et al.* [59] (Proposition 3.3.3) provides a practical condition to prove that a given manifold is an embedded submanifold:

Theorem 5 (Submersion Theorem). *Let $\mathcal{M}_1$ and $\mathcal{M}_2$ be two manifolds whose dimensions satisfy $d_1 > d_2$ and let $f : \mathcal{M}_1 \mapsto \mathcal{M}_2$ be a smooth function. If for a given point $x \in \mathcal{M}_2$ the rank of f is equal to d_2 on the whole preimage $f^{-1}(x)$, then $f^{-1}(x)$ is a closed embedded submanifold of $\mathcal{M}_1$ and has dimension $d_1 - d_2$.*

With this result at hand, we can now turn to an example.

The complex Stiefel manifold

The complex Stiefel manifold can be defined via the set

$$\text{St}(n, p) := \{X \in \mathbb{C}^{n \times p} : X^\dagger X = \mathbb{1}_p\}. \quad (2.15)$$

In the real case $X \in \mathbb{R}^{n \times p}$ the hermitian conjugate is replaced by the transpose. Stiefel manifolds turn up in optimization problems of several fields, such as signal processing, quantum chemistry and machine learning [65–69]. In Absil *et al.* [59] it was shown via Theorem 5 that the real Stiefel manifold is an embedded submanifold of the standard matrix manifold $\mathbb{R}^{n \times p}$. We will now show in an analogous manner how the complex Stiefel manifold is an embedded submanifold of $\mathbb{R}^{2n \times p}$. For this we first define a basis on $\mathbb{R}^{2n \times p}$ via the standard unit matrices $\{E_{ij}, iE_{ij}\}_{i \in [n], j \in [p]}$, which are orthonormal with respect to the inner product

$$\langle X, Y \rangle := \text{Re Tr}[X^\dagger Y]. \quad (2.16)$$

Any matrix in $\mathbb{C}^{n \times p}$ can be represented in this basis with real coefficients, thus making it a real vector space equivalent to taking the real and imaginary part in the standard identification $\mathbb{C}^n \simeq \mathbb{R}^{2n}$.

If we write $X \in \mathbb{C}^{n \times p} = A + iB$ with $A, B \in \mathbb{R}^{n \times p}$ we get $\mathbb{1} = X^\dagger X = A^T A - iB^T A + iA^T B + B^T B$. Let now $\text{Sym}(p)$ be the set of real symmetric $p \times p$ matrices and $\text{Asym}(p)$ be the set of real antisymmetric $p \times p$ matrices. The (smooth) function $f : \mathbb{R}^{2n \times p} \mapsto \text{Sym}(p) \oplus \text{Asym}(p)$:

$$f(A, B) = \begin{pmatrix} A^T A + B^T B - \mathbb{1} \\ A^T B - B^T A \end{pmatrix} \quad (2.17)$$

now encodes the Stiefel constraints as the (complex) Stiefel manifold is given by pairs of matrices (A, B) that satisfy $f(A, B) = 0$, such that $\text{St}(n, p) = f^{-1}(0)$. According to Theorem 5, what we

now have to check is that f is of full rank on $f^{-1}(0)$. For this we take the differential using the rules of matrix differentiation to find that

$$Df(A, B)[Z_1, Z_2] = \begin{pmatrix} A^T Z_1 + Z_1^T A + B^T Z_2 + Z_2^T B \\ Z_1^T B + A^T Z_2 - (Z_2^T A + B^T Z_1) \end{pmatrix}. \quad (2.18)$$

If for all (A, B) and all $\tilde{Z}_1 \in \text{Sym}(p)$, $\tilde{Z}_2 \in \text{Asym}(p)$ we can find a pair $Z_1, Z_2 \in \mathbb{R}^{2n \times p}$ such that $Df(A, B)[Z_1, Z_2] = (\tilde{Z}_1, \tilde{Z}_2)$, we know that rank of the differential is equal to the dimension of $\text{Sym}(p) \oplus \text{Asym}(p)$. It can then be straightforwardly verified using $A^T A + B^T B = \mathbb{1}$ and $A^T B - B^T A = 0$, that setting $Z_1 = \frac{1}{2}(A\tilde{Z}_1 + B\tilde{Z}_2)$ and $Z_2 = \frac{1}{2}(B\tilde{Z}_1 - A\tilde{Z}_2)$ satisfies this condition and hence $\text{St}(n, p)$ is an embedded submanifold of $\mathbb{R}^{2n \times p}$.

The real dimension of $\text{St}(n, p)$ can also be computed using Theorem 5 as $\dim(\text{St}(n, p)) = \dim(\mathbb{R}^{2n \times p}) - \dim(\text{Sym}(p) \oplus \text{Asym}(p)) = 2np - (p(p+1)/2 + p(p-1)/2) = 2np - p^2$. Additional properties of the complex Stiefel manifold, such as the tangent spaces, the natural metric and geodesics with respect to this metric are derived throughout this section and Appendix A. For now we will first define these terms and provide the background for their derivation.

Tangent spaces

The concept of a tangent vector to a curve $\gamma : \mathbb{R} \mapsto \mathcal{M} : t \mapsto \gamma(t)$ is well-defined for an embedded submanifold of a vectors space, since we can add elements of $\mathcal{M}$ by identifying them with elements in the parent vector space:

$$\dot{\gamma}(0) := \lim_{t \rightarrow 0} \frac{\gamma(t) - \gamma(0)}{t}, \quad (2.19)$$

provided the limit exists. When the vector space structure is not available, one can generalize the concept of a tangent vector as a linear map from the set of smooth functions $\mathcal{F}_X := \{f : \mathcal{U}_X \mapsto \mathbb{R}\}$, where $X \in \mathcal{U} \subset \mathcal{M}$, to the real numbers $\mathbb{R}$. Let $\gamma(t)$ again be a curve on $\mathcal{M}$ satisfying $\gamma(0) = X$. The tangent vector realized by the curve $\gamma(t)$ is then the linear map $\dot{\gamma} : \mathcal{F}_X \mapsto \mathbb{R}$ with

$$\dot{\gamma}f := \left. \frac{df(\gamma(t))}{dt} \right|_{t=0}. \quad (2.20)$$

This general definition can readily be reconciled with Eq. (2.19) by noting that whenever $\mathcal{M}$ is a submanifold of a vector space, we get

$$\dot{\gamma}f = D\hat{f}(X)[\dot{X}_0], \quad (2.21)$$

where $\hat{f}$ is the extension of f on the ambient space. The interpretation of the tangent 'vector' $\dot{X}$ is thus that it is a direction determined by the curve $\gamma(t)$, either given as vector (Eq. (2.19)), or a map (Eq. (2.20)).

The *tangent space* $T_X \mathcal{M}$ of the manifold $\mathcal{M}$ at position X is defined to be the set of all tangent vectors $\dot{\gamma}$ realized by curves $\gamma(t)$ with $\gamma(0) = X$. A very useful property of the tangent space is that it is a vector space. This allows us to stretch, add and transform elements of the tangent space, which we can then map back to the manifold in the context of iterative optimization algorithms.

If the tangent space for a manifold $\mathcal{M}$ at position X is defined via the set $f^{-1}(0)$ of a constant rank function such as in Eq. (2.17), it can be determined in a simple way. Note that since the resulting manifold is an embedded submanifold of a vector space, Eq. (2.21) and Eq. (2.20) imply $Df(X)[\dot{\gamma}(0)] = \left. \frac{df(\gamma(t))}{dt} \right|_{t=0}$. For any curve $\gamma(t)$ on $\mathcal{M}$ we have that $f(\gamma(t)) = 0$ and hence $Df(X)[\dot{\gamma}(0)] = 0$. Therefore $\dot{\gamma}(0) \in \text{Ker}(Df(X))$ and via parameter counting arguments it can be further shown that indeed $T_X \mathcal{M} = \text{Ker}(Df(X))$ (see Section 3.5.7 in [59]).

If $T_X\mathcal{M}$ is a vector space over either the real or complex numbers, we can define an inner product $\langle \cdot, \cdot \rangle : T_X\mathcal{M} \times T_X\mathcal{M} \mapsto \mathbb{R}$. The natural choice is the restriction of the inner product defined on an ambient vector space $\bar{\mathcal{M}}$ to the tangent space:

$$\langle \Delta, \Omega \rangle_{T_X\mathcal{M}} := \langle \bar{\Delta}, \bar{\Omega} \rangle_{\bar{\mathcal{M}}}. \quad (2.22)$$

where $\Delta, \Omega \in T_X\mathcal{M}$ and $\bar{\Delta}, \bar{\Omega}$ denote their respective inclusions in $\bar{\mathcal{M}}$. The inner product is only defined for a fixed $T_X\mathcal{M}$ and thus depends on the position X , we will however drop the subscript $T_X\mathcal{M}$ since the position is usually clear from the context.

Within the ambient space $\bar{\mathcal{M}}$ one can also define the *normal space* $N_X\mathcal{M}$ of $\mathcal{M}$ at X , which is the orthogonal complement of $T_X\mathcal{M}$:

$$N_X\mathcal{M} = \{\Delta_\perp \in \bar{\mathcal{M}} : \langle \bar{\Delta}, \Delta_\perp \rangle = 0 \ \forall \Delta \in T_X\mathcal{M}\} \quad (2.23)$$

One can further write any element of $T_X\bar{\mathcal{M}}$ as $\bar{\Delta} = \Delta + \Delta_\perp$ for some $\Delta \in T_X\mathcal{M}$ and some $\Delta_\perp \in N_X\mathcal{M}$. If $\mathcal{M}$ is a vectors space, it holds that $T_X\mathcal{M} \simeq \mathcal{M}$ and in that case every element of $\bar{\mathcal{M}}$ itself can be written via elements of $T_X\mathcal{M}$ and $N_X\mathcal{M}$. From now on we will use X for both $X \in \mathcal{M}$ and its inclusion $\bar{X} \in \bar{\mathcal{M}}$.

The inner product canonically induces a norm $\|\Delta\| := \sqrt{\langle \Delta, \Delta \rangle}$ on $T_X\mathcal{M}$. If the inner product is a smooth function, then the pair $(\mathcal{M}, \langle \cdot, \cdot \rangle)$ is called a *Riemannian manifold* and $\langle \cdot, \cdot \rangle$ is called a *Riemannian metric*. The Riemannian metric with the induced norm gives us a tool to measure distances on the tangent space, but we would also like to measure distances on the manifold. One can generalize the standard way of measuring the length of a curve on Euclidean space to a curve $\gamma : [a, b] \mapsto \mathcal{M}$ on the manifold via

$$l(\gamma) := \int_a^b \sqrt{\langle \dot{\gamma}(t), \dot{\gamma}(t) \rangle} dt, \quad (2.24)$$

since $\langle \dot{\gamma}(t), \dot{\gamma}(t) \rangle$ is well-defined. This induces a distance between points $X, Y \in \mathcal{M}$ via $\text{dist}(X, Y) = \inf_{\gamma} l(\gamma)$, where $\gamma(t)$ are curves that satisfy $\gamma(a) = X$ and $\gamma(b) = Y$.

In analogy to the gradient defined for smooth functions over vector spaces, we can define the Riemannian gradient as a generalization to manifolds. Given again a smooth function $f : \mathcal{M} \mapsto \mathbb{R}$, the *Riemannian gradient* of f at position X is defined to be the unique tangent space element $G \in T_X\mathcal{M}$ that satisfies

$$\langle G, \Delta \rangle = Df(X)[\Delta] \quad \forall \Delta \in T_X\mathcal{M}. \quad (2.25)$$

For the purpose of optimization, G gives the tangent space direction of fastest increase of the function at position X , since

$$G/\|G\| = \underset{\Delta \in T_X\mathcal{M} : \|\Delta\|=1}{\operatorname{argmax}} Df(X)[\Delta]. \quad (2.26)$$

For the Stiefel manifold we can use the identification $T_X\mathcal{M} = \text{Ker}(Df(X))$ together with Eq. (2.17) and Eq. (2.18) to determine its tangent space. The condition $Df(X)[\Delta] = 0 \ \forall \Delta \in T_X\mathcal{M}$ then translates into

$$\begin{pmatrix} A^T Z_1 + Z_1^T A + B^T Z_2 + Z_2^T B \\ Z_1^T B + A^T Z_2 - (Z_2^T A + B^T Z_1) \end{pmatrix} = 0. \quad (2.27)$$

Since the characterization is much simpler in terms of complex matrices, we rewrite the condition as

$$A^T Z_1 + Z_1^T A + B^T Z_2 + Z_2^T B + i(Z_1^T B + A^T Z_2 - (Z_2^T A + B^T Z_1)) = 0. \quad (2.28)$$

Setting $X = A + iB$ and $\Delta = Z_1 + iZ_2$, a short calculation reveals the concise characterization of the tangent space for the $n \times p$ complex Stiefel manifold:

$$T_X \text{St}(n, p) = \{\Delta \in \mathbb{C}^{n \times p} : X^\dagger \Delta + \Delta^\dagger X = 0\}. \quad (2.29)$$

Let us define $X_\perp$ such that the matrix $[X \ X_\perp]$ is unitary, i.e.

$[X \ X_\perp]^\dagger [X \ X_\perp] = \mathbb{1}_n$. Let further $\text{Skew}(p)$ denote the set of $p \times p$ skew hermitian matrices and $\text{Herm}(p)$ denote the set of $p \times p$ hermitian matrices. The standard way of parametrizing elements of the tangent space [58, 60] is then $\Delta = XS + X_\perp C$, where $S \in \text{Skew}(p)$ and C is an arbitrary $(n-p) \times p$ complex matrix. Since by definition of $X_\perp$ its columns are orthogonal to the columns of X , and it holds that $X^\dagger X_\perp = 0$ one can easily see that $\Delta = XS + X_\perp C$ satisfies the tangent space condition of Eq. (2.29).

The obvious choice of metric for the Stiefel manifold would be given by the standard inner product defined in Eq. (2.16), where the smoothness condition is satisfied if we treat $\langle \cdot, \cdot \rangle$ as a function on $\mathbb{R}^{2n \times p}$. This is however not the most natural choice for $\text{St}(n, p)$. Note that the standard metric is given in terms of the entries of S and C as

$$\text{Re Tr}[\Delta^\dagger \Delta] = \text{Re Tr}[(S^\dagger X^\dagger + C^\dagger X_\perp^\dagger)(XS + X_\perp C)] = \text{Re Tr}[S^\dagger S] + \text{Re Tr}[C^\dagger C] = \sum_{ij} |S_{ij}|^2 + \sum_{ij} |C_{ij}|^2,$$

where we used $X_\perp^\dagger X_\perp = \mathbb{1}_{n-p}$ and $X^\dagger X_\perp = 0$. As $S_{ij} = -S_{ji}^*$, the standard metric weights the independent variables $\{S_{ij}\}_{i>j}$ and $\{C_{ij}\}_{i,j}$ unevenly. One thus defines the *canonical inner product* or *canonical metric* for the Stiefel manifold via

$$\langle \Delta, \Omega \rangle_c := \text{Re Tr}[\Delta^\dagger (\mathbb{1} - XX^\dagger / 2) \Omega]. \quad (2.30)$$

It is straightforward to show that $\langle \Delta, \Delta \rangle_c = \sum_{i>j} |S_{ij}|^2 + \sum_{ij} |C_{ij}|^2$ and all degrees of freedom are weighted equally. The canonical inner product also ensures that the natural basis vectors on the tangent space $\{X_\perp E_{ij}, X(E_{ij} - E_{ji}), X i(E_{ij} - E_{ji})\}$ are orthonormal (see also [60]).

After settling for the canonical metric, we determine the normal space to $T_X \mathcal{M}$ to be

$$N_X \text{St}(n, p) = \{XH : H \in \text{Herm}(p)\}. \quad (2.31)$$

Using $\text{Re Tr}[S^\dagger H] = 0$ for any hermitian H and skew hermitian S , one can verify that indeed $\langle XS + X_\perp C, XH \rangle_c = 0$ for all $S \in \text{Skew}(p)$, $H \in \text{Herm}(p)$, $C \in \mathbb{C}^{(n-p) \times p}$. Finally, one would also like to be able to project from the ambient space to the tangent space and to the normal space. We define

$$P_{N_X}(Y) = X(Y^\dagger X + X^\dagger Y)/2, \quad (2.32)$$

which is a projector since $P_{N_X}^2(Y) = P_{N_X}(Y)$ and moreover $P_{N_X}(Y) \in N_X \text{St}(n, p)$ since $Y^\dagger X + X^\dagger Y$ is hermitian for all X, Y . The projector onto the tangent space is then simply $P_{T_X}(Y) = Y - P_{N_X}(Y)$.

Retractions and affine connections

We define the set of all tangent vectors to $\mathcal{M}$ as the *tangent bundle* $T\mathcal{M} = \bigcup_{X \in \mathcal{M}} T_X \mathcal{M}$. Let $R : T\mathcal{M} \mapsto \mathcal{M}$ which is a map that lets us relate elements from any given tangent space to elements on the manifold. Let further $R_X = R|_{T_X \mathcal{M}}$ and $\text{id}_X : T_X \mathcal{M} \mapsto T_X \mathcal{M} : \Delta \mapsto \Delta$. We call R a *retraction* if for all X the restrictions R_X are smooth functions that satisfy

- (i) $R_X(0) = X$ (with $0 \in T_X \mathcal{M}$),
- (ii) For every curve $\gamma(t) := R_X(t\Delta)$ it holds that $\dot{\gamma}(0) = \Delta$.

The first condition tells us that with respect to R_X , the tangent space is centered at X . The second condition ensures that the curve traced by the retraction is faithful to the tangent space direction at its origin, i.e.

$\dot{R}_X(t\Delta) = \text{id}_X$ at the origin. A retraction can also be used in an optimization algorithm to trace a cost function $f : \mathcal{M} \mapsto \mathbb{R}$ through the tangent space, by defining $\tilde{f}_X : T_X \mathcal{M} \mapsto \mathbb{R} : \tilde{f}_X(\Delta) = f(R_X(\Delta))$ in a neighborhood of X . The function $\tilde{f}$ is then called the *pullback* of f .

This is especially useful for the line search problem: Imagine you are given an update direction $\Delta \in T_X \mathcal{M}$ and want to follow it until a local minimum is reached. The resulting sub-problem

$$\hat{t} = \underset{t \in \mathbb{R}}{\operatorname{argmin}} \tilde{f}(t\Delta) \quad (2.33)$$

leads to the optimal update step $X \rightarrow R_X(\hat{t}\Delta)$.

An affine connection is a structure on a tangent bundle that generalizes the directional derivative of a vector field on $\mathcal{M}$ and is essential to have a notion of a second order derivative on a manifold. A vector field on $\mathcal{M}$ is a smooth function $\Omega : \mathcal{M} \mapsto T\mathcal{M} : X \mapsto \Omega(X)$. The addition of vector fields is defined entry-wise as $(\Omega_1 + \Omega_2)(X) := \Omega_1(X) + \Omega_2(X)$. A vector field can be multiplied by a scalar field $f : \mathcal{M} \mapsto \mathbb{R}$ via the rule $(f\Omega)(X) := f(X)\Omega(X)$. Let $\gamma^\Omega(t)$ a curve that realizes the tangent vector $\Omega(X)$ as in Eq. (2.20). Using this notion we can define

$$(\Omega f)(X) := \Omega(X)f = \left. \frac{df(\gamma^\Omega(t))}{dt} \right|_{t=0}, \quad (2.34)$$

meaning $(\Omega f)(X) : \mathcal{M} \mapsto \mathbb{R}$ is a function that assigns to each X the derivative of f in the direction determined by the vector field Ω at X .

The directional derivative ∇_Ξ on $\mathcal{M} = \mathbb{R}^n$ of a vector field Ω in direction $\Xi \in T_X \mathbb{R}^n \equiv \mathbb{R}^n$ at position X is just

$$(\nabla_\Xi \Omega)(X) = \lim_{t \rightarrow 0} \frac{\Omega(X + t\Xi) - \Omega(X)}{t}. \quad (2.35)$$

On a manifold this definition is a priori ill-defined since we can neither perform the addition $X + t\Xi$ nor $\Omega(X + t\Xi) - \Omega(X)$ since $\Omega(X + t\Xi)$ and $\Omega(X)$ would belong to different tangent spaces. However, the operator ∇_Ξ can be generalized as follows. We first define $\mathcal{V}(\mathcal{M})$ to be the set of smooth vector fields on $T\mathcal{M}$. Then $\nabla : \mathcal{V}(\mathcal{M}) \times \mathcal{V}(\mathcal{M}) \mapsto \mathcal{V}(\mathcal{M})$ is an affine connection if it satisfies the properties

$$\begin{aligned} (i) \quad & \nabla_{f\Omega + g\Xi} = f\nabla_\Omega + g\nabla_\Xi \quad \text{for all } f, g : \mathcal{M} \mapsto \mathbb{R}, \\ (ii) \quad & \nabla_{a\Omega + b\Xi} = a\nabla_\Omega + b\nabla_\Xi \quad \text{for all } a, b \in \mathbb{R}, \\ (iii) \quad & \nabla_\Xi(f\Omega) = (\Xi f)\Omega + f\nabla_\Xi \Omega. \end{aligned} \quad (2.36)$$

The product rule (iii) can be hard to parse at first sight but becomes more clear once the spaces each object belongs to are marked:

$$\nabla_\Xi \left(\underbrace{f\Omega}_{\mathcal{M} \mapsto T\mathcal{M}} \right) = \underbrace{\left(\underbrace{\Xi f}_{\mathcal{M} \mapsto \mathbb{R}} \right)}_{\mathcal{M} \mapsto \mathbb{R}} \underbrace{\Omega}_{\mathcal{M} \mapsto T\mathcal{M}} + \underbrace{f}_{\mathcal{M} \mapsto \mathbb{R}} \underbrace{\nabla_\Xi \Omega}_{\mathcal{M} \mapsto T\mathcal{M}}. \quad (2.37)$$

In order to specify an affine connection one can define a basis or coordinate vector field, which associates with each point on the manifold an elementary direction on its tangent space. This depends on the local chart, so let $X \in \mathcal{U} \subset \mathcal{M}$ with a chart $\varphi : \mathcal{M} \mapsto \mathbb{R}^n$ defined on $\mathcal{U}$. Given a basis vector e_i on $\mathbb{R}^n$, we define a corresponding basis vector on $T_X \mathcal{M}$ via the curve $\gamma(t) = \varphi^{-1}(\varphi(X) + t \cdot e_i)$. The tangent vector $E_i(X)$ defined by the curve then according to Eq. (2.20) acts as

$$E_i(X)f = \left. \frac{df(\gamma(t))}{dt} \right|_{t=0} = \frac{d}{dt} (f \circ \varphi^{-1})(\varphi(X) + te_i)|_{t=0} = \partial_i(f \circ \varphi^{-1}) \quad (2.38)$$

Here ∂_i is the partial derivative in direction e_i on $\mathbb{R}^n$ given as $\partial_i g := \lim_{t \rightarrow 0} \frac{g(x + te_i) - g(x)}{t}$. To obtain basis vectors on all $X \in \mathcal{U} \subset \mathcal{M}$ we define the vector field $E_i : \mathcal{M} \mapsto T_X \mathcal{M} : X \mapsto E_i(X)$. Upon repeating this procedure for all basis vectors e_i we get the set of vector fields $\{E_i\}_{i=1}^n$ that we use to represent an arbitrary vector field Ω via $\Omega = \sum_{i=1}^n \Omega_i E_i$. Here $\Omega_i : \mathcal{M} \mapsto \mathbb{R}$ are coordinate functions that contain the information about the whole vector field.

A given affine connection ∇_{Ξ} can now be represented via the coordinate vector fields as follows:

$$\nabla_{\Xi}\Omega = \nabla_{\sum_i \Xi_i E_i} \sum_j \Omega_j E_j = \sum_{ij} \Xi_i \Omega_j \nabla_{E_i} E_j + \Xi_i (E_i \Omega_j) E_j, \quad (2.39)$$

where we used the linearity condition (i) and the product rule (iii) in Eq. (2.36). We thus obtained the affine connection in terms of the elementary actions $\nabla_{E_i} E_j$. These can again be expressed in our basis $\{E_k\}$ and the resulting coefficients Γ_{ij}^k are called *Christoffel symbols*:

$$\nabla_{E_i} E_j = \sum_k \Gamma_{ij}^k E_k. \quad (2.40)$$

Using the Christoffel symbols, the action of the affine connection is expressed as

$$\nabla_{\Xi}\Omega = \sum_{ij} \left(\Xi_i (\partial_i \Omega_j) E_j + \Xi_i \Omega_j \sum_k \Gamma_{ij}^k E_k \right). \quad (2.41)$$

An affine connection given by Christoffel symbols is not unique in two ways. First, the underlying chart φ determines the coefficients Γ_{ij}^k and second, there are infinitely many affine connections that satisfy Eq. (2.36) to begin with. However, the following fundamental theorem asserts that under reasonable additional conditions, there exists a unique affine connection.

Theorem 6 (Theorem 5.3.1 in Absil *et al.* [59]). *There exists a unique affine connection called the Riemannian connection on a Riemannian manifold $\mathcal{M}$ with a given atlas $\mathcal{A}$ that satisfies*

$$\begin{aligned} (i) \quad & \Gamma_{ij}^k = \Gamma_{ji}^k \\ (ii) \quad & \Omega \langle \Xi, \Theta \rangle = \langle \nabla_{\Omega} \Xi, \Theta \rangle + \langle \Xi, \nabla_{\Omega} \Theta \rangle \end{aligned} \quad (2.42)$$

where Ω, Ξ, Θ are vector fields on $\mathcal{M}$.

The first condition ensures symmetry, i.e.

$\nabla_{E_i} E_j = \nabla_{E_j} E_i$. The second condition is a product rule with respect to the Riemannian metric, which on vector fields acts as $\langle \cdot, \cdot \rangle : \mathcal{V}(\mathcal{M}) \times \mathcal{V}(\mathcal{M}) \mapsto \mathcal{F}(\mathcal{M})$, where $\mathcal{F}(\mathcal{M})$ is the set of real valued functions on $\mathcal{M}$ and $\langle \Xi, \Theta \rangle(X) \mapsto \langle \Xi(X), \Theta(X) \rangle \in \mathbb{R}$.

We can also define the metric with respect to the basis $\{E_i\}$, by noting that $\langle \cdot, \cdot \rangle$ is determined via a matrix g with entries $g_{ij} := \langle E_i, E_j \rangle$. It can then be shown [59] that the Christoffel symbols for the unique affine connection on $\mathcal{M}$ are related to the metric coefficients via

$$\Gamma_{ij}^k = \frac{1}{2} \sum_l g_{kl}^{-1} (\partial_i g_{lj} + \partial_j g_{li} - \partial_l g_{ij}). \quad (2.43)$$

In case we are dealing with an embedded submanifold to a Riemannian manifold $\bar{\mathcal{M}}$, Proposition 7 below shows that evaluating the Riemannian connection amounts to evaluating the (potentially simpler) Riemannian connection on the embedding space. Just as we saw for the Stiefel manifold, if the embedding space $\bar{\mathcal{M}}$ is a vector space, its elements can be uniquely written as $\bar{X} = \Delta + \Delta_{\perp} = P_{T_X}(X) + P_{N_X}(X)$ with $\Delta \in T_X \mathcal{M}, \Delta_{\perp} \in N_X \mathcal{M}$ and P_{T_X}/P_{N_X} being the projector onto the tangent space/normal space.

Proposition 7 (Proposition 5.3.2 in Absil *et al.* [59]). *Let ∇ and $\bar{\nabla}$ be Riemannian connections on $\mathcal{M}$ and $\bar{\mathcal{M}}$ respectively, then ∇ can be evaluated at $X \in \mathcal{M}$ as*

$$\nabla_{\Omega(X)} \Xi = P_{T_X} \bar{\nabla}_{\Omega(X)} \Xi. \quad (2.44)$$

Furthermore, if $\bar{\mathcal{M}}$ is a Euclidean space, then $\nabla_{\Omega(X)} \Xi = P_{T_X} D \Xi(X)[\Omega(X)]$ with D being the standard directional derivative.

The latter statement provides a considerable simplification for practical purposes, since for embedded submanifolds of a Euclidean space, only the tangent space projection needs to be determined and no Riemannian connection needs to be derived.

Geodesics

Geodesics are curves on a manifold that connect points via their shortest path. They are consequently the most natural paths to follow in an iterative optimization algorithm on $\mathcal{M}$. As we will see shortly, Geodesics also constitute a special type of retraction, one that does not only follow a given tangent space direction but additionally has zero curvature at the origin. They are thus locally the most faithful to update directions computed on the tangent space. A disadvantage is that there are usually other retractions which are computationally cheaper to evaluate. In Euclidean space, geodesics are just straight lines, which coincide with curves $\gamma(t)$ that have vanishing acceleration $\frac{d^2}{dt^2}\gamma(t)$.

We start by defining the velocity vector field $\dot{\gamma}(t)$ on a curve $\gamma : [a, b] \mapsto \mathcal{M}$ in a straightforward way by setting $\dot{\gamma} : [a, b] \mapsto T\mathcal{M} : t \mapsto \dot{\gamma}(t)$, where $\dot{\gamma}(t) \in T_X\mathcal{M}$ is the tangent vector to the curve at position $\gamma(t)$. We can write a general vector field on a curve to be the restriction of a vector field Ω to the curve via $\Omega_{\gamma(t)} = \Omega \circ \gamma(t)$. For the acceleration we need a way to quantify changes from $\dot{\gamma}(t)$ to $\dot{\gamma}(t + dt)$, but since they live on different tangent spaces, they are not directly comparable. To remedy this issue one defines a new function $\frac{D}{dt} : \mathcal{V}(\mathcal{M}) \mapsto \mathcal{V}(\mathcal{M})$ which acts as

$$\frac{D}{dt}(\Omega \circ \gamma)(t) := \nabla_{\dot{\gamma}(t)}\Omega. \quad (2.45)$$

The map $\frac{D}{dt}$ is also called the *covariant derivative* induced by the connection ∇ . If we require linearity $\frac{D}{dt}(a\Omega + b\Xi) = a\frac{D}{dt}\Omega + b\frac{D}{dt}\Xi \forall a, b \in \mathbb{R}$ and the product rule $\frac{D}{dt}(f\Omega) = f'\Omega + f\frac{D}{dt}\Omega \forall f : [a, b] \mapsto \mathbb{R}$, then it can be shown that the covariant derivative $\frac{D}{dt}$ is unique. The acceleration vector field for a curve $\gamma(t)$ is then defined as

$$\frac{D^2}{dt^2}\gamma(t) := \frac{D}{dt}\dot{\gamma}(t). \quad (2.46)$$

We are thus equipped to define a *geodesic*, which is any curve $\gamma : [a, b] \mapsto \mathcal{M}$ with vanishing acceleration according to the covariant derivative:

$$\frac{D^2}{dt^2}\gamma(t) = \frac{D}{dt}\dot{\gamma}(t) = 0. \quad (2.47)$$

This is also called the geodesic equation. Akin to Eq. (2.41), we can write the acceleration in terms of the basis vector fields $\{E_i\}$ with the velocity parametrized as $\dot{\gamma}(t) = (\dot{\gamma}(t))_i E_i$. We can use Eq. (2.45) to write the acceleration as $\frac{D^2}{dt^2}\gamma(t) = \nabla_{\dot{\gamma}(t)}\dot{\gamma}(t)$, whereafter the product rule and the linearity of the connection allow us to compute the acceleration as

$$\begin{aligned} \nabla_{\dot{\gamma}(t)}\dot{\gamma}(t) &= \nabla_{\dot{\gamma}(t)} \sum_i (\dot{\gamma}(t))_i E_i \\ &= \sum_i \left((\ddot{\gamma}(t))_i E_i + (\dot{\gamma}(t))_i \nabla_{\sum_j (\dot{\gamma}(t))_j E_j} E_i \right) \\ &= \ddot{\gamma}(t) + \sum_{ij} (\dot{\gamma}(t))_i (\dot{\gamma}(t))_j \nabla_{E_j} E_i \\ &= \ddot{\gamma}(t) + \sum_{i,j,k} \Gamma_{ij}^k (\dot{\gamma}(t))_i (\dot{\gamma}(t))_j E_k. \end{aligned} \quad (2.48)$$

In Appendix A we derive the geodesic equation for the Stiefel manifold by using the above form on the ambient space and invoking Proposition 7. This leads to the condition

$$P_{T_{\gamma(t)}} \left(\ddot{\gamma}(t) + \dot{\gamma}(t)\dot{\gamma}(t)^\dagger \gamma(t) - \gamma(t)\dot{\gamma}(t)^\dagger \dot{\gamma}(t) - \dot{\gamma}(t)\gamma(t)^\dagger \dot{\gamma}(t) \right) = 0, \quad (2.49)$$

which we use to determine the geodesics.

The Riemannian Hessian

For standard optimization algorithms that also take into account second order approximations of a cost function, such as the Newton method, one requires the Hessian matrix in Euclidean space. The Hessian can be generalized for a real valued function f at position $X \in \mathcal{M}$ via the linear map $Hf(X) : T_X\mathcal{M} \mapsto T_X\mathcal{M}$, with

$$Hf(X)[\Omega] = \nabla_\Omega Gf, \quad (2.50)$$

which is called the *Riemannian Hessian*. Here $\Omega \in T_X\mathcal{M}$ and Gf is the Riemannian gradient (Eq. (2.25)) of f . In analogy to the Hessian in Euclidean space, the Riemannian Hessian can be shown to be symmetric with respect to the Riemannian metric: $\langle Hf(X)[\Omega], \Delta \rangle = \langle \Omega, Hf(X)[\Delta] \rangle$. Just as the Hessian matrix on $\mathbb{R}^n$ can be seen as a bilinear map $\tilde{H} : \mathbb{R}^n \times \mathbb{R}^n \mapsto \mathbb{R} : \tilde{H}(u, v) \mapsto u^T \tilde{H} v$, we can define $\text{Hess} : T_X\mathcal{M} \times T_X\mathcal{M} \mapsto \mathbb{R}$ which acts as

$$\text{Hess}(\Delta, \Omega) = \langle Hf(X)[\Delta], \Omega \rangle. \quad (2.51)$$

A useful trick to compute the Riemannian Hessian is given by the following proposition, which uses the derivative of the cost function along a certain type of retraction.

Proposition 8. *Let $R : T_X\mathcal{M} \mapsto \mathcal{M}$ be a retraction that satisfies*

$$\left. \frac{D^2}{dt^2} R(t\Delta) \right|_{t=0} = 0 \quad (2.52)$$

for all $\Delta \in T_X\mathcal{M}$ then

$$\text{Hess}(\Delta, \Omega) = \frac{1}{2} \left. \frac{d^2}{dt^2} [f(R(t(\Delta + \Omega))) - f(R(t\Delta)) - f(R(t\Omega))] \right|_{t=0} \quad (2.53)$$

For the proof see Proposition 5.5.5 and its discussion in Absil *et al.* [59]. An obvious example of a retraction that satisfies Eq. (2.52) (called a second order retraction) is given by a geodesic $\dot{\gamma}(t)$ with initial condition $\dot{\gamma}(0) = \Delta$.

2.2 Characterization of quantum dynamics

A high degree of control over quantum dynamics is a fundamental requirement to achieve a quantum advantage in computational tasks. Although improvements of experimental control are largely made through the experimentalists understanding of the device physics, there is a need to validate these physical models and to find potentially unknown error contributions just from measurement data with minimal prior assumptions.

There is generally a correlation between the stringency of assumptions made and the efficiency of a characterization protocol. Similarly, the more expressive a theoretical model of the dynamics is, the higher the minimal measurement and classical post-processing complexity required to attain the parameters of the model from data. In the following we summarize many commonly taken routes in the literature to square these counteracting requirements. While benchmarking protocols aim to provide a single or a small number of error measures to quantify the performance of the device dynamics, device identification or learning aims to provide a mathematical model of the dynamics that can be interpreted regarding error sources. The obvious use is for device calibration, but a full model can serve several other purposes. First, rigorous worst case error measures can be computed from the mathematical model, which serve as thresholds for fault-tolerant quantum computation [70, 71]. Second, a concrete noise model can be used to actively design quantum circuits that mitigate said noise, for instance via probabilistic error cancellation [72]. In Chapter 3 we further demonstrate an example of error mitigation via post-processing of data from a noisy device. A third and often overlooked additional use of a fully characterized noise

model is for the testing of proposed quantum algorithm or benchmarking protocols with classical simulations. Since accurate noise models from characterization experiments are unfortunately rare in the literature, simulations are often done with the most simplistic noise models, which do not accurately capture real world devices.

In the following we first familiarize ourselves with the fundamentals of quantum dynamics from a quantum information theory point of view in Section 2.2.1 and Section 2.2.2. We further review different approaches to the characterization of quantum dynamics, starting with quantum process tomography in Section 2.2.3 and later focusing on self-consistent approaches in Section 2.2.4. This will serve as a basis to understand and contextualize our research on gate set tomography which is presented in Chapter 3.

2.2.1 Quantum operations and the superoperator formalism

Textbook quantum mechanics describes dynamics on quantum systems in terms of unitary evolutions on pure states $|\psi\rangle$, which are elements of the Hilbert space $\mathcal{H}$. Here we instead use the formalism deployed in quantum information theory, where the derivations presented in the following are based on the standard works by Nielsen and Chuang [73] and Watrous [74]. In this setting, one separates the Hilbert space into the system $\mathcal{H}_S$ and the environment $\mathcal{H}_E$, where we assume that the initial state on $\mathcal{H}_S \otimes \mathcal{H}_E$ is given by the product state $\rho_0 \otimes |E_0\rangle\langle E_0|$. The environment is chosen to be in a pure state, since this case is sufficiently general for our purposes. After a unitary evolution U on the combined space takes place, one ‘forgets’ about the environment by taking a partial trace, resulting in the final state

$$\rho = \sum_{i=1}^{\dim(\mathcal{H}_E)} \langle i|U(\rho_0 \otimes |E_0\rangle\langle E_0|)U^\dagger|i\rangle, \quad (2.54)$$

where $\{|i\rangle\}_{i=1}^{\dim(\mathcal{H}_E)}$ is an orthonormal basis on $\mathcal{H}_E$. We then define $K_i := \langle i|U|E_0\rangle \in L(\mathcal{H}_S)$ in order to write the final state as

$$\rho = \mathcal{C}(\rho_0) := \sum_{i=1}^{\dim(\mathcal{H}_E)} K_i \rho_0 K_i^\dagger. \quad (2.55)$$

It follows that we can describe physical dynamics on the system by a linear map $\mathcal{C} : L(\mathcal{H}_S) \mapsto L(\mathcal{H}_S)$, which is given by the set $\{K_i\}$, whose elements are commonly referred to as *Kraus operators*. Since we don’t require the environment in the following, we use just $\mathcal{H}$ for $\mathcal{H}_S$. One can immediately see from Eq. (2.55) that $1 = \text{Tr}[\rho] = \text{Tr}[\sum_i K_i \rho_0 K_i^\dagger] = \text{Tr}[\sum_i K_i^\dagger K_i \rho_0]$ has to hold for all ρ_0 . We thus get the condition

$$\sum_i K_i^\dagger K_i = \mathbb{1}. \quad (2.56)$$

Furthermore, it can directly be seen that $\sum_i (K_i \otimes \mathbb{1}_{d'}) \rho_0 (K_i^\dagger \otimes \mathbb{1}_{d'}) \succeq 0$ for all positive semidefinite matrices $\rho_0 \in \mathcal{H} \otimes \mathcal{H}'$ and $d' = \dim(\mathcal{H}')$. A linear map $\mathcal{C}$ that can be written in terms of Kraus operators is thus always completely positive, i.e.

, $(\mathcal{C} \otimes \text{id})(\rho) \succeq 0$ for all $\rho \succeq 0$, with id being the identity channel. In fact, the converse also holds: Any completely positive trace preserving linear map can be written in the form of Eq. (2.55), as we will see below. Due to this equivalence, a completely positive trace preserving (CPT) linear map is also called a *quantum channel*, since it constitutes the most general form of physical dynamics according to the above assumptions.

Before we delve deeper, we will first familiarize ourselves with a tensor network notation for quantum circuits, which allows for an intuitive understanding of the above concepts.

A tensor is a multilinear map $T : V_1 \otimes \cdots \otimes V_r \mapsto V_{r+1} \otimes \cdots \otimes V_{r+s}$, which can be written with respect to bases on the vector spaces $V_1, \dots, V_{r+s}$. The basis coefficients $T_{i_1, \dots, i_r}^{i_{r+1}, \dots, i_{r+s}}$ determine

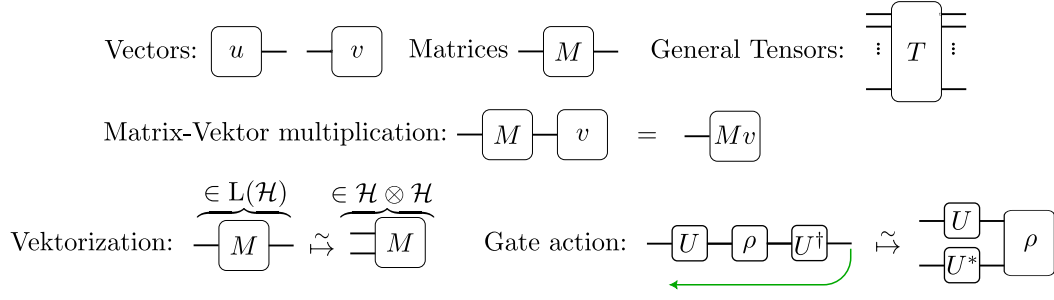


Figure 2.1: Basic elements of the tensor network representation. In our convention, matrices are always applied to to right, i.e.

for $M = \sum_{i,j} M_{ij} e_i^T \otimes e_j$ with $\{e_i\}_{i=1}^n$ being the standard basis vectors, the leg corresponding to the index j is always on the right and the leg corresponding to i on the left. Vectorization is always done in row-major order, i.e.

$\text{vec}(M) = \sum_{i,j} M_{ij} e_i \otimes e_j$. The green arrow indicates the rearrangement of tensor legs, which encapsulates the isomorphism in the bottom right.

the tensor and its action on a set of vectors $\{v_1, \dots, v_n\}$ with $n \leq r + s$, which is given by

$$(T(v_1, \dots, v_n))^{i_{n+1} \dots i_{r+s}} = \sum_{i_1, \dots, i_n} T_{i_1, \dots, i_r}^{i_{r+1} \dots i_{r+s}} (v_1)_{i_1} \cdots (v_n)_{i_n}. \quad (2.57)$$

This is termed a *contraction* of T and $v_1, \dots, v_n$. To avoid excessive use of indices, these contractions can be pictorially represented in tensor network diagrams. A tensor is therein associated with a box from which $r + s$ lines (called 'legs') emanate. Contractions between tensors are represented by joining lines together. For a representation of the basic building blocks that we need for our purposes, see Figure 2.1. These tensor network diagrams are commonly used in the literature, and a more thorough introduction can be found for instance in Wood *et al.* [75].

A useful concept for understanding the structure of quantum channels is the Choi–Jamiołkowski isomorphism, which identifies a quantum channel $\mathcal{C} : L(\mathcal{H}) \mapsto L(\mathcal{H})$ with a quantum state $\mathcal{J}(\mathcal{C}) \in L(\mathcal{H}^{\otimes 2})$. Let $|\phi^+\rangle = 1/\sqrt{d} \sum_{i=1}^d |\mathbb{1}\rangle$ with $d = \dim(\mathcal{H})$. Then the Choi-state $\mathcal{J}(\mathcal{C})$ is given as

$$\mathcal{J}(\mathcal{C}) = \mathcal{C} \otimes \text{id}(|\phi^+\rangle\langle\phi^+|). \quad (2.58)$$

A tensor network representation of the isomorphism given in Figure 2.2 provides a simple understanding of what the map $\mathcal{J}$ does: a reordering and recombining of tensor indices.

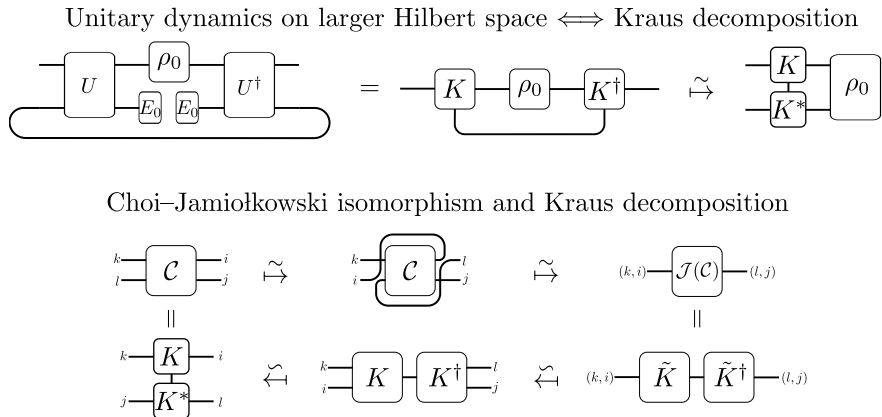


Figure 2.2: Tensor network representations for equivalent descriptions of quantum channels.

Since $|\phi^+\rangle\langle\phi^+| \succeq 0$ and $\mathcal{C}$ is completely positive, one can immediately see that $\mathcal{J}(\mathcal{C}) \succeq 0$. Any positive semidefinite matrix can be decomposed as $\mathcal{J}(\mathcal{C}) = A^\dagger A$ for some matrix $A \in L(\mathcal{H}^{\otimes 2})$

(for instance via the Cholesky decomposition). We can now relate this decomposition of the Choi matrix to the Kraus decomposition of the original channel $\mathcal{C}$ as follows. The matrix entries of $A \in \mathbb{C}^{d^2 \times d^2}$ can be reordered into a 3-legged tensor by reshaping each of its rows into a $d \times d$ matrix. We then identify the resulting set of d^2 matrices with the set of Kraus operators $\{K_i\}_{i=1}^{d^2}$. Upon reordering the indices (reversing the first step of the Choi isomorphism), we arrive at the Kraus decomposition of $\mathcal{C}$ (see Figure 2.2). We have thus seen that there is a one-to-one correspondence between completely positive linear maps and maps that admit a Kraus decomposition.

It is important to note that the Kraus decomposition is not unique, as can be easily seen by considering that the decomposition of $\mathcal{J}(\mathcal{C})$ is not unique: $\mathcal{J}(\mathcal{C}) = A^\dagger A = A^\dagger U^\dagger U A$. On the level of the Kraus operators $\{K_i\}_{i=1}^{d^2}$, this amounts to a reshuffling, where $K_i \mapsto \sum_j K_j U_{ji}$.

We treat measurements in the usual way, by defining a positive operator valued measure (POVM), consisting of positive operators $\{E_i\}$ that satisfy $\sum_i E_i = \mathbb{1}$. The probability for obtaining outcome i given a quantum system in the state ρ is then computed via the so-called Born rule: $\mathbb{P}[i] = \text{Tr}[E_i \rho]$. Positivity of E_i and $\sum_i E_i = \mathbb{1}$ ensure that all probabilities are positive and sum to one, respectively.

Just as we wrote the action of individual quantum channels via tensors, we can do so for an entire quantum experiment. From a quantum information theory point of view, an experiment consists of the preparation of an initial state ρ (most typically $\rho = |0\rangle\langle 0|$), and a sequence of m gates followed by a POVM measurement. Figure 2.3 shows the tensor networks of an idealized experiment with unitary dynamics, as well as a general experiment with quantum channels in Kraus decomposition. For the latter, the tensor network diagram provides a considerably simplified picture of the underlying circuit. Learning a general quantum channel becomes infeasible

$$\begin{aligned} \text{Tr}[E U_n \cdots U_1 \rho U_1^\dagger \cdots U_n^\dagger] &= \text{Diagram 1} = \text{Diagram 2} \\ &= \sum_{i_1, \dots, i_m} \text{Tr}[E K_{m, i_m} \cdots K_{1, i_1} \rho K_{1, i_1}^\dagger \cdots K_{m, i_m}^\dagger] = \text{Diagram 3} \end{aligned}$$

The figure contains three tensor network diagrams. The first diagram represents the trace of a sequence of unitaries $U_1, \dots, U_n$ acting on an initial state ρ_0 , with adjoint unitaries $U_1^\dagger, \dots, U_n^\dagger$ following. The second diagram shows the same trace, but the unitaries are rearranged into a product of $U_i \otimes U_i^*$ acting on $\text{vec}(\rho_0)$. A green arrow indicates the rearrangement of the adjoint unitaries. The third diagram shows the trace as a sum over Kraus operators $K_{i,j}$ and their adjoints $K_{i,j}^\dagger$, with a POVM element E at the end.

Figure 2.3: Unitary quantum circuits and general circuits composed of CPT maps visualized as tensor networks. The green arrow indicates how the matrices $U_1^\dagger, \dots, U_n^\dagger$ are rearranged so that the circuit acts as a product of matrices $U_i \otimes U_i^*$ on $\text{vec}(\rho_0)$.

very quickly for increased system size, since the number of free (real) parameters of a linear map in $\mathbb{C}^{d^2 \times d^2}$ is $2d^4$. Even the physicality constraints do not help much, since for the Kraus decomposition we have $\sum_i K_i^\dagger K_i = \mathbb{1} \in \mathbb{C}^{d \times d}$, which amounts to d real constraints (the diagonal entries of $\mathbb{1}$) and $d(d-1)/2$ complex constraints (off-diagonal entries of $\mathbb{1}$). In total, we then have $2d^4 - d^2$ free parameters in a quantum channel given via the Kraus representation. If we further remove the unitary freedom from the non-uniqueness of the Kraus operators we end up with $d^4 - d^2$ real parameters (see the discussion in Section 3.1). There are, however, simple physically motivated noise models that only require one or two Kraus operators and thus have considerably fewer free parameters. This leads us to define the *Kraus rank* of a quantum channel,

$$r_K(\mathcal{C}) = \min_r \left\{ r \in \mathbb{N} : \exists \{K_i\}_{i=1}^r : \mathcal{C} = \sum_{i=1}^r K_i \otimes K_i^* \text{ and } \sum_{i=1}^r K_i^\dagger K_i = \mathbb{1} \right\}, \quad (2.59)$$

which coincides with the rank of the Choi state $\mathcal{J}(\mathcal{C})$.

In the following, we highlight a different channel parametrization in terms of its action on Pauli operators, which is sometimes more intuitive to understand and directly incorporates the trace preservation constraint.

The Pauli transfer matrix representation

Let $\mathcal{X} : L(\mathcal{H}) \mapsto L(\mathcal{H})$ be a linear map (also called superoperator). We use the Dirac notation for elements of $L(\mathcal{H})$ (see Section 2.1.1) to write $\mathcal{X}$ with respect to any orthogonal basis on $L(\mathcal{H})$, where the basis elements $|B_i\rangle$ satisfy $(B_i|B_j) = \delta_{i,j}$. This could be the standard basis with $\{B_0 = |0\rangle\langle 0|, B_1 = |0\rangle\langle 1|, \dots\}$, or the set of normalized Pauli matrices $\{\check{\sigma}_a := \sigma_a/\sqrt{d}\}_{a \in \mathbb{F}_2^{2n}}$. Then $\mathcal{X}$ can be written with respect to this basis as $\mathcal{X} = \sum_{a,a'} X_{a,a'} |\check{\sigma}_a\rangle\langle \check{\sigma}_{a'}|$.

If a quantum channel is written in this basis as $\mathcal{C} = \sum_{a,a'} C_{a,a'} |\check{\sigma}_a\rangle\langle \check{\sigma}_{a'}|$, the matrix C is called the *Pauli transfer matrix*, or PTM for short. Since a quantum channel is hermiticity preserving, and since any hermitian matrix admits a decomposition in the normalized Pauli basis with real coefficients, we can see that the coefficients $C_{a,a'}$ must be real themselves: $C_{a,a'} = (\check{\sigma}_a|\mathcal{C}|\check{\sigma}_{a'}) \in \mathbb{R}$. The trace preservation condition is also very simple in the PTM-representation, since

$$1 = \text{Tr}[\mathcal{C}(\rho)] = (\mathbb{1}|\mathcal{C}|\rho) = \sqrt{d}(\check{\sigma}_0|\mathcal{C}|\rho) = \sqrt{d} \sum_a C_{0,a}(\check{\sigma}_a|\rho) = C_{0,0} + \sqrt{d} \sum_{a \neq 0} C_{0,a}(\check{\sigma}_a|\rho) \quad (2.60)$$

has to hold for every ρ . This implies that $C_{0,0} = 1$ and $C_{0,a} = 0$ for $a \neq 0$, meaning the first row of the matrix C is the first unit vector. We can thus write C as

$$C = \begin{pmatrix} 1 & 0 & \dots & 0 \\ | & & & \\ u & & T & \\ | & & & \end{pmatrix}, \quad (2.61)$$

where $T \in \mathbb{R}^{(d^2-1) \times (d^2-1)}$ and $u \in \mathbb{R}^{d^2-1}$, with a total of $d^4 - d^2$ free parameters. A channel $\mathcal{C}$ is commonly called *unital* if it satisfies $\mathcal{C}(\mathbb{1}) = \mathbb{1}$, which translates to $u = 0$ in our parametrization. The PTM representation can also be used to make general statements about the eigenvalues $\{\lambda_i\}$ of the channel. We first note that the eigenvalues of C are the eigenvalues of T in addition to the eigenvalue $\lambda_0 = 1$. This holds since for a block matrix as in Eq. (2.61) we have $\det(C - \lambda \mathbb{1}_{d^2}) = (1 - \lambda) \det(T - \lambda \mathbb{1}_{d^2-1})$. Since T is a real matrix, its eigenvalues are either real or for every complex eigenvalue the conjugate is an eigenvalue as well, i.e.

$|\lambda|e^{i\phi}$ and $|\lambda|e^{-i\phi}$ are both in the spectrum. Additionally, it was shown that all eigenvalues satisfy $|\lambda_i| \leq 1$ [76].

Unfortunately, the condition that $\mathcal{C}$ needs to be completely positive is not immediately apparent in the PTM representation, as is the case for the Kraus representation. Nevertheless, many ubiquitous channels in quantum information are particularly simple when written in terms of their action on Pauli matrices. Prime examples are elements of the Clifford group Cl_n , which is the normalizer of the Pauli group $\mathcal{P}_n$ under the familiar unitary representation: $\omega(g)(\sigma_a) = g\sigma_a g^\dagger \propto \sigma_{a'}$ for $g \in \text{Cl}_n$ and some $a' \in \mathbb{F}_2^{2n}$. Therefore Clifford group elements have a sparse PTM matrix: Each row contains only one non-zero entry, which is in $\{-1, +1\}$. Elements of the Clifford group are thus signed permutation matrices in the PTM representation.

We will now give a brief overview over noise models that are frequently encountered in the literature. These provide us with a reference for the analysis of characterization results of quantum processes.

Common noise processes

Unitary errors, also known as coherent errors, are unwanted unitary processes occurring during the implementation of a gate. They can range from simple overrotations, where instead of $U = e^{i\theta H}$, $\tilde{U} = e^{i(\theta+\delta)H}$ is applied, to more complex coherent errors where unwanted terms in the Hamiltonian are present: $\tilde{U} = e^{i\theta(H+\Delta)}$. Unitary errors have a comparatively sparse representation, since they are only given by at most d^2 independent real parameters. Moreover, they are often easier to understand in terms of a physical model of the device, and can hence be reduced via calibration.

Pauli channels are defined as channels that can be written as

$$C_P(\rho) = \sum_a p_a \sigma_a \rho \sigma_a, \quad (2.62)$$

where $p_a \geq 0$ and $\sum_a p_a = 1$. From this definition we can gather the PTM coefficients:

$$(C_P)_{a,a'} = \sum_{b \in \mathbb{F}_2^{2n}} p_b \text{Tr}[\sigma_a \sigma_b \sigma_{a'} \sigma_b] = \sum_{b \in \mathbb{F}_2^{2n}} p_b (-1)^{\langle a', b \rangle} \text{Tr}[\sigma_a \sigma'_a] = \sum_{b \in \mathbb{F}_2^{2n}} p_b (-1)^{\langle a', b \rangle} \delta_{a,a'}. \quad (2.63)$$

The matrix C_P is thus diagonal with entries in the interval $[-1, 1]$. The diagonal elements of any diagonal matrix in the PTM representation are also called Pauli eigenvalues, since they satisfy the eigenvalue equation $C_P|\sigma_a\rangle = (C_P)_{a,a'}|\sigma_a\rangle$. The transformation $W : \mathbb{R}^{d^2-1} \mapsto \mathbb{R}^{d^2-1}$ that maps the Pauli probabilities to the Pauli eigenvalues is called the *Walsh-Hadamard* transform, given by

$$W = \sum_{a,a'} (-1)^{\langle a,a' \rangle} |\sigma_a\rangle \langle \sigma'_a| \quad \text{with} \quad W^{-1} = \frac{1}{d^2-1} \sum_{a,a'} (-1)^{\langle a,a' \rangle} |\sigma_a\rangle \langle \sigma'_a|. \quad (2.64)$$

Note that even though all Pauli channels have a diagonal PTM representation, the converse is not true, since the conditions $W^{-1} \text{diag}(C_P) \geq 0$ and $\sum_a (a|W^{-1} \text{diag}(C_P) = 1$ need to be additionally satisfied. Pauli channels are thus a very restricted subset of possible noise processes. They are nevertheless a useful model since they encompass common noise channels such as depolarizing, amplitude damping and dephasing noise, are more efficiently learnable than general noise [77], and are sufficiently general for quantum error correction [78]. Moreover, general noise in an experiment can be reduced to Pauli noise via randomized compiling [79–81].

In the following we give a brief description of depolarizing, amplitude damping and dephasing noise. Additional details can be found in [73]. The most prominent of the three is the depolarizing noise, which models the loss of information about the system. The associated noise channel acts as $\mathcal{D}_p(\rho) := (1-p)\rho + \frac{p}{d}\mathbb{1}$, describing a process where with probability p , the state is replaced by the completely mixed state. Using the twirl over the Pauli group given in Eq. (2.4), we can give the following description of the depolarizing channel:

$$\begin{aligned} \mathcal{D}_p(\rho) &= (1-p)\rho + p \frac{1}{|\mathcal{P}_n|} \sum_{U \in \mathcal{P}_n} U \rho U^\dagger \\ &= (1-p)\rho + \frac{p}{4^n} \sum_{a \in \mathbb{F}_2^{2n}} \sigma_a \rho \sigma_a \\ &= (1-p(1-4^{-n}))\rho + \frac{p}{4^n} \sum_{a \neq 0} \sigma_a \rho \sigma_a. \end{aligned} \quad (2.65)$$

The channels is thus represented by the d^2 Kraus operators $\{\sqrt{c_a} \sigma_a\}_{a \in \mathbb{F}_2^{2n}}$ with $c_0 = 1-p(1-4^{-n})$ and $c_a = p4^{-n}$ for $a \neq 0$.

Energy loss from a qubit (such as spontaneous emission of a photon) can be modeled by what is called an *amplitude damping* channel $\mathcal{E}_{\text{ad}}$, given by the Kraus operators

$$K_0 = |0\rangle\langle 0| + \sqrt{1-\lambda}|1\rangle\langle 1|, \quad K_1 = \sqrt{\lambda}|0\rangle\langle 1|. \quad (2.66)$$

The parameter $\lambda \in [0, 1]$ can be interpreted as the probability for the occurrence of an energy loss event that is described by the action of K_1 , which takes the excited state to the ground state. One may also consider amplitude damping as the result of more general open system dynamics [82], where Eq. (2.66) describes a noise process on the system resulting from the presence of memory effects in the environment, in which case λ is not necessarily confined to the interval $[0, 1]$. The action of the amplitude damping channel on a single qubit density matrix is given by

$$\mathcal{E}_{\text{ad}} \left(\begin{pmatrix} a & b \\ b^* & 1-a \end{pmatrix} \right) = \begin{pmatrix} 1 - \frac{(1-a)(1-\lambda)}{b^* \sqrt{1-\lambda}} & b \sqrt{1-\lambda} \\ (1-a)(1-\lambda) & \end{pmatrix}, \quad (2.67)$$

$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1-p & 0 & 0 \\ 0 & 0 & 1-p & 0 \\ 0 & 0 & 0 & 1-p \end{pmatrix}$	$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1-2p & 0 \\ 0 & 0 & 0 & 1-2p \end{pmatrix}$	$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1-2p & 0 & 0 \\ 0 & 0 & 1-2p & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$
Depolarizing	X-dephasing/Bit-flip	Z-dephasing/Phase-flip
$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \sqrt{1-\lambda} & 0 & 0 \\ 0 & 0 & \sqrt{1-\lambda} & 0 \\ \lambda & 0 & 0 & 1-\lambda \end{pmatrix}$	$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & \cos(\delta) & -\sin(\delta) \\ 0 & 0 & \sin(\delta) & \cos(\delta) \end{pmatrix}$	$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos(\delta) & -\sin(\delta) & 0 \\ 0 & \sin(\delta) & \cos(\delta) & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$
Amplitude damping	X - Rotation	Z - Rotation

Table 2.1: Examples for common single qubit noise processes in the PTM representation.

where $a \in \mathbb{R}, b \in \mathbb{C}$. This channel affects both diagonal and off-diagonal elements of the density matrix and incorporates loss of amplitude and phase information.

A different process that only affects the off-diagonal 'phase' information is given by the *dephasing channel* $\mathcal{E}_{\text{ph}}$ with Kraus operators

$$K_0 = \sqrt{1-p} \mathbb{1}, \quad K_1 = \sqrt{p} \sigma_z, \quad (2.68)$$

for $p \in [0, 1]$. This can occur when the $|0\rangle$ and $|1\rangle$ states accumulate different phases from interactions with the environment. As a result, we lose information about the relative phase, since it is perturbed according to an unknown process.

A different parametrization of both amplitude damping and dephasing processes is defined through the density matrix evolution

$$\begin{pmatrix} a_\infty + (a - a_\infty) e^{-t/T_1} & b e^{-t/T_2} \\ b^* e^{-t/T_2} & (1 - a_\infty)(a_\infty - a) e^{-t/T_1} \end{pmatrix}, \quad (2.69)$$

where the parameters are called $T_1(T_2)$ relaxation or $T_1(T_2)$ error. The amplitude damping channel Eq. (2.67) contains both T_1 and T_2 errors with $\sqrt{1-\lambda} = e^{t/T_2(\mathcal{E}_{\text{ad}})}$ and $T_2(\mathcal{E}_{\text{ad}}) = T_1(\mathcal{E}_{\text{ad}})/2$. The dephasing channel only contains T_2 errors, with $p = (1 - e^{-t/T_2(\mathcal{E}_{\text{ph}})})/2$.

In Table 2.1 the PTM representations of common single qubit coherent and incoherent noise processes are shown. These matrices will be useful for the interpretation of characterization results in quantum experiments (Section 3.2).

Representing quantum channels via their generators

Single qubit unitaries are routinely associated with rotations on the Bloch sphere, exploiting $\text{SU}(2)/\mathbb{Z}_2 \simeq \text{SO}(3)$. Given a $U \in U(2)$, we can uniquely write it as $U = e^{iH}$, with the Hamiltonian $H = \sum_{a \in \mathbb{F}_2^2} h_a \sigma_a$, $h_a \in \mathbb{R}$. Since we do not care about global phases, we can set $h_0 = 1$ and are left with a vector $\mathbf{h} := (h_{01}, h_{10}, h_{11})$, where $\mathbf{h}/\|\mathbf{h}\|_p$ is the rotation axis on the Bloch sphere and $2\|\mathbf{h}\|_p$ corresponds to the angle of rotation around this axis. Going from matrix entries to weights of Paulis in the Hamiltonian has the advantage that the latter parametrization is easier to interpret in terms of the underlying physics of the experiment. One can go a step further and describe channels in terms of generators, setting $\mathcal{C} = e^{\mathcal{L}}$. To ensure that $\mathcal{C}$ is CPTP, $\mathcal{L}$ can be chosen as a Lindbladian parametrized operator given by

$$\mathcal{L}(\rho) = i[\rho, H] + \sum_{i,j} \beta_{ij} \left(B_i \rho B_j - \frac{1}{2} (B_j B_i \rho + \rho B_j B_i) \right), \quad (2.70)$$

with a Hamiltonian H and a hermitian and orthonormal basis $\{B_i\}_{i=1}^{d^2}$, such as the normalized Pauli basis. In this form, coherent errors in terms of a mismatch between the implemented and the target Hamiltonian are separated from incoherent errors caused by dissipative processes. A caveat of this parametrization is that it only covers a subset of CPTP maps [83]. For example, all channels of this form are necessarily divisible, which does not have to be the case for arbitrary quantum channels.

2.2.2 Distance measures

For the purpose of the characterization and benchmarking of quantum systems, distance measures are required to assess discrepancies between outcome probabilities of measurements, as well as directly between the building blocks of mathematical model for quantum states, quantum channels and POVMs. The right choice of distances measure for a given task is crucial, since objects which are close with respect to one might be far apart with respect to another (as can be seen when comparing average case and worst case measures). Often but not exclusively, distance measures are induced by norms, i.e.

$\text{dist}_\alpha(A, B) = \|A - B\|_\alpha$ for some choice of norm $\|\cdot\|_\alpha$. We will now start by defining the most relevant distances for our purposes on probability vectors.

Distances between discrete probability distributions

Let $p, q \in [0, 1]^n$ with $\sum_i p_i = \sum_i q_i = 1$. The *total variation distance* (or trace distance/ ℓ_1 distance) $\delta d(p, q)$ is defined via the vector ℓ_1 -norm as

$$\delta d(p, q) := \frac{1}{2} \|p - q\|_{\ell_1} = \max_{S \subset [n]} \left| \sum_{i \in S} p_i - q_i \right|. \quad (2.71)$$

The latter form in the above equation gives an operational interpretation: The maximal discrepancy between probabilities over all events S . Another norm-induced distance is the mean squared distance or *mean squared error*

$$\mathcal{L}(p, q) = \frac{1}{n} \|p - q\|_{\ell_2}. \quad (2.72)$$

This is the preferred distance to minimize in optimization problems since, as opposed to the total variation distance, it is differentiable with respect to p, q and strictly convex. The latter property ensures that for a given p and a convex set $C \subset [0, 1]^n$, there is only a single $q \in C$ that minimizes $\mathcal{L}(p, q)$.

Another distance measure on probability distributions that is relevant to us is the *Kullback-Leibler divergence*

$$D_{\text{KL}}(p||q) = \sum_i p_i \log \left(\frac{p_i}{q_i} \right). \quad (2.73)$$

This is strictly speaking not a distance, since it does not satisfy the triangle inequality and is not symmetric in its arguments. For our purposes we will set p to be the empirically determined probability distribution of measurement data, i.e.

$p_i = k_i/n$, where k_i is the number of times outcome i is observed out of n total measurements. If we set $q(\theta)$ to be the probability distribution determined by a mathematical model with internal parameters θ , then the *likelihood function*

$$L(\{k_i\}||q(\theta)) = \prod_{i \in [n]} q_i^{k_i} \quad (2.74)$$

gives the probability of the event where each outcome i was observed exactly k_i times. It is thus natural to maximize $L(\{k_i\}||q(\theta))$ over θ to obtain the model that is most likely to have produced

the data. Since the logarithm is concave, we can just as well maximize $\log L(\{k_i\}||q(\theta)) = \sum_i k_i \log(q_i) = n \sum_i p_i \log(q_i(\theta))$, which is easier to handle. The latter is called the *log-likelihood* function. It is equivalent to the Kullback-Leibler divergence up to constants, since

$$D_{\text{KL}}(p||q(\theta)) = \sum_i p_i \log(p_i) - \sum_i p_i \log(q_i(\theta)) = c_1 \cdot \log L(p||q(\theta)) + c_2, \quad (2.75)$$

with c_1 and c_2 independent of θ . A strong argument for using the Kullback-Leibler divergence is given by the Neyman-Pearson lemma. It states that the log-likelihood ratio $\log(p_i/q_i)$ is the key figure to distinguish between the distributions p, q given an observation i , in the sense that it minimizes the type II error for a fixed type I error in a hypothesis test. The Kullback-Leibler divergence $D_{\text{KL}}(p||q)$ is then the expected value of $\log(p_i/q_i)$, assuming the observation was drawn from the distribution p .

Distances between quantum states

Distinguishing quantum states is usually done via one of the three measures given below. Let $\mathcal{S}(H) = \{\sigma \in L(\mathcal{H}) : \sigma \succeq 0, \text{Tr}[\sigma] = 1\}$ be again the state space and let $\rho, \sigma \in \mathcal{S}(H)$. We define

$$D_{\text{Tr}}(\rho, \sigma) := \frac{1}{2} \|\rho - \sigma\|_1 \quad \text{and} \quad F(\rho, \sigma) = \|\sqrt{\rho}\sqrt{\sigma}\|_1^2, \quad (2.76)$$

where $D_{\text{Tr}}(\rho, \sigma)$ is called the *trace distance* and $F(\rho, \sigma)$ the *fidelity*. Since $F(\rho, \rho) = 1$ while $D_{\text{Tr}}(\rho, \rho) = 0$, often the *infidelity* $1 - F(\rho, \sigma)$ is given instead. If one state is pure, the fidelity can be simplified: $F(|\psi\rangle\langle\psi|, \sigma) = \langle\psi|\sigma|\psi\rangle$. In quantum state tomography [84–87], the Frobenius norm $\|\rho - \sigma\|_F$ is also commonly used to distinguish states, mainly for its differentiability and the fact that there is a closed form expression [88] for $\arg\min_{\rho \in \mathcal{S}(H)} \|X - \rho\|_F$ with $X \in \text{Herm}(n)$.

Similar to the total variation distance between probability distributions, the trace distance admits the following alternative form

$$D_{\text{Tr}}(\rho, \sigma) = \sup_{0 \leq E \leq 1} \text{Tr}[E(\rho - \sigma)]. \quad (2.77)$$

It is thus the highest probability of distinguishing ρ and σ via any POVM element E , which gives it an operational interpretation and makes it the measure of choice in most of the literature. The trace distance is related to fidelity and Frobenius norm via the inequalities

$$\begin{aligned} 1 - \sqrt{F(\rho, \sigma)} &\leq D_{\text{Tr}}(\rho, \sigma) \leq \sqrt{1 - F(\rho, \sigma)}, \\ \|\rho - \sigma\|_F &\leq D_{\text{Tr}}(\rho, \sigma) \leq \sqrt{r} \|\rho - \sigma\|_F, \end{aligned} \quad (2.78)$$

where r is the rank of $\rho - \sigma$. Note that the rank can scale with the dimension in the worst case, while inequalities relating the trace distance and average gate fidelity contain no large factors. Knowledge of the fidelity thus gives a usable upper bound on the trace distance even in high dimensions, although the square root makes bounding the trace distance via the fidelity suboptimal for mixed states.

Distances between quantum channels

The most widely reported distance measure between quantum channels $\mathcal{C}, \mathcal{U}$, where $\mathcal{U}$ is unitary and $\mathcal{C}$ arbitrary, is the *average gate fidelity* defined as

$$F_{\text{avg}}(\mathcal{U}, \mathcal{C}) = \int \text{Tr}[\mathcal{U}(|\psi\rangle\langle\psi|)\mathcal{C}(|\psi\rangle\langle\psi|)] d\mu(\psi), \quad (2.79)$$

with $d\mu(\psi)$ being the unitary invariant measure on state vectors. As the name suggests, the average gate fidelity is the average over all quantum states of the fidelity between the pure

state $\mathcal{U}(|\psi\rangle\langle\psi|)$ and the mixed state $\mathcal{C}(|\psi\rangle\langle\psi|)$, where typically $\mathcal{U}$ is the target gate and $\mathcal{C}$ a noisy implementation. In analogy to the fidelity for states, sometimes the *average error rate* $r(\mathcal{U}, \mathcal{C}) = 1 - F_{\text{avg}}(\mathcal{U}, \mathcal{C})$ is given. Let $\mathcal{J}(\mathcal{U})$ and $\mathcal{J}(\mathcal{C})$ be the Choi states of $\mathcal{U}$ and $\mathcal{C}$, respectively. Then the average gate fidelity is related to the *entanglement fidelity* $F_e(\mathcal{U}, \mathcal{C}) := F(\mathcal{J}(\mathcal{U}), \mathcal{J}(\mathcal{C})) = \text{Tr}[\mathcal{U}^\dagger \mathcal{C}]/d^2$ (see [27, 89]) via

$$F_{\text{avg}}(\mathcal{U}, \mathcal{C}) = \frac{d F_e(\mathcal{U}, \mathcal{C}) + 1}{d + 1}. \quad (2.80)$$

The main selling point for the use of the average gate fidelity as a gate quality measure is that under some assumptions [32, 90], it is linked to the decay parameter in randomized benchmarking [31, 52, 53], which can be efficiently estimated.

A more stringent error measure is defined via the diamond norm

$$\|\mathcal{C}\|_\diamond := \sup_{\rho \in \mathcal{S}(\mathcal{H} \otimes \mathcal{H})} \|(\text{id} \otimes \mathcal{C})(\rho)\|_1 = \sup_{0 \leq E \leq \mathbb{1}} \sup_{\rho \in \mathcal{S}(\mathcal{H} \otimes \mathcal{H})} \text{Tr}[E(\text{id} \otimes \mathcal{C})(\rho)], \quad (2.81)$$

where $E \in \mathcal{L}(\mathcal{H} \otimes \mathcal{H})$. The *diamond distance* between channels $\mathcal{C}$ and $\tilde{\mathcal{C}}$ is then defined as $\frac{1}{2}\|\mathcal{C} - \tilde{\mathcal{C}}\|_\diamond$, where the factor $1/2$ ensures that it lies in the interval $[0, 1]$. The operational interpretation is straightforward: The diamond distance gives the worst case discrepancy in outcome probability over all states and all POVM-elements between two gates acting on a subsystem. For this reason it has traditionally been used to provide provable error thresholds, below which quantum error correcting codes achieve fault tolerance [91]. To date there is no method known that can efficiently estimate the diamond distance between a target gate $\mathcal{U}$ and its implementation. Unfortunately the average gate fidelity does not give a useful upper bound on the diamond distance either, since without additional assumptions, no better bound than

$$\frac{d+1}{d}r(\mathcal{U}, \mathcal{C}) \leq \frac{1}{2}\|\mathcal{C} - \tilde{\mathcal{C}}\|_\diamond \leq \sqrt{d(d+1)r(\mathcal{U}, \mathcal{C})} \quad (2.82)$$

is known. It is therefore necessary to acquire a full description of $\mathcal{C}$ via process or gate set tomography first, whereafter the diamond distance can be computed via a semidefinite program [92]. Interestingly, as Kueng *et al.* [71] have shown, purely unitary errors exhibit the worst case scaling and no better bound than in Eq. (2.82) can be obtained there. It was further shown in Wallman [93] that if we define the worst case infidelity

$$r_{\text{max}} := \max_{|\psi\rangle\langle\psi| \in \mathcal{S}(H)} (1 - \langle\psi|\mathcal{C}(|\psi\rangle\langle\psi|)|\psi\rangle), \quad (2.83)$$

then it holds that $r(\mathcal{C}) \leq r_{\text{max}} \leq (d+1)r(\mathcal{C})$ and the bound is tight since there exist channels $\mathcal{C}$ for which $r_{\text{max}} = \mathcal{O}(dr(\mathcal{C}))$. This implies that the dimensional factor in the upper bound of Eq. (2.82) is due to the worst-case to average-case conversion, while only the square root scaling with $\sqrt{r(\mathcal{C})}$ is due to the different distance measure. However, for non-unitary (incoherent) noise, a better scaling of the diamond norm bound in the infidelity can be found. First we need to define the unitarity introduced in Wallman *et al.* [94], which is an efficiently estimable quantity that measures how 'unitary' a process is. Let $\mathcal{C}'$ be a linear map that satisfies $\mathcal{C}'(\mathbb{1}) = 0$ and $\mathcal{C}'(X) = \mathcal{C}(X) - \text{Tr}[\mathcal{C}(X)]\mathbb{1}/\sqrt{d}$ for all traceless X . Then the unitarity is defined to be the quantity

$$u(\mathcal{C}) = \frac{d}{d-1} \int \text{Tr} \left[\mathcal{C}'(|\psi\rangle\langle\psi|)^\dagger \mathcal{C}'(|\psi\rangle\langle\psi|) \right] d\mu(\psi), \quad (2.84)$$

which satisfies $u(\mathcal{U}) = 1$ for any unitary channel $\mathcal{U}$. It was shown in [71] that if $\mathcal{C}$ is unital and its unitarity scales as $u(\mathcal{C}) = (1 - \frac{dr}{d-1})^2 + \mathcal{O}(r^2)$, then $\frac{1}{2}\|\mathcal{C} - \tilde{\mathcal{C}}\|_\diamond = \mathcal{O}(d^2r)$. This bound is only useful for small systems, where the factor d^2 is less important than the quadratic improvement in r . For example if $d = 2$ and $r = 10^{-4}$, which is an average error rate routinely achieved in current experiments. In Wallman [93], an alternative bound of the diamond distance in terms of unitarity and infidelity is also given.

For Pauli channels $\mathcal{C}_p$ (see Eq. (2.62)) the diamond norm and average gate fidelity are equivalent [95]:

$$\|\mathcal{C}_P\|_\diamond = (1 + 1/d) r(\text{id}, \mathcal{C}_P) = \|e_1 - p\|_{\ell_1}, \quad (2.85)$$

where p is the vector of Pauli error probabilities and e_1 the first unit vector.

As will be discussed in Section 2.2.4, sometimes channels are only given up to similarity transformations TCT^{-1} for $B \in \mathbb{C}^{d^2 \times d^2}$. It would therefore be a desirable property of a given distance measure to be invariant under similarity transformations in one of its arguments, i.e. $\text{dist}(\mathcal{U}, \tilde{\mathcal{C}}) = \text{dist}(\mathcal{U}, T\tilde{\mathcal{C}}T^{-1})$. Unfortunately this is not generally satisfied by the diamond distance or the average error rate. One exception is the average error rate to the identity channel, which is computed via the entanglement fidelity (Eq. (2.80)), satisfying $F_e(\text{id}, \mathcal{C}) = \text{Tr}[\mathbb{1} \mathcal{C}]/d^2 = \text{Tr}[TCT^{-1}]/d^2$. One way to get around the problem is to define distance measures that only depend on the eigenvalues of $\mathcal{U}$ and $\mathcal{C}$, since those are invariant under similarity transformations. One example of such an error measure is what we call the *spectral 1-distance*, defined by

$$D_{s1}(\mathcal{C}, \tilde{\mathcal{C}}) := \frac{1}{d^2} \min_{\pi \in \text{Perm}(d^2)} \sum_{i=1}^{d^2} |\lambda_i(\mathcal{C}) - \lambda_{\pi(i)}(\tilde{\mathcal{C}})|, \quad (2.86)$$

where $\text{Perm}(d^2)$ is the set of Permutations of d^2 numbers. Computing the above distance can be phrased as an instance of a linear assignment problem, for which there exist polynomial time algorithms in the problem size d^2 . It is thus far not clear if the spectral 1-distance admits an operational interpretation in general. However, for the special case where $\mathcal{C} = \text{id}$ and $\tilde{\mathcal{C}}$ is a Pauli channel with Pauli eigenvalues $\lambda_i(\tilde{\mathcal{C}}) \in [-1, 1]$, we get

$$D_{s1}(\text{id}, \tilde{\mathcal{C}}) = \frac{1}{d^2} \sum_{i=1}^{d^2} (1 - \lambda_i(\tilde{\mathcal{C}})) = 1 - p_{\mathbb{1}}(\tilde{\mathcal{C}}), \quad (2.87)$$

where $p_{\mathbb{1}}(\tilde{\mathcal{C}})$ is the probability of $\tilde{\mathcal{C}}$ acting as the identity.

The distances between quantum channels in process tomography protocols are also routinely measured via distances on the corresponding Choi states, defined via either the Frobenius norm $\|\mathcal{J}(\mathcal{C}) - \mathcal{J}(\tilde{\mathcal{C}})\|_F$ or the trace norm $\|\mathcal{J}(\mathcal{C}) - \mathcal{J}(\tilde{\mathcal{C}})\|_1$.

Distances between POVMs

Let $E = \{E_i\}_{i=1}^k$ and $\tilde{E} = \{\tilde{E}_i\}_{i=1}^k$ be two k -outcome POVMs. We can define the analogous distance to the trace distance for quantum states via the so-called *operational distance* [96, 97]

$$D_{\text{op}}(E, \tilde{E}) := \max_{\rho \in \mathcal{S}(H)} \delta d(p(\rho, E), p(\rho, \tilde{E})), \quad (2.88)$$

where $p(\rho, E) = (\text{Tr}[\rho E_1], \dots, \text{Tr}[\rho E_k])$ is the vector of outcome probabilities of E given ρ . Since it is defined via the total variation distance, the operational distance can also be formulated as the maximum difference in probabilities of two events occurring, over all states. The operational distance can be computed [97] via

$$D_{\text{op}}(E, \tilde{E}) = \max_{I \subseteq [n]} \left\| \sum_{i \in I} (E_i - \tilde{E}_i) \right\|_\infty. \quad (2.89)$$

Another way to define distance measures on POVMs is to consider them as measure and prepare channels $\mathcal{E} : \mathcal{S}(H) \mapsto \mathbb{C}^k : \mathcal{E}(\rho) = \sum_{i=1}^k \text{Tr}[E_i \rho] |i\rangle\langle i|$. Any distance measures defined for channels can thus be defined for POVMs, in particular the diamond distance

$$D_\diamond(E, \tilde{E}) = \frac{1}{2} \|\mathcal{E} - \tilde{\mathcal{E}}\|_\diamond. \quad (2.90)$$

Since the output state of the prepare and measure channel is essentially a classical state, we can see that the operational distance is equivalent to the half $1 \rightarrow 1$ norm defined as

$$\frac{1}{2} \|\mathcal{E} - \tilde{\mathcal{E}}\|_{1 \rightarrow 1} = \frac{1}{2} \max_{\rho \in \mathcal{S}(\mathcal{H})} \|(\mathcal{E} - \tilde{\mathcal{E}})(\rho)\|_1 = \max_{\rho \in \mathcal{S}(\mathcal{H})} \delta d(p(\rho, E), p(\rho, \tilde{E})). \quad (2.91)$$

The diamond distance between POVMs, in contrast, treats the measurement as an action on a larger Hilbert space and thus allows for more powerful protocols to distinguish E and $\tilde{E}$ via initial states in $\mathcal{S}(\mathcal{H}^{\otimes 2})$ and entanglement between the original system and ancillas.

2.2.3 Process tomography

Quantum process tomography (QPT) considers the problem of reconstructing a mathematical model of a quantum process from measurement data. Since a real quantum system is never perfectly isolated from the environment, a model is naturally just an approximation to the real dynamics in the system. Loosely speaking, every scheme that extracts a model which can be used to predict measurement outcomes for different initial conditions can be considered QPT. This distinguishes between tomography and certification/benchmarking, where methods in the latter group aim to extract performance metrics which can not be used to make fine-grained predictions.

In this section we aim to give a (non-exhaustive) overview over different approaches to QPT, which differ in the mathematical models assumed to describe the dynamics, as well as in the strategies to reconstruct said models. The variant which is central to the work presented in this thesis, gate set tomography, is then explored in more detail in section Section 2.2.4.

The principal setting of QPT is that a single process is assumed to be unknown, while arbitrary and fully known input states and measurements on the system can be realized. This is sometimes extended to include input states and measurements on a larger system including ancillas. Assumptions on the process are that it is (i) time stationary, (ii) context independent and (iii) described by a quantum channel acting only on the system Hilbert space. These conditions are essential for quantum tomography schemes, since outputs are probabilistic and many measurement settings are required, thus tomography can only function given a sequence of observations of an internally consistent process. The conditions can be relaxed by including environment interactions to the model in what is termed non-Markovian tomography, which we will summarize at the end of this section. A distinction also has to be made between coherent and incoherent measurement strategies, where the former assumes that identical copies of a channel can be applied to a larger quantum state, and that the output can be measured with measurement operators that are entangled over the Hilbert spaces of different copies. This assumes access to a large number of qubits and sufficient control to ensure the channel copies are identical, which is hard to come by in the current era of noisy intermediate sized quantum devices. We therefore focus on incoherent measurements, where only a single copy of the channel is applied per measurement round.

Standard tomography of quantum channels

We first consider the textbook case where a quantum operation is modeled by a quantum channel $\mathcal{C} : L(\mathcal{H}) \mapsto L(\mathcal{H})$ as defined in Section 2.2.1. The tomographic procedure to determine $\mathcal{C}$ from repeated cycles of state preparation, the application of $\mathcal{C}$ and subsequent state tomography of the output state was first formalized in 1997 [98, 99]. In the meantime, QPT has been applied to different platforms such as nuclear magnetic resonance, photonic, trapped ion, solid state and superconducting qubit systems [100–104]. The standard version of QPT can be summarized as follows. Let $\{B_i\}_{i=1}^{d^2}$ with $B_i \in L(\mathcal{H})$ be an orthonormal basis for the space $L(\mathcal{H})$ and let $\mathcal{C}(\rho) = \sum_l K_l \rho K_l^\dagger$ be given in Kraus representation. If we express the Kraus operators in this basis as $K_l = \sum_i \alpha_i B_i$, we can write the channel action as $\mathcal{C}(\rho) = \sum_{ij} C_{ij} B_i \rho B_j^\dagger$, where

$C_{ij} = \alpha_i \alpha_j^*$. Now suppose we can prepare the d^2 operators $\{B_i\}$ as input states to $\mathcal{C}$, then the action of $\mathcal{C}$ on these inputs can be written as

$$\mathcal{C}(B_k) = \sum_{ij} C_{ij} B_i B_k B_j^\dagger =: \sum_l \gamma_{lk} B_l. \quad (2.92)$$

Here $\sum_l \gamma_{lk} B_l$ is the output state of $\mathcal{C}$ on the input B_k , whose coefficients γ_{lk} can be procured via quantum state tomography. If we define $B_i B_k B_j^\dagger =: \beta_{ijk} B_m$, then Eq. (2.92) becomes

$$\sum_{ij} C_{ij} B_i B_k B_j^\dagger = \sum_{ijm} C_{ij} \beta_{ijk} B_m = \sum_l \gamma_{lk} B_l. \quad (2.93)$$

Since the operators B_l form an orthonormal basis, the final equation reads

$$\sum_{ij} \beta_{kmij} C_{ij} = \gamma_{km}, \quad (2.94)$$

which is a linear system of equations that in the vectorized form $\beta \text{vec}(C) = \text{vec}(\gamma)$ is just solved via $\text{vec}(C) = \beta^{-1} \text{vec}(\gamma)$. An alternative but equivalent formulation is to define a measurement map

$$\mathbf{M}_{ij}(C) := \text{Tr}[E_i \mathcal{C}(\rho_j)] = P_{ij} \quad (2.95)$$

with a set of POVM elements $\{E_i\}$ and initial states $\{\rho_j\}$. The outcomes probabilities P_{ij} are experimentally determined by repeating the experiment for a given ρ_j and E_i a number of times to collect statistics (see left side of Figure 2.4). The channel is then recovered by computing the pseudoinverse $\mathbf{M}^+$ of the linear map $\mathbf{M}$ to get

$$\hat{C} = \mathbf{M}^+(P). \quad (2.96)$$

Here the set of POVMs or the set of initial states can again be an orthonormal basis, or just any frame on $L(\mathcal{H})$.

A commonly chosen basis on $L(\mathcal{H})$ is the Pauli basis $\mathcal{P}$. Note that a Pauli σ_a with $a \in \mathbb{F}_2^{2n}$ is not a valid quantum state, but we can write it in spectral decomposition as $\sigma_a = \sum_i s_i |\psi_i\rangle\langle\psi_i|$ with eigenstates $|\psi_i\rangle$ and eigenvalues $s_i \in \{-1, 1\}$. Since everything is linear,

$$\text{Tr}(E_j \mathcal{C}(\sigma_a)) = \sum_i s_i \text{Tr}(E_j \mathcal{C}(|\psi_i\rangle\langle\psi_i|)) \quad (2.97)$$

and we get the expectation value by preparing all eigenstates of σ_a . Other commonly considered frames on $L(\mathcal{H})$ are the set of mutually unbiased bases (MUBs) [105] and the set of symmetric informationally complete POVMs (SIC-POVMs) [106], both of which constitute examples of QPT with complex projective 2-designs [107, 108].

An even simpler approach is given by what is termed *direct characterization* [109, 110]. Using rounded bracket notation (Section 2.1.1) that lets us write the identity channel as $\text{id} = \sum_{i=1}^{d^2} |B_i\rangle(B_i|$, we can express $\mathcal{C}$ in this basis as $\mathcal{C} = \sum_{i,j=1}^{d^2} (B_i|\mathcal{C}|B_j) |B_i\rangle(B_j|$. Let $C_{ij} = (B_i|\mathcal{C}|B_j)$, then assuming we can prepare a state $|B_j\rangle$ (see argument above) and measure $(B_i|$, we have direct access to the entry C_{ij} .

A different but from a theoretical standpoint equivalent method to the above protocols relies on representing $\mathcal{C}$ via its Choi-state $\mathcal{J}(\mathcal{C})$ [111, 112]. Since by Eq. (2.58) the Choi state is given by the application of $\mathcal{C}$ to one half of a maximally entangled state, it can be prepared in an experiment and subsequently measured (see right side of Figure 2.4). The disadvantage of this method is that the preparation of a maximally entangled state on a larger system is experimentally much more demanding than preparing and measuring states locally. Other approaches where the channel is applied to a state on a larger Hilbert space that not necessarily require entanglement are known as ancilla assisted QPT [113]. In these settings one can also make

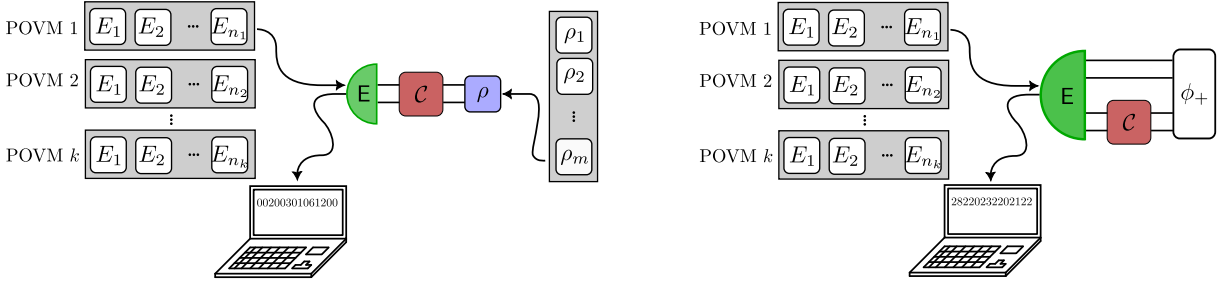


Figure 2.4: Schematic depiction of quantum process tomography with standard QPT on the left and ancilla-assisted tomography of the Choi state on the right.

the distinction between product- and entangled measurements, where the former assumes access to outcome probabilities $\text{Tr}[(E \otimes \tilde{E})\mathcal{J}(C)]$ with $E, \tilde{E} \in \mathcal{L}(\mathcal{H})$, while the latter allows access to $\text{Tr}[E\mathcal{J}(C)]$ with $E \in \mathcal{L}(\mathcal{H} \otimes \mathcal{H})$ including POVM elements that do not factorize. Note that

$$\text{Tr}(E\mathcal{C}(\rho)) = \text{Tr}[(\mathbb{1} \otimes E)\mathcal{J}(C)(\rho^T \otimes \mathbb{1})] = \text{Tr}[(\rho^T \otimes E)C], \quad (2.98)$$

which implies that state tomography on the Choi state with measurement operators that factorize is formally equivalent to QPT. The Choi Isomorphism thus enables us to apply protocols developed for quantum state tomography to process tomography (see e.g. [114, 115]).

If errors due to imperfect state preparation and measurement (SPAM) are assumed to be negligible, the main estimation error in QPT is due to what is called *shot noise*—the statistical errors incurred by determining the outcome probabilities $P_{ij} = \text{Tr}[E_i\mathcal{C}(\rho_j)]$ from finitely many repetitions (shots). The linear inversion step in Eq. (2.96) then further leads to channel estimates $\hat{C}$ not necessarily being completely positive and trace preserving. This issue is commonly resolved by either projecting the linear inversion estimate to the manifold of CPT-maps (see e.g. [116]) or via iterative optimization approaches that we will summarize in the following. Let $f(\mathcal{X}) : \mathcal{L}(\mathcal{H}) \mapsto \mathbb{R}$ be a cost function such as the least-squares error

$$f(\mathcal{X}) = \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} (\text{Tr}[E_i\mathcal{X}(\rho_j)] - P_{ij})^2. \quad (2.99)$$

Then a physical channel estimate is given by the solution to the following optimization problem

$$\begin{aligned} & \underset{\mathcal{X}}{\text{minimize}} && f(\mathcal{X}) \\ & \text{subject to} && \mathcal{X} \text{ CPT} \Leftrightarrow \mathcal{J}(\mathcal{X}) \succeq 0, \text{Tr}(\mathcal{J}(\mathcal{X})) = 1. \end{aligned} \quad (2.100)$$

In Knee *et al.* [117] this problem was solved via an optimization algorithm of the log-likelihood cost function (Eq. (2.75)) that alternates between gradient descent updates and projections onto the set of CPT-maps. The authors further present numerical evidence that this projected gradient descent algorithm leads to final estimates with lower trace distance error on the Choi matrix compared to unconstrained gradient descent followed by a single final projection.

In Surawy-Stepney *et al.* [116] the *projected least-squares* method to QPT is presented and analyzed, generalizing a previous method for state tomography developed by some of the authors. It first solves the least-squares optimization problem of Eq. (2.99), which can be written as

$$\hat{\mathcal{X}} = \underset{\mathcal{X} \in \mathcal{L}(\mathcal{H})}{\arg \min} \|\mathbf{M}(\mathcal{X}) - P\|_F, \quad (2.101)$$

where $\mathbf{M}$ is again a measurement map as defined in Eq. (2.95). One can show that the pseudo-inverse $\hat{\mathcal{X}} = \mathbf{M}^+(P)$ minimizes the above least-squares error, and admits a closed form solution, which the authors derive for different scenarios: Direct or ancilla assisted tomography with Pauli measurements or mutually unbiased measurements/input states. The significance of the

existence of a closed form solution of the pseudoinverse is that no matrix inversion of $\mathbf{M}$ needs to be performed, which would be particularly costly since given a total of N measurement settings, $\mathbf{M} : \mathcal{L}(\mathcal{L}(\mathcal{H})) \mapsto \mathbb{R}^N$ can be prohibitively large. After linear inversion, the Choi matrix of the estimate $\hat{\mathcal{X}}$ is then projected onto the set of positive unit trace matrices. Crucially, the authors prove concrete error bounds on the Frobenius and trace norm between the estimated and the true Choi matrix. Let us assume we are given N measurements of an n -qubit channel $\mathcal{C}$, whose Choi matrix is of rank r (Kraus rank). The scaling of the total number of measurements N for a fixed error is also called the *sample complexity*. Let further $g(n)$ be a measurement strategy dependent scaling function which is $\mathcal{O}(3^{-2n})$ for local Pauli measurements and $\mathcal{O}(2^{-2n})$ for mutually unbiased bases measurements, then according to [116], Theorem 1, the error bounds are given as

$$\begin{aligned} \mathbb{P}(\|\mathcal{J}(\hat{\mathcal{X}}) - \mathcal{J}(\mathcal{C})\|_{\text{F}} \geq \epsilon) &\leq d^2 \exp\left(-\frac{3N\epsilon^2 g(n)}{64r}\right), \\ \mathbb{P}(\|\mathcal{J}(\hat{\mathcal{X}}) - \mathcal{J}(\mathcal{C})\|_1 \geq \epsilon) &\leq d^2 \exp\left(-\frac{3N\epsilon^2 g(n)}{256r^2}\right). \end{aligned} \quad (2.102)$$

For an error ϵ achieved with probability $1 - \delta$, this results in a sample complexity of

$$N = \mathcal{O}\left(\frac{r \log(d^2/\delta)}{g(k)\epsilon^2}\right) \quad \text{in } \|\cdot\|_{\text{F}} \quad \text{and} \quad N = \mathcal{O}\left(\frac{r^2 \log(d^2/\delta)}{g(k)\epsilon^2}\right) \quad \text{in } \|\cdot\|_1. \quad (2.103)$$

Since the rank r of the Choi matrix can be up to d^2 in general, the achieved scaling in r gives a massive reduction in sample complexity for low rank channels. This even extends to channels which are only approximately of low rank, as further shown in Surawy-Stepney *et al.* [116]. It can moreover be argued that the error bounds for mutually unbiased bases are essentially optimal, since they scale linearly in the number of free parameters of $\mathcal{C}$, which is given by rd^2 .

An inconvenience about the bounds in Eq. (2.102) is that they are phrased in terms of norms on the Choi matrix, which are not as well-motivated from an operational point of view as the diamond distance. In a follow-up work by Oufkir [118] it was shown via the norm inequality $\|\mathcal{X} - \mathcal{C}\|_{\diamond} \leq d^2 \|\mathcal{J}(\mathcal{X}) - \mathcal{J}(\mathcal{C})\|_{\infty}$ that $N = \mathcal{O}(d^6 \log(d^2)/\epsilon^2)$ measurements suffice to reconstruct a channel $\mathcal{C}$ within error ϵ in diamond norm. Even more importantly, the author also proved a matching lower bound, which states that for $\epsilon \leq 1/16$ and $d \geq 4$, $N = \Omega(d^6/\epsilon^2)$ measurements are required for reconstruction to error ϵ in diamond norm.

In the remainder of this section we want to give an overview of methods that, as opposed to full tomography, reconstruct lower parameter count models which aim to reduce the measurement overhead, while still providing a reasonably descriptive model of the dynamics.

Compressed sensing process tomography

We previously saw that for error bounds on the projected least-squares method, a low Kraus rank r allows for substantially better bounds, even when the low rank is not explicitly enforced in the estimate $\hat{\mathcal{X}}$. Using the prior information that a given channel is of at most rank r and thereby reducing the model complexity falls under the umbrella of *compressed sensing* (CS) [119]. CS-methods were originally developed for classical signal processing tasks, using the sparsity of a signal in some basis to significantly reduce the sampling cost compared to lower bounds for the general case. Note that for unitaries it holds that the Choi matrix is of unit rank, thus in the eigenbasis of its Choi matrix, any unitary is sparse. Since quantum channels tend to aim for the implementation of a unitary map, we can usually assume that if the implementation error is low, the resulting channel is also sparse. In the following, we call a tensor s -sparse, if at most s of its entries are nonzero.

The first account of using compressed sensing for QPT was not framed for low Kraus rank channels, but sparsity in a known basis by Kosut [120]. This could for instance be the Pauli basis in which Clifford gates are sparse, since they contain only one nonzero entry per row. This

approach can be formalized as follows. Let channels $\mathcal{X}$ be parametrized in the basis where we assume our target channel to be sparse, then the compressed sensing formulation of the QPT estimation problem in Eq. (2.100) is given by

$$\hat{\mathcal{X}} = \arg \min_{\mathcal{X} \in CPT} \|\mathcal{X}\|_{\ell_1} \quad \text{subject to} \quad \|\mathbf{M}(\mathcal{X}) - P\|_F \leq \epsilon. \quad (2.104)$$

Here the ℓ_1 -norm minimization leads to a sparse solution, since the ℓ_1 -norm is the convex relaxation of the sparsity measure $\sum_{ij} \mathcal{X}_{ij}^0$, i.e. the number of non-zero elements of $\mathcal{X}$. This approach was later equipped with performance guarantees [121] via a common technique in compressed sensing, which relies on the *restricted isometry property* (RIP) of the measurement map $\mathbf{M}$ that we define in the following. We say $\mathbf{M}$ satisfies the RIP with isometry constant δ_s if

$$1 - \delta_s \leq \frac{\|\mathbf{M}(\mathcal{X}_1) - \mathbf{M}(\mathcal{X}_2)\|_F^2}{\|\mathcal{X}_1 - \mathcal{X}_2\|_F^2} \leq 1 + \delta_s \quad (2.105)$$

for all channels $\mathcal{X}_1, \mathcal{X}_2$ that are s -sparse. Intuitively this condition ensures that for channels that are close in Frobenius norm, their measurements with respect to $\mathbf{M}$ are also close in Frobenius norm. Once a RIP is proven for a given constant δ_s , standard bounds are applied in Shabani *et al.* [121] to show that only $\mathcal{O}(s \log(d^4/s))$ measurement settings are required to recover an s -sparse channel $\mathcal{C}$ in Frobenius norm. This also extends to approximately s -sparse channels in the sense that the additional error incurred is quantifiable as

$$\|\hat{\mathcal{X}} - \mathcal{C}\|_F \leq \frac{c_1}{\sqrt{s}} \|\mathcal{C}_s - \mathcal{C}\|_{\ell_1} + c_2 \epsilon, \quad (2.106)$$

where $\mathcal{C}_s$ is the best s -sparse approximation to $\mathcal{C}$, c_1, c_2 are constants and ϵ is the shot noise error. The main disadvantage to this approach via sparse matrices is that the basis in which the channel is sparse has to be known in advance. Additionally, a measurement operator has to be constructed that satisfies the RIP with respect to this basis. Even if such a measurement operator can be constructed, it is not clear whether the prescribed measurements are easily implementable in practice.

This problem was remedied by concurrent works on quantum state tomography via compressed sensing that also apply to process tomography [122–124]. Here the concept of sparsity is applied more naturally to the quantum setting, by considering low rank density matrices and low rank Choi matrices. The recovery guarantees use a different variant of RIP, where the isometry constant δ_r is rank-dependent, and it holds that

$$1 - \delta_r \leq \frac{\|\mathbf{M}(X)\|_F}{\|X\|_F} \leq 1 + \delta_r. \quad (2.107)$$

for all rank r matrices X . Crucially, RIP can be achieved in this setting for low Kraus rank r with easily implementable local Pauli measurements that lead to a guaranteed recovery with $\mathcal{O}(rd^2 \log(d))$ settings [124]. However, as was discussed before, local Pauli measurements of the Choi matrix which determine $\text{Tr}[\sigma_a \otimes \sigma_{a'} \mathcal{J}(\mathcal{C})] = \text{Tr}(\sigma_a \mathcal{C}(\sigma_{a'}))$ require the preparation of d eigenstates of $\sigma_{a'}$ (see Eq. (2.97)). This problem effectively introduces an additional factor of d in the number of measurement settings, meaning (not including shot noise) a total of $\mathcal{O}(rd^3 \log(d))$ settings are required in practice. It needs to be emphasized that the number of settings does not correspond to the sample complexity, since the number of required shots for each setting might also scale with d . In Flammia *et al.* [124] the sample complexity for process tomography was not analyzed, but we can extrapolate it based on the given sample complexity for quantum states. In analogy to Eq. (2.106), the incurred error due to the rank condition not being strictly satisfied can be quantified as

$$\|\hat{\rho} - \rho\|_1 \leq c \|\rho_r - \rho\|_{\ell_1} + \epsilon, \quad (2.108)$$

where again c is a constant, ϵ is due to shot noise and ρ_r is the best rank r approximation to ρ . Disregarding the model mismatch error $\|\rho_r - \rho\|_{\ell_1}$, it was further shown that the total sample complexity for learning a quantum state via Pauli measurements to error $\|\hat{\rho} - \rho\|_1 \leq \epsilon$, is $\mathcal{O}(r^2 d^2 \log(d)/\epsilon^2)$. With the extra dimensional factor from the preparation of Pauli-eigenstates, and the Choi matrix being of dimension d^2 , this would suggest a sample complexity of $\mathcal{O}(r^2 d^5 \log(d)/\epsilon^2)$ for learning a Choi-matrix in trace distance.

A direct compressed sensing method for QPT without the use of maximally entangled states was developed by Kliesch *et al.* [125]. To define the measurement setting used therein, we first define a complex projective k -design as a probability distribution ν that reproduces k th order moments of the Haar measure μ on the complex unit sphere

$$\int (|\psi\rangle\langle\psi|)^{\otimes k} d\nu(\psi) = \int (|\psi\rangle\langle\psi|)^{\otimes k} d\mu(\psi). \quad (2.109)$$

The definition is analogous to the definition of unitary k -designs in Eq. (2.13), and it can be seen that a unitary k -design induces a complex projective k -design since we can write $|\psi\rangle = U|0\rangle$ for some U and $d\nu(\psi) = d\nu(U)$. The protocol in [125] then uses random measurement settings given by pairs $(E, |\psi\rangle\langle\psi|)$ where $|\psi\rangle$ is drawn from a complex projective 4-design and $E = UE_0U^\dagger$ with U drawn from a unitary 4-design with E_0 fixed, traceless and of unit spectral norm. Apart from the standard compressed sensing approach of minimizing the trace norm (as in Eq. (2.104)), minimization of the diamond norm

$$\hat{\mathcal{X}} = \arg \min_{\mathcal{X} \in \text{HT}} \|\mathcal{X}\|_{\diamond} \quad \text{subject to} \quad \|\mathbf{M}(\mathcal{X}) - P\|_F \leq \epsilon \quad (2.110)$$

is also analyzed, where $\text{HT} \subset \text{L}(\text{L}(\mathcal{H}))$ is the space of hermiticity and trace preserving superoperators. Furthermore, a CPT-constrained least-squares problem with objective function defined in Eq. (2.101) is also considered. Interestingly, in numerical tests the CPT-constrained least-squares, HT-constrained diamond norm and HT-constrained ℓ_1 minimization reach the same high success probability $\mathbb{P}(\|\hat{\mathcal{X}} - \mathcal{C}\|_F \leq 10^{-5})$ for a low number of measurement settings. In contrast, without HT constraint, diamond norm minimization outperforms ℓ_1 norm minimization. A superior performance for the diamond norm as opposed to the trace norm for general compressed sensing problems has also been demonstrated analytically and numerically in Kliesch *et al.* [126]. Note that compressed sensing estimation problems such as the one defined in Eq. (2.110) are instances of convex optimization problems and can thus be recast as semidefinite programs and solved via standard libraries.

Using the 4-design properties it was further shown that for the measurement ensemble defined above, learning a rank r channel in Frobenius norm, i.e.

$\|\hat{\mathcal{X}} - \mathcal{C}\|_F \leq \epsilon$, can be done with a total sample complexity of $\mathcal{O}(rd^5/\epsilon^2)$ [125]. Other works on QPT in a CS-setting include [127, 128] where a focus is set on reconstructing unitary ($r = 1$) channels.

Another intriguing approach is given in Roth *et al.* [129], which, generalizing earlier results by Kimmel *et al.* [130], introduces robustness to state preparation and measurement errors to the compressed sensing QPT setting. The general idea is that access to average gate fidelities (AGFs) $\mathcal{F}(\mathcal{U}, \mathcal{C})$ for many unitaries $\mathcal{U}$, provides enough information about $\mathcal{C}$ for a reconstruction, provided $\mathcal{C}$ is unital. This approach can already be motivated by noting that any unital channel can be written as $\mathcal{C}(\rho) = \sum_i \alpha_i \mathcal{U}_i(\rho)$, where $\mathcal{U}_i$ are unitary channels and $\alpha_i \in \mathbb{R}$, $\sum_i \alpha_i = 1$ [131]. In fact, Roth *et al.* [129] show that if the set $\{\mathcal{U}_i\}_{i=1}^N$ forms a unitary 2-design, then

$$\mathcal{C}(\rho) = \frac{1}{N} \sum_i \tilde{\alpha}_i \mathcal{U}_i(\rho) \quad \text{with} \quad \tilde{\alpha}_i = d(d+1)(d^2-1)(\mathcal{F}(\mathcal{U}_i, \mathcal{C}) - 1/d) + 1 \quad (2.111)$$

for any unital channel $\mathcal{C}$ (see also Scott [132]). For unitary channels, it was shown that the required number of AGFs is essentially optimal and given by $\mathcal{O}(d^2)$. Combined with efficient protocols to estimate the average gate fidelities [43, 133], this leads to a total sample complexity

of $\mathcal{O}(d^4/\epsilon^2)$ for reconstruction in Frobenius norm, which was shown to be optimal for rank 1 measurements (such as AGFs) on the Choi state. Alternatively, the AGFs can be bounded [130] via interleaved randomized benchmarking [134, 135], which is robust to state preparation and measurement errors. The reason that the average gate fidelities can not directly be determined via RB is that interleaved RB gives access to $\mathcal{F}(\mathcal{C}\mathcal{E}, \mathcal{U})$, where $\mathcal{E}$ is the gate-independent noise channel assumed to act on Clifford gates applied in RB. Standard randomized benchmarking gives access to $\mathcal{F}(\mathcal{E}, \text{id})$ which can then be used to bound $\mathcal{F}(\mathcal{C}, \mathcal{U})$. The original bounds were later improved in Carignan-Dugas *et al.* [136]. It must also be noted that randomized benchmarking actually measures a decay parameter, which can be linked to the average gate fidelity in a specific gauge (see Section 2.2.4) [90]. This gauge, however, can not always be chosen such that the noise is described by a CPT map [32]. The AGF has further been linked to RB decay parameters by Helsen *et al.* [32] via additional parameters, which are however as of yet not experimentally accessible in an efficient manner. Another practical issue that arises with the estimation of AGFs through RB is that many of the channels $\mathcal{U}$, for which $\mathcal{F}(\mathcal{C}, \mathcal{U})$ needs to be known, have a small value of $\mathcal{F}(\mathcal{C}, \mathcal{U})$. The resulting exponential decay in the RB data is thus rapid and hard to fit.

Pauli channel tomography and spectral tomography

Pauli channels as defined in Eq. (2.62) constitute a ubiquitous and well-motivated subclass of channels, which are defined by $d^2 - 1$ real parameters as opposed to the $\mathcal{O}(d^4)$ parameters of a full quantum channel. In the work by Flammia and Wallman [77], the authors give provably sample efficient and SPAM-robust protocols for learning Pauli-noise with custom version of RB over the Pauli group. The protocol uses initial states and measurements in a stabilizer basis with varying lengths of random Paulis in between. As opposed to the two-outcome POVM of standard RB, each stabilizer measurement provides n outcome bits, where the outcome probability for each bit string is given by a sum of exponential decays. A central idea of the protocol is then the isolation of the exponential decays via filter functions, allowing for the estimation of up to 2^n decay parameters from a single measurement setting (see also Helsen *et al.* [137]). We will give an overview of the results in the following. Let $\omega(g)$ be the unitary channel given by the adjoint action of Pauli $g \in \mathcal{P}_n$ and $\tilde{\omega}(g) = \omega(g)\Lambda(g)$ be a noisy implementation. The protocol in [77] works under the usual assumptions of time-stationary and context-independent noise, where in addition the noise needs to be gate-independent, i.e.

$\tilde{\omega}(g) = \omega(g)\Lambda$ for a fixed Λ and all $g \in \mathcal{P}_n$. Furthermore, the noise Λ is assumed to be weak in the sense that the Pauli-twirled noise channel

$$\sum_{g \in \mathcal{P}_n} \omega(g)\Lambda\omega^\dagger(g) := \Lambda_{\text{Ptw}} \quad (2.112)$$

satisfies $\|\text{id} - \Lambda_{\text{Ptw}}\|_{\text{op}} \leq c$. Note that the Pauli twirl is the projection of Λ to the set of Pauli channels. Let now p be the vector of Pauli error probabilities of the channel Λ_{op} , where its first element p_0 is the probability of no error occurring. Then the p can be estimated up to relative error

$$\|\hat{p} - p\|_{\ell_2} \leq \mathcal{O}(\epsilon)(1 - p_0), \quad (2.113)$$

using $\mathcal{O}(nd/\epsilon^2)$ samples with high probability. The relative error scaling is particularly advantageous since in current experiments, Pauli-gates can often be implemented with very high fidelities in which case $(1 - p_0) \rightarrow 0$. Thus small errors can still be resolved without prohibitively many samples. Although the scaling in the dimension is optimal, estimating all Pauli error probabilities still requires exponentially many samples. The authors then further give two settings in which a constant number of samples is required. The first setting considers the task of estimating only the Pauli error probabilities with respect to a subset $S \in \mathcal{P}$ of Paulis with $|S| = s$, like for instance all Paulis with limited support. The guarantee is then given as: $\|\hat{p}_{|S} - p_{|S}\|_{\ell_\infty} \leq \mathcal{O}(\epsilon)(1 - p_0)$ with high probability given $\mathcal{O}(\log(s) \log(s/\epsilon^2)/\epsilon^4)$ samples.

The second setting uses the ansatz that p can be factorized as over subsets $C \in [d^2]$ of Pauli-probabilities via

$$p(a) = \frac{1}{Z} \prod_j \phi_j(a|_{C_j}), \quad (2.114)$$

where $\phi_j : a|_{C_j} \mapsto \mathbb{R}^+$, $a \in \mathbb{Z}_2^{2n}$ and Z is a normalization constant. A probability distribution factorized in this way is also called a Gibbs random field or factor graph model. The latter name stems from the identification of the sets C_j and the indices $i \in [d^2]$ with a bipartite graph, each index and each set C_j are represented by a vertex and edges (i, j) are present if $i \in C_j$. A bounded degree factor graph of degree k is then a graph where each C_j is at most connected to k indices. Note that an immediate ansatz model for the factor graph can often be made according to the proximity of qubits or their connectivity in terms of multi-qubit gates. To determine the error distribution p , only the factor potentials ϕ_j corresponding to each set C_j have to be determined. The sample complexity for estimating a factor graph model is then shown to be polynomial, i.e.

$\|\hat{p} - p\|_{\ell_1} \leq \mathcal{O}(\epsilon) \|e_1 - p\|_{\ell_\infty}$ is achievable with high probability using $\mathcal{O}(n^2 \log(n)/\epsilon^2)$ many samples.

In a follow-up work by Harper *et al.* [138], the protocol was extended to the task of learning noise channels which are twirled over the local Clifford group. Twirling a single qubit channel Λ over the Clifford group according to Eq. (2.7) results in

$$\sum_{g \in \text{Cl}_1} \omega(g) \Lambda \omega^\dagger(g) = \text{Tr}[\Lambda \Pi_0] \Pi_0 + \frac{1}{3} \text{Tr}[\Lambda \Pi_{ad}] \Pi_{ad}, \quad (2.115)$$

where $\Pi_0 = |\mathbf{1}\rangle\langle\mathbf{1}|$ and $\Pi_{ad} = \sum_{a \neq 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a|$ are projectors onto the irreps of Cl_1 . Similarly, twirling an n -qubit channel over the Clifford group leads to a sum over the irreps of Cl_1^n , and the resulting channel is described by 2^n parameters. In Harper *et al.* [138] it is then demonstrated how these can be learned in an analogous way to the 4^n Pauli error probabilities using a factor graph model. The authors further show how the learned factor graph model can be used to compute correlations between the error probabilities of different qubits, which enables them to provide an informative visualization of the average crosstalk on a device of 14 superconducting qubits.

Remaining practical shortcomings of these methods are first the strong assumption of gate-independent noise, as well as the problem that the factor graph has to be known beforehand and is not learned in the procedure. In a later work by Rouzé and Stilck França [139], the second shortcoming was solved by a method that can learn the factor graph efficiently using $\mathcal{O}(\log(n))$ samples. Even though it is as of yet unclear from a theoretical point of view how to interpret the reconstructed error probabilities for experiments where the noise is generally gate-dependent (or even time- and context-dependent), these methods offer a scalable estimation of error rates and crosstalk. They are thus a very promising benchmarking tool for increasingly larger devices, while the error correlations as a measure of crosstalk can aide in device calibration.

Since a Pauli channel is diagonal in the Pauli-basis, learning it amounts to learning all its eigenvalues. For general channels the eigenbasis is not known beforehand, and the question arises if all the eigenvalues of a general channel can still be learned efficiently, i.e.

with a sample complexity that scales linearly in the number of eigenvalues. In Helsen *et al.* [140] a SPAM-robust method to estimating the eigenvalues, termed *spectral quantum tomography*, was introduced. It builds on the observation that if a channel $\mathcal{C}$ given by the PTM matrix C is diagonalizable via $C = VDV^\dagger$, it holds that

$$\sum_{a \in \mathbb{F}_2^{2n}} (\sigma_a | \mathcal{C}^k | \sigma_a) = \text{Tr}(C^k) = \sum_{i=1}^{d^2} \lambda_i^k. \quad (2.116)$$

This gives a prescription on how to measure the signal $g(k) := \sum_{i=1}^{d^2} \lambda_i^k$: Choose a Pauli σ_a , prepare eigenstates of σ and measure $\mathcal{C}^k$ applied to each eigenstate in the basis of σ_a to reconstruct

$(\sigma_a | \mathcal{C}^k | \sigma_a)$. State preparation and measurement noise can w.l.o.g be modeled by noise channels $|\tilde{\sigma}_a\rangle = \Lambda_{\text{prep}}|\sigma_a\rangle$ and $\langle\tilde{\sigma}_a| = \langle\sigma_a| \Lambda_{\text{meas}}$. The noisy measurement of $\mathcal{C}^k$ is therefore equivalent to the noise-free measurement of $\Lambda_{\text{meas}} \mathcal{C}^k \Lambda_{\text{prep}}$ and

$$\text{Tr}(\Lambda_{\text{meas}} \mathcal{C}^k \Lambda_{\text{prep}}) = \text{Tr}(\Lambda_{\text{prep}} \Lambda_{\text{meas}} \mathcal{C}^k) = \sum_{i=1}^{d^2} \alpha_i \lambda_i^k =: \tilde{g}(k), \quad (2.117)$$

where α_i are SPAM parameters. After repeating the measurement for all d^2 Paulis and different values of k , one has access to the series $(\tilde{g}(k_1), \dots, \tilde{g}(k_K))$. Standard signal analysis techniques such as ESPIRIT [141] can then reconstruct the eigenvalues from this series. Although SPAM-robust with easy to implement measurements, the protocol needs d Pauli eigenstates to determine each of the $d^2 - 1$ values $(\tilde{\sigma}_a | \mathcal{C}^k | \tilde{\sigma}_a)$, which leads to the total number of measurement settings $Kd(d^2 - 1)$.

Learning channel properties with classical shadows

The classical shadow formalism, which we describe in Section 2.3, can be applied to channel estimation tasks as follows. Consider an experimental protocol where a random sequence of gates $g = (U_1, \dots, U_l)$ is applied to a fixed initial state ρ , followed by a POVM $\{E_i\}$. Let $p(i, g) = \text{Tr}[E_i U_l \circ \dots \circ U_1(\rho)]$ and assume the unitaries are classically efficiently representable, such as for instance multi-qubit Cliffords. Then the set

$$\{p(i_1, g_1), \dots, p(i_N, g_N)\} \quad (2.118)$$

is called a *gate set shadow* [43]. It can be interpreted as a classical model of the device, where term 'classical' refers to the fact that each E_i and g_i is represented with $\text{poly}(n)$ parameters. A key difference between shadow based methods and standard QPT is that in most cases the goal is not the reconstruction of a channel, but the simultaneous estimation of channel properties such as average gate fidelities or crosstalk metrics [43].

A different scenario is considered in Huang *et al.* [142], where the goal is to estimate $\text{Tr}[OC(\rho)]$ with a fixed channel $\mathcal{C}$ over a distribution of input states, for different observables. The classical shadow from which these expectation values are estimated is given by

$$\left\{ \left(E_k = \bigotimes_{i=1}^n |s_{ik}^{\text{out}}\rangle \langle s_{ik}^{\text{out}}|, \rho_k = \bigotimes_{i=1}^n |s_{ik}^{\text{in}}\rangle \langle s_{ik}^{\text{in}}| \right) \right\}_{k=1}^N, \quad (2.119)$$

where $|s_{ik}\rangle$ are stabilizer states and thus classically efficiently representable. Let a bounded observable be defined as an observable that can be written as the sum over local observables, such that only a constant number of these local observables have support on any given qubit. Using sophisticated proof techniques, it was shown in Huang *et al.* [142] that for any distribution p_{inv} over quantum states that are invariant under single qubit unitaries, an average prediction error

$$\mathbb{E}_{\rho \sim p_{\text{inv}}} |\widehat{\langle O \rangle}_\rho - \text{Tr}[OC(\rho)]|^2 \leq \epsilon \quad (2.120)$$

is achieved with high probability using $2^{\mathcal{O}(\log(\epsilon^{-1}) \log(n))}$ samples.

Two concurrent works by Kunjummen *et al.* [143] and Levy *et al.* [144] give a more direct generalization of classical shadow to quantum processes by the standard identification of a target channel and its Choi state. This results in an analogous bound to shadow estimation for states, in that the prediction errors $|\widehat{\langle O_i \rangle}_{\rho_i} - \text{Tr}[O_i \mathcal{C}(\rho_i)]| \leq \epsilon$ over any set $\{\Omega_i = \rho_i^T \otimes O_i\}_{i=1}^M$ of observable-state combinations can be realized with probability $1 - \delta$ using

$$\frac{\log(2M/\delta)}{\epsilon^2} 4^n \max_{i \in [M]} \|\Omega_i - \text{Tr}[\Omega_i] \mathbb{1}/2\|_{\text{shadow}}^2 \quad (2.121)$$

many samples. A definition of the shadow norm can be found in Section 2.3. The key difference to classical shadows for states is the additional exponential factor 4^n , which makes the method not scalable.

Hamiltonian and Lindbladian tomography

Parametrizing a unitary gate via its generating Hamiltonian and a general CPT map via a time-dependent or time-independent Lindbladian Eq. (2.70) directly provides a physical model which is easier to relate to an experimental setup and its control. Although it is possible to construct a Lindblad parametrization of a CPT map determined via GST or QPT (see e.g. [145]), most works consider the problem of directly reconstructing Hamiltonians or Lindbladians from measurement data. This has the advantage that sparse models are often well-motivated. For instance if a Hamiltonian is given as

$$H = \sum_{a \in \mathbb{F}_2^{2n}} h_a \sigma_a, \quad (2.122)$$

it can be argued in many scenarios that a system is described by at most k -local interactions, i.e.

only coefficients h_a with $\text{supp}(a) \leq k$ are non-vanishing. Several recent works have proposed algorithms to efficiently estimate the Hamiltonian parameters or more generally the parameters of a Lindbladian [146–152]. The Lindbladian parameters can be accessed from the time evolution according to the master equation

$$\dot{\rho} = \mathcal{L}(\rho) = i[\rho, H] + \sum_{a, a' \in \mathbb{F}_2^{2n}} \beta_{a, a'} \left(\sigma_a \rho \sigma_{a'} - \frac{1}{2}(\sigma_{a'} \sigma_a \rho + \rho \sigma_{a'} \sigma_a) \right), \quad (2.123)$$

where we have written the dissipative terms with respect to the Pauli basis. The time derivative of an observable expectation value is thus given by

$$\frac{d}{dt} \text{Tr}[O\rho] = \text{Tr}[O\mathcal{L}(\rho)]. \quad (2.124)$$

One can further choose a set of probe states and observable combinations $\{(O_i, \rho_i)\}_{i=1}^N$ for which the time derivatives $d \text{Tr}[O_i \rho_i]/dt$ are estimated. As seen from the parametrization of H and the dissipative terms, each estimated time derivative is a linear function of the parameters h_a and $\beta_{a, a'}$. If the set of probe settings $\{(O_i, \rho_i)\}_{i=1}^N$ is informationally complete, $\mathcal{L}$ can thus be determined by solving a system of linear equations. In Stilck França *et al.* [148], $d \text{Tr}[O_i \rho_i(t)]/dt$ was estimated for a time series $t_1, \dots, t_l$ and a low order polynomial in t was fitted to the resulting data. From this polynomial, the time derivative can be accurately extracted. The authors were able to prove that a total sample count of $\mathcal{O}(\epsilon^{-2} \text{poly} \log(n, \epsilon^{-1}))$ suffices to estimate all parameters of a local Lindbladian to errors $|\hat{\beta}_{a, a'} - \beta_{a, a'}| \leq \epsilon$ and $|\hat{h}_a - h_a| \leq \epsilon$ with high probability using local Pauli measurements and Pauli eigenstates. It was further shown in numerical examples that the time derivative based on polynomial interpolation can greatly outperform methods that use a finite differences approach to determine derivatives.

It is also worth mentioning the work by Huang *et al.* [151], where a Heisenberg scaling in learning the parameters of a Hamiltonian was achieved, meaning that only $\mathcal{O}(\text{poly} \log(\epsilon^{-1}))$ measurements and time evolutions for times no longer than $\mathcal{O}(\epsilon^{-1})$ are used in the protocol. The method builds on robust phase estimation [153], which can estimate gate parameters up to error $1/N$, where N is the number of times the gate is applied. Moreover, the protocol only uses single qubit Cliffords, local measurements and works in a SPAM robust fashion, although it assumes the ability to interleave Cliffords in a time evolution, which is a higher degree of control than what is assumed in most other works and difficult to achieve in practice.

Tensor network models

Using low bond dimension tensor network models to drastically reduce the number of parameters required to represent a quantum states has been an established method in quantum information and condensed matter physics [154, 155] under the name of matrix product states or tensor trains. The tensor network formalism has also been used to great effect in quantum state tomography [85, 156, 157], with extensions to ancilla assisted QPT [158]. It has recently been applied to standard QPT and to models of non-Markovian dynamics, for both of which we give a brief overview in this section. Representing quantum circuits as tensor network is furthermore a central idea of the original work presented in this thesis (see Chapter 3 and Appendix A).

In Torlai *et al.* [159] the Choi matrix $\mathcal{J}(\mathcal{C})$ of a channel acting on an n -qubit system is parametrized via k tensors $\{A_i\}_{i=1}^n$ of dimensions $r_{K,i} \times r_{S,i} \times 4$. Here i labels the qubit subsystems and the number $r_{K,i}$ can be understood as a local Kraus rank, while the rank $r_{S,i}$ is the Schmidt rank of the bipartition of the first i subsystems versus the rest of the system. The ranks $r_{S,i}$ thus quantify how entangling the gate is along the bipartition, with $r_{S,i} = 1$ for gates that factorize along the subsystem boundary. The resulting parametrization of $\mathcal{J}(\mathcal{C})$ reads

$$\mathcal{J}(\mathcal{C}) = \sum_{l_1=1}^{r_{S,1}} \cdots \sum_{l_n=1}^{r_{S,n}} \sum_{k_1=1}^{r_{K,1}} \cdots \sum_{k_n=1}^{r_{K,n}} \prod_{i=1}^n (A_i)_{l_{i-1}, l_i, k_i}^{i'_1, j'_1} (A_i^*)_{l_{i-1}, l_i, k_i}^{i_1, j_1}, \quad (2.125)$$

where l_0 and l_{n+1} are dummy indices. The graphical representation of the tensor network is given in Figure 2.5. As can be understood from the decomposition shown in the Figure, the

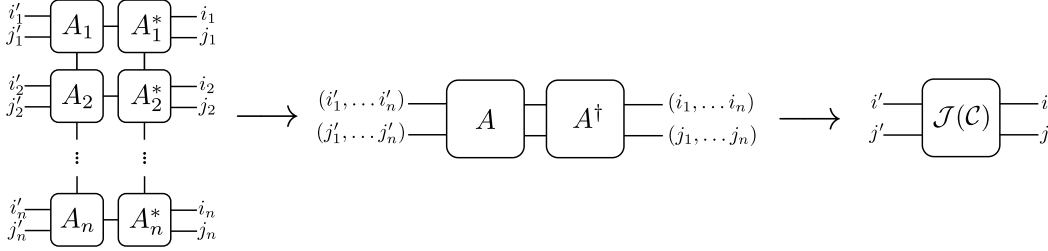


Figure 2.5: Tensor network factorization the Choi matrix over Kraus-indices (horizontal) and subsystem-indices (vertical).

parametrization automatically yields a positive Choi state. Choi states $\mathcal{J}(\mathcal{C})$ written in this parametrization have also been studied under the name of *locally purified density operators* [160]. Torlai *et al.* [159] then use a gradient descent algorithm and techniques from supervised learning to find an optimal fit to the log-likelihood cost function (Eq. (2.75)), evaluated on measurement outcome probabilities. Trace preservation is achieved by adding a penalty term to the cost function which goes to zero for trace preserving models. The input states are randomly chosen Pauli eigenstates, whereas the POVM is the uniform POVM over all Pauli-eigenstates, i.e.

the set of POVM elements

$$\left\{ \frac{1}{6^n} \bigotimes_{i=1}^n \frac{1}{2} (\mathbb{1} + (-1)^{x_i} \sigma_a) \right\}_{a \in (\mathbb{Z}_2^2 \setminus 0)^n, x \in \{0,1\}^n}. \quad (2.126)$$

Since it is a priori not clear what the correct Kraus and Schmidt ranks are, progressively higher ranks have to be fitted to the data until a desired level of convergence in the cost function is reached. The method is numerically tested to generate estimates of channels on up to 10 qubits. The key takeaway is that although the method remains heuristic, low reconstruction errors measured by the infidelity are achievable with greatly reduced measurement overhead as compared to full tomography, which would for 10 qubits require a total of 2^{40} measurement settings and thus be entirely infeasible.

Tensor networks have also recently been employed to the growing field of non-Markovian quantum process tomography [161–166]. What is termed non-Markovianity in the literature typically refers to dynamics arising from initial conditions where the initial state of system and environment are entangled, as opposed to the assumption in Eq. (2.54). Thus far we have described the QPT setting where gates are assumed to be time-stationary and independent of previous control operations and measurements. Furthermore, state preparation as well as measurement were largely assumed to be controllable. In the non-Markovian setting, none of these assumptions are made, since unitary dynamics on entangled system-environment states incorporate the mentioned effects of time- and context dependence. This could manifest in straightforward ways such as slow drifts in time, or in more complex system-environment interactions where gates applied at an earlier time step affect the background noise at a later time step in a circuit. Modelling additional degrees of freedom arising from these effects greatly increases the computational cost of tomography procedures, and simplified models that capture the most relevant correlations via tensor networks are essentially required to make reconstruction practical. In White *et al.* [167], the authors introduce such an ansatz model, which we briefly explain in the following. Let $\mathcal{H} = \mathcal{H}_S \otimes \mathcal{H}_E$ be the combined Hilbert space of system and environment, which is assumed to be in the initial state ρ_0 at the beginning of the experiment. The core component of the model is a so-called process tensor, which is simply a multilinear map $\mathcal{T}_l$ that takes in an ordered set of l control operations described by CPT maps $(\mathcal{A}_{l-1}, \dots, \mathcal{A}_0)$ and produces an output state. Note that in this formulation, the process tensor already incorporates the fixed initial state ρ_0 . Experimentally accessible outcome probabilities for a given POVM $\{E_i\}$ are then given as

$$p_j(\mathcal{A}_{l-1}, \dots, \mathcal{A}_0) = \text{Tr}[E_j \mathcal{T}_l(\mathcal{A}_{l-1}, \dots, \mathcal{A}_0)]. \quad (2.127)$$

In Figure 2.6 the action of the process tensor is visualized. Standard non-Markovian tomography assumes the control operations to be known, meaning they can be used to probe the process tensor. In [167] they are learned in a self-consistent manner, with a formalism based on fiducial sequences adapted from the GST literature (see Section 2.2.4). The process tensor incorporates correlations between the control sequences, background correlations in the environment and the influence of control on the environment. Since for high fidelity implementations of quantum circuits most of the correlations are by now fairly small, a tensor network decomposition of $\mathcal{T}$ with small bond dimensions can in practice be expected explain experimental data well. The

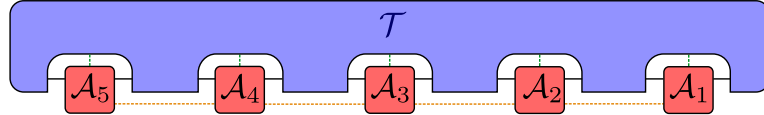


Figure 2.6: Contraction of the process tensor with CPT maps $\mathcal{A}_i$, where black bonds represent the physical indices, red dashed bonds are due to correlations across time steps between control operations, while dashed green bonds are due to interactions between the environment and the control operations.

model is then fit to measurement data by standard local optimization on the log-likelihood cost function with regularization to ensure causality and trace preservation. In numerical simulations, the method is shown to produce significantly lower generalization errors than non-self-consistent process tensor tomography, quantified by how measured output states for different sequences deviate from predicted output states using the estimated model. A low bond dimension tensor network was also applied to reconstruct the process tensor of four qubits across four time steps ($l = 4$) from experimental data of a superconducting qubit device. In comparison, GST is typically restricted to two qubits [168] or three qubits (Chapter 3). The achievement of self-consistent process tensor tomography on 4-qubit processes, albeit with a low bond dimension model, is thus quite significant and speaks for the usefulness of tensor network parametrizations for the characterization of quantum dynamics.

2.2.4 Gate set tomography

A common drawback of most QPT protocols is the lack of device independence and self-consistency. Consider a standard QPT experiment where an initial state is prepared, the gate in question applied and a POVM measured. Typically, there is an initial state and a measurement that is native to the experiment, for instance preparation in the ground state and a measurement in the computational basis. The full set of initial states and measurements required for QPT then have to be prepared using specific operations. In the simplest case only local Clifford unitaries need to be applied to prepare eigenstates of the Pauli matrices, but many protocols required general stabilizer states or a more specialized set of states. But even local Clifford unitaries, which are of high fidelity in most modern experiments, can limit the accuracy of QPT. For instance if these local Cliffords each come with small depolarizing noise, then in QPT this noise would turn up in the tomographic reconstruction of the gate in question. This limits the accuracy of QPT, even for single qubits, where the target gate is often of the same quality as the gates used for state preparation and measurement. The same holds true for multi qubit QPT with methods that require general stabilizer states, which in turn need global Clifford unitaries to be prepared. Global Cliffords generally require $\mathcal{O}(n^2)$ entangling gates to implement. This leads to high state preparation and measurement errors in current experiments, where two qubit gates are often much noisier than their single qubit counterparts. Moreover, most recent attempts to make QPT SPAM-robust use variations of RB, where some of the gates are assumed to be trusted, or only average noise parameters are learned.

An alternative solution is given by self-consistent process tomography protocols [28–30, 169, 170], which circumvent the SPAM noise problem by simultaneously learning all system parameters. This includes the native initial state, the native POVM and all gates applied during the protocol. The assumptions made on the experiment are otherwise the same as in QPT: time stationary and Markovian dynamics.

Let $\mathfrak{G} = (\mathcal{G}_1, \dots, \mathcal{G}_k)$ be the tuple of gates to be used for self-consistent process characterization. Let further $(\rho_j)_j$ and $(E_i)_i$ be tuples of initial states and POVM elements, respectively. The gates are usually chosen to be native gates of the experimental platform, but can in principle be composed of arbitrary circuits. The task to self-consistently estimate the system parameters

$$\mathcal{X} = ((E_i)_i, \mathfrak{G}, (\rho_j)_j) \quad (2.128)$$

is then termed Gate Set Tomography (GST) after the work by Blume-Kohout *et al.* [30] and subsequent advancements by the same group [168, 171], which includes a comprehensive Python package called pyGSTi [172, 173]. Because GST as implemented in pyGSTi is by far the dominant self-consistent QPT method, the term GST has become genericized and is often used for both, the pyGSTi implementation and self-consistent QPT in general.

The experimental model of GST for a given gate set and an experiment with the native POVM given by $E = (E_1, \dots, E_M)$ is summarized in Figure 2.7. A GST experiment is defined by a predetermined set of gate sequences $I \subseteq \bigcup_{l=1}^L [k]^l$, where we allow varying sequence lengths up to a maximum of L . The GST estimation problem is given by

$$\begin{aligned} & \underset{\mathfrak{G}, E, \rho_0}{\text{minimize}} && f(\mathfrak{G}, E, \rho_0 | \hat{p}) \\ & \text{subject to} && \mathcal{G}_i \text{ CPT}, \quad \rho_0 \in \mathcal{S}, \quad E_i \leq 0, \quad \sum_i E_i = \mathbb{1}, \end{aligned} \quad (2.129)$$

where $\hat{p} = (\hat{p}_s(j))_{s \in I, j \in [M]}$ is the full matrix of outcome probabilities. We write the minimizers of f as $\hat{\mathfrak{G}} = (\hat{\mathcal{G}}_1, \dots, \hat{\mathcal{G}}_k)$, $\hat{E}$ and $\hat{\rho}_0$. Previously discussed cost functions such as the sum of squares or the Kullback-Leibler divergence depend on the Born rule probabilities $p_{s,j}(\mathfrak{G}, E, \rho_0) = \text{Tr}[E_j \mathcal{G}_{s_1} \circ \dots \circ \mathcal{G}_{s_L}(\rho_0)]$. Those are arbitrary degree polynomials in the gate parameters, for instance the sequence set I might contain the sequence $(1, 1, \dots, 1)$ for which $p_{s,j} = \text{Tr}[E_j \mathcal{G}_1^{\circ L}(\rho_0)]$.

There are several factors which complicate the GST estimation problem. First and foremost, the cost function is non-convex and can be riddled with local minima and saddle points, requiring

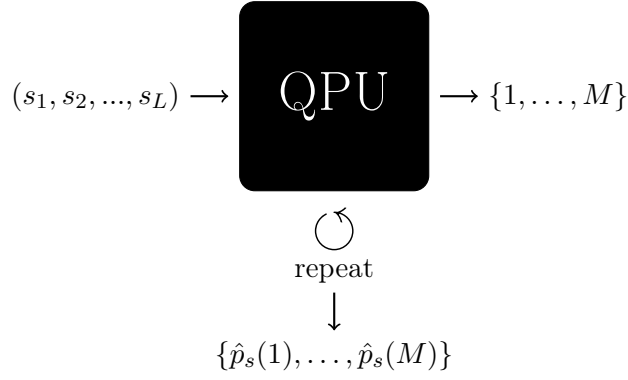


Figure 2.7: The measurement scenario considered in GST: A gate sequence described by a tuple of gate indices $s = (s_1, \dots, s_L)$ is send to the black box quantum processing unit (QPU). The QPU returns the measurement outcome in the native POVM, and the process is repeated for the same gate sequence. After sufficient statistics are gathered, outcome probability estimates for the given sequence are given by the observed outcome frequencies.

iterative optimization methods together with a good starting point to find the global minimum. Second, the physicality constraints have to be incorporated, which is usually done via projection, regularization, or as in Chapter 3, via manifold optimization. The third problem is that the gate set and the set of sequences have to be chosen to ensure tomographic completeness. In other words the measurement map $p : (\mathfrak{G}, E, \rho_0) \mapsto [0, 1]^{I \times M}$ has to be injective, meaning that no distinct $(\mathfrak{G}, E, \rho_0)$ produce the same measurement data. In practice this problem is eased by the fact that one has a reasonably founded belief on what process each gate should implement and sequences can be designed accordingly.

In the following we will summarize the approaches taken in the literature to tackle the GST estimation problem, with a focus on the standard GST implementation pyGSTi. In what is called linear GST [168, 174] the route of standard QPT is followed, i.e.

we aim to get access to $(E_i | \mathcal{G}_j | \rho_k)$, where both $\{E_i\}$ and $\{\rho_k\}$ span $L(\mathcal{H})$. For simplicity we now assume that the native measurement is described by a two-outcome POVM $\{E_0, \mathbb{1} - E_0\}$. Let $\text{circ}(\mathfrak{G})$ be the set of all circuits composed of elements of $\mathfrak{G}$, then $\mathfrak{G}$ can generate spanning sets $\{E_i\}$ and $\{\rho_k\}$ if there exist $\mathcal{F}_i^{\text{out}} \in \text{circ}(\mathfrak{G})$ and $\mathcal{F}_k^{\text{in}} \in \text{circ}(\mathfrak{G})$ such that

$$(E_i | = (E_0 | \mathcal{F}_i^{\text{out}} \quad \text{and} \quad | \rho_k) = \mathcal{F}_k^{\text{in}} | \rho_0). \quad (2.130)$$

The circuits $\mathcal{F}_i^{\text{out}}$ and $\mathcal{F}_k^{\text{in}}$ are commonly called fiducials. Oftentimes state preparation fiducials $\mathcal{F}_k^{\text{in}}$ and measurement fiducials $\mathcal{F}_i^{\text{out}}$ coincide, for instance if $\rho_0 = E_0$. Let us now assume we are given sets of exactly d^2 of each, state preparation and measurement fiducials. A *linear GST* experiment estimates the probabilities

$$p_{ijk} = (E_0 | \mathcal{F}_i^{\text{out}} \mathcal{G}_j \mathcal{F}_k^{\text{in}} | \rho_0) = \sum_{l,m} (E_0 | \mathcal{F}_i^{\text{out}} | l) (l | \mathcal{G}_j | m) (m | \mathcal{F}_k^{\text{in}} | \rho_0) = \sum_{l,m} A_{il} (G_j)_{l,m} B_{mk}, \quad (2.131)$$

where for the last equality we defined $A = \sum_i |i\rangle (E_0 | \mathcal{F}_i^{\text{out}}$ and $B = \sum_k \mathcal{F}_k^{\text{in}} | \rho_0) \langle k|$ for some orthonormal basis $\{|i\rangle\}_{i=1}^{d^2}$ of $L(\mathcal{H})$. It can then be straightforwardly shown that the matrices A, B defined in this way satisfy Eq. (2.131). The matrix A has just the measurement effects $(E_0 | \mathcal{F}_i^{\text{out}}$ as its rows, while B has the initial states $\mathcal{F}_k | \rho_0$ as its columns. These matrices thus define a measurement map in the same way as in process tomography (Eq. (2.95)), with the difference that now they are unknown. From directly measuring the fiducial sequences without an inserted gate set in the middle, we gain access to

$$(E_0 | \mathcal{F}_i^{\text{out}} \mathcal{F}_k^{\text{in}} | \rho_0) = (AB)_{ik}, \quad R_k := (E_0 | \mathcal{F}_k^{\text{in}} | \rho_0) \quad \text{and} \quad L_i := (E_0 | \mathcal{F}_i^{\text{out}} | \rho_0). \quad (2.132)$$

Here AB is the just gram matrix of the set of POVM vectors $(E_i|$ and the initial state vectors $|\rho_i\rangle$. Let P_j be the matrix with entries $(P_j)_{ik} = p_{ijk}$. If A and B are invertible, we can reconstruct the full gate set up to a similarity transformation:

$$(AB)^{-1}P_j = B^{-1}A^{-1}AG_jB = B^{-1}G_jB. \quad (2.133)$$

Since B is unknown, we do not obtain a single estimate of G_j , but an equivalence class $B^{-1}G_jB$ for $B \in \text{GL}(d^2)$. This is known as *gauge freedom* in the GST literature and fundamentally unavoidable in self-consistent tomography.

An analogous argument can be made to reconstruct ρ_0 and E_0 up to the gauge matrix B . Let $R = \sum_k |k\rangle\langle k| (E_0 | \mathcal{F}_k^{\text{in}} | \rho_0)$ and $L = \sum_i (E_0 | \mathcal{F}_i^{\text{out}} | \rho_0) \langle i|$. Then we find that

$$(AB)^{-1}R = B^{-1}A^{-1}A|\rho_0\rangle = B^{-1}|\rho_0\rangle, \quad L = (E_0|B. \quad (2.134)$$

We will now discuss the gauge freedom in detail before we define tomographic completeness in the context of linear GST.

The gauge problem of gate set tomography

A gauge freedom is present in any circuit model of a quantum experiment where no prior information is given. The only way to gather information about the system is then via measurement outcomes governed by the Born-rule. One can immediately see that the following equality

$$p(E) = (E | \mathcal{G}_L \cdots \mathcal{G}_1 | \rho) = (E | B^{-1}B\mathcal{G}_LB^{-1}B\mathcal{G}_{L-1} \cdots B^{-1}B\mathcal{G}_1B^{-1}B | \rho) \quad (2.135)$$

holds and the probability $p(E)$ of observing outcome E is the same for any invertible matrix B . In Figure 2.8 the gauge freedom is depicted for the tensor network representation of a GST experiment. In the theory of matrix product states, this gauge freedom is also well known and

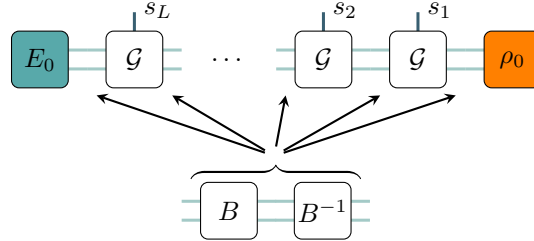


Figure 2.8: The gauge freedom depicted in the tensor network representation of a GST experiment where individual gates $G_{s_1}, \dots, G_{s_L}$ are represented as channels.

used to bring the matrix product state into a canonical form [175]. Without any restrictions the set of gauge transformations is equal to the group $\text{GL}(d^2)$. Note that any physicality constraints on the gates do restrict the gauge, since if $\mathcal{G}_i$ is CPT, $B^{-1}\mathcal{G}_iB$ is not necessarily CPT for all $B \in \text{GL}(d^2)$. In Rudnicki *et al.* [176] it was shown using the PTM representation of a channel from Eq. (2.61), that for $B^{-1}\mathcal{G}_iB$ to be trace preserving for any Gate $\mathcal{G}_i$, the gauge matrix also has to be in block form

$$B = \begin{pmatrix} 1 & 0 \\ u & T \end{pmatrix} \quad (2.136)$$

in the Pauli basis, where $u \in \mathbb{R}^{d^2-1}$ and $T \in \mathbb{R}^{d^2-1 \times d^2-1}$. In other words, B has to be trace preserving as well and the set of gauge matrices in Eq. (2.136) again form a group. Furthermore, in [45] we show that if the gate set is universal, i.e.

any pure state can be prepared via an arbitrarily long sequence of gates from the gate set, then gauge transformations that preserve the CPT condition are limited to unitary and anti-unitary channels. In practice the gauge problem becomes more complicated, since B does not need

to map all possible CPT maps to CPT maps, but only those accessible through sequences of noisy gates $\{\tilde{\mathcal{G}}_1, \dots, \tilde{\mathcal{G}}_k\}$ performed in the experiment. This is due to the fact that if the gates $\{B^{-1}\tilde{\mathcal{G}}_1B, \dots, B^{-1}\tilde{\mathcal{G}}_kB\}$ are CPT, so is any sequence $B^{-1}\tilde{\mathcal{G}}_sB$, since

$$B^{-1}\tilde{\mathcal{G}}_sB = B^{-1}\tilde{\mathcal{G}}_{s_L} \dots \tilde{\mathcal{G}}_{s_1}B = B^{-1}\tilde{\mathcal{G}}_{s_L}B \dots B^{-1}\tilde{\mathcal{G}}_{s_1}B. \quad (2.137)$$

Consider the example where $B = \Lambda_p$, the depolarizing channel. Since Λ_p commutes with any unital CPT map, we have for all unital gates in the gate set that $\mathcal{G}_i = \Lambda_p^{-1}\tilde{\mathcal{G}}_i\Lambda_p$. But if the initial state ρ_0 is pure, then its gauge transformed version $\Lambda_p^{-1}|\rho_0\rangle$ would lie outside the Bloch sphere and be unphysical. However, Λ_p is a valid gauge transformation if ρ_0 is mixed and satisfies $\Lambda_p^{-1}|\rho_0\rangle \in \mathcal{S}(H)$.

Gauge transformations for noisy gate sets that preserve physicality conditions do not generally have a group structure and even vary at different points in the parameters space of the gate set (see Nielsen *et al.* [168]). A gate set at the boundary of the CPT constraint region then has fewer allowed gauge directions than a gate set in the interior, just as we saw with the depolarizing channel example.

An argument can be made that since gauge transformations do not change measurement outcomes, the quality of gate sets should be assessed in a completely gauge invariant manner. For example a unitary gauge freedom is just a basis change on the Hilbert space, and from the point of view of quantum information processing, we do not care in which basis an algorithm is run, as long as everything is consistent. However, typically there is a very specific target implementation, and we are interested in whether the actual gates, as measured by GST, correspond to the target. In Section 2.2.2 we saw that typical distance measures like the average gate fidelity and the diamond distance are not gauge invariant. Furthermore, in order to detect overrotations or other noise sources for calibration purposes, the GST estimate has to be given in the gauge of the target gate set. This means that after an initial GST estimate is given, another optimization step is required called the gauge fit:

$$\hat{B} = \underset{B}{\operatorname{argmin}} \left(\sum_{i=1}^k \|B^{-1}\hat{\mathcal{G}}_iB - \mathcal{G}_i^{\text{target}}\|_F + \|B^{-1}|\hat{\rho}\rangle - |\rho\rangle^{\text{target}}\|_F + \|(\hat{E}|B - (E|^{\text{target}}\|_F \right), \quad (2.138)$$

where B can be restricted to different gauge sets depending on physicality constraints. Different norms could be chosen to measure the distance in gauge optimization, but the Frobenius norm is sufficient in practice [30]. An alternative to gauge optimization at the end of a GST protocol is to add the gauge error Eq. (2.138) as a regularization term in the cost function of the GST estimation problem Eq. (2.129). This route was taken in Sugiyama *et al.* [177], where the authors prove that giving the regularization term a decreasing weight throughout the optimization does not lead to a bias, or more concretely: The final estimate also minimizes the unregularized cost function in the infinite sample limit.

But what can be said about the quality of a gate set in a gauge invariant way? In Section 2.2.2 we introduced the spectral 1-distance, a measure that depends only on the spectrum of a channel and is thus gauge invariant. Other measures that depend on the spectrum could be defined analogously. The most direct route is to compare gate sets on the space of Born probabilities, which are naturally gauge invariant. Assume we are given $p_{s,j}(\mathfrak{G}, E, \rho)$ and $p_{s,j}(\tilde{\mathfrak{G}}, \tilde{E}, \tilde{\rho})$, where the probabilities can be calculated either from a classically stored gate set or determined in a quantum experiment. Then the gate sets can be compared with either the sum of squares error (Eq. (2.72)), the total variation error (Eq. (2.71)), the Kullback-Leibler divergence (Eq. (2.73)) or any other distance measure between probability distributions. A normalized variant of the total variation error called the *mean variation error* (MVE) was proposed in Lin *et al.* [178] together with an estimation protocol. The MVE is the main gauge invariant error measure we use in Chapter 3. In addition, worst case variants of the aforementioned distance measures can

be considered, for example

$$\max_{s,j} |p_{s,j}(\mathfrak{G}, E, \rho) - p_{s,j}(\tilde{\mathfrak{G}}, \tilde{E}, \tilde{\rho})| \quad \text{and} \quad \max_{s,j} p_{s,j}(\mathfrak{G}, E, \rho) \log \left(\frac{p_{s,j}(\mathfrak{G}, E, \rho)}{p_{s,j}(\tilde{\mathfrak{G}}, \tilde{E}, \tilde{\rho})} \right). \quad (2.139)$$

To quantify how a target implementation differs from its experimental implementation, a direct approach is to compute the circuit probabilities of the target circuits (up to the system size limit of classical computation) and compare them to the experimentally estimated probabilities. This is often done in current experiments via *cross entropy benchmarking* [179, 180]. There the cross entropy difference

$$\Delta H = H_0 - \sum_j p_j(\tilde{\mathcal{U}}) \log \left(\frac{1}{p_j(\mathcal{U})} \right), \quad (2.140)$$

for a classically simulated unitary circuit $\mathcal{U}$ and its implementation $\tilde{\mathcal{U}}$ is given. Here H_0 is the cross entropy of uniform sampling, i.e.

the implementation outputs measurement outcomes uniformly at random. Given a set of measurement outcomes $\{j_1, \dots, j_m\}$ for a fixed circuit $\mathcal{U}$, the cross entropy difference can be empirically estimated via

$$\widehat{\Delta H} = H_0 - \frac{1}{m} \sum_{i=1}^m \log \left(\frac{1}{p_{j_i}(\mathcal{U})} \right) = \Delta H + \mathcal{O}(1/\sqrt{m}). \quad (2.141)$$

A related variant defines the linear cross entropy fidelity $\mathcal{F}_{\text{XEB}} = d \sum_i p_{j_i}(\tilde{\mathcal{U}})/m - 1$ as a quality metric for $\tilde{\mathcal{U}}$, and this measure was used for the first quantum advantage demonstration [15].

Tomographic completeness

The question now arises as to how the fiducials need to be chosen in order for linear GST to be tomographically complete. First we observe that for the method to work, we need that the Gram matrix AB is invertible. Since it holds that $\text{rank}(AB) \leq \min(\text{rank}(A), \text{rank}(B))$, it follows that if AB is invertible, A and B have to be invertible as well. That means even though the fiducials are in general not known, we can immediately check for tomographic completeness by determining if the experimentally measured gram matrix AB is invertible. In Sugiyama *et al.* [177] the authors formally prove informational completeness up to gauge under the following conditions: The set of measurement effects $\{(E_0 | \mathcal{F}_i^{\text{out}})\}_{i=1}$ and the set of states $\{\mathcal{F}_k^{\text{in}} | \rho_0\}_k$ are both frames of $L(\mathcal{H})$ and the outcomes

$$x_{ijk}^{\text{data}}(\mathfrak{G}, E_0, \rho_0) = (p_{ijk}(\mathfrak{G}, E_0, \rho_0), (AB)_{ik}(\mathfrak{G}, E_0, \rho_0), R_k(\mathfrak{G}, E_0, \rho_0), L_i(\mathfrak{G}, E_0, \rho_0)) \quad (2.142)$$

are measured for all fiducials i, k and gates j in the gate set. Furthermore, the gates in $\mathfrak{G}$ and $\tilde{\mathfrak{G}}$ are assumed to be invertible. For these measurements to be informationally complete up to gauge we need that

$$x_{ijk}^{\text{data}}(\mathfrak{G}, E_0, \rho_0) = x_{ijk}^{\text{data}}(\tilde{\mathfrak{G}}, \tilde{E}_0, \tilde{\rho}_0) \quad (2.143)$$

holds if and only if the gate sets $(\mathfrak{G}, E_0, \rho_0)$ and $(\tilde{\mathfrak{G}}, \tilde{E}_0, \tilde{\rho}_0)$ are in the same gauge equivalence class. We already confirmed the forward direction by noting that gauge equivalent gate sets produce the same measurement data. The idea in Sugiyama *et al.* [177] to prove the other direction will be given in the following. Let $\mathcal{F}_i^{\text{in}}, \mathcal{F}_i^{\text{out}}$ be the fiducials for gates $\mathfrak{G}$ and let $\tilde{\mathcal{F}}_i^{\text{in}}, \tilde{\mathcal{F}}_i^{\text{out}}$ be the fiducials for $\tilde{\mathfrak{G}}$. First since both sets of states and measurement effects are frames,

$$(E_0 | \mathcal{F}_i^{\text{out}} \mathcal{F}_k^{\text{in}} | \rho_0) = (\tilde{E}_0 | \tilde{\mathcal{F}}_i^{\text{out}} \tilde{\mathcal{F}}_k^{\text{in}} | \tilde{\rho}_0) \quad (2.144)$$

implies that there exists a unique invertible matrix B such that

$$\mathcal{F}_k^{\text{in}} | \rho_0 = B \tilde{\mathcal{F}}_k^{\text{in}} | \tilde{\rho}_0 \quad \text{and} \quad (E_0 | \mathcal{F}_i^{\text{out}} B^{-1} = (\tilde{E}_0 | \tilde{\mathcal{F}}_i^{\text{out}}. \quad (2.145)$$

We can use this and Eq. (2.143) to write

$$(E_0 | \mathcal{F}_i^{\text{out}} \mathcal{G}_j \mathcal{F}_k^{\text{in}} | \rho_0) = (\tilde{E}_0 | \tilde{\mathcal{F}}_i^{\text{out}} \tilde{\mathcal{G}}_j \tilde{\mathcal{F}}_k^{\text{in}} | \tilde{\rho}_0) = (E_0 | \mathcal{F}_i^{\text{out}} B^{-1} \tilde{\mathcal{G}}_j B \mathcal{F}_k^{\text{in}} | \rho_0). \quad (2.146)$$

Since again the measurements $(E_0 | \mathcal{F}_i^{\text{out}}$ and states $\mathcal{F}_k^{\text{in}} | \rho_0)$ form frames, it must hold that

$$B^{-1} \tilde{\mathcal{G}}_j B = \mathcal{G}_j \quad (2.147)$$

for all gates $j \in [k]$. Now we know that the gates are gauge equivalent and the states and measurement effects prepared by the fiducials are gauge equivalent. It remains to show that also ρ_0, E_0 and $\tilde{\rho}_0, \tilde{E}_0$ are gauge equivalent. This is easily verified, since $\mathcal{F}_k^{\text{in}}$ is itself a sequence of gates in $\mathfrak{G}$ and thus by Eq. (2.147) and Eq. (2.145) we can write

$$\mathcal{F}_k^{\text{in}} | \rho_0) = B^{-1} \tilde{\mathcal{F}}_k^{\text{in}} B | \rho_0) = B \tilde{\mathcal{F}}_k^{\text{in}} | \tilde{\rho}). \quad (2.148)$$

Since we assumed the gates are invertible, then the gate sequences $\tilde{\mathcal{F}}_k^{\text{in}}$ are also invertible and the above equation gives us $B | \rho_0) = | \tilde{\rho}_0)$. An analogous argument leads to $(E_0 | B^{-1} = (\tilde{E}_0 |$, which ends the informational completeness proof.

The full gate set tomography procedure in pyGSTi

After a tomographically complete setting is identified, and linear GST is performed, several additional steps are taken in the standard GST implementation pyGSTi [168, 173]. The first has to do with enforcing physicality constraints, since a linear GST estimate is not constrained to be CPT. This is done via a CPT projection step. Since the output of linear GST is a linear inversion estimate, it minimizes the least-squares error to the data. However, it is often desirable to find the minimum of a more well-motivated cost function, such as the log-likelihood error Eq. (2.75). Consequently, the result of the CPT projection step is subsequently used as an initial point for the non-convex optimization problem given by the log-likelihood cost function. The linear GST estimate thus serves the purpose of a good initialization, which can significantly ease finding the global minimum of the optimization problem. Furthermore, the similarity of the linear GST sequence design and standard QPT allows for a straightforward experiment design with provable informational completeness and experimental verification through checking invertibility of the Gram matrix. This is in contrast to random experiment designs, which although expected to ensure informationally completeness if enough sequences are given, are difficult to analyze and thus far do not come with informational completeness proofs.

An essential technique in pyGSTi is the use of carefully designed long sequences and the ensuring protocol which uses a linear GST estimate as a starting point and then optimizes the log-likelihood on long sequences is called *long sequence GST*. The principal motivation behind long sequences is error amplification. Consider a single qubit Bloch rotation $e^{i(\varphi/2+\delta)\sigma_x}$, where φ is the desired rotation angle and δ is a calibration error. Repeating the unitary M times leads to a rotation given as $e^{iM\varphi/2\sigma_x} e^{iM\delta\sigma_x}$. This means that if a protocol can estimate the rotation angle $M\delta$ up to accuracy $\mathcal{O}(1)$, it can estimate δ up to accuracy $\mathcal{O}(1/M)$. This is termed as *Heisenberg* scaling of the error, and gives a considerable advantage compared to the $\mathcal{O}(1/\sqrt{N})$ shot noise scaling where N is the number of samples. The trade-off is that longer sequences have to be applied, which can only be done until the total circuit time reaches the coherence time of the experiment, i.e.

$Mt_{\text{Gate}} \approx T_2$. Note also that not every gate parameter in a given gate $\mathcal{G}$ is equally amplified by just measuring gate sequences $(E_0 | \mathcal{F}_i^{\text{out}} \mathcal{G}_j^M \mathcal{F}_k^{\text{in}} | \rho_0)$. In general, if $\mathcal{G} = \mathcal{G}^{\text{ideal}} \Lambda$ and Λ commutes with $\mathcal{G}^{\text{ideal}}$, then we get $\mathcal{G}^M = (\mathcal{G}^{\text{ideal}})^M \Lambda^M$ and errors are amplified as desired. A GST experiment design which includes sequences that amplify every non-gauge parameter in all the gates is then termed *amplificationally complete*. As reported in Nielsen *et al.* [168], simply applying random long sequences does not yield the desired Heisenberg scaling and thus the GST sequences have to be determined for each gate set. With a properly designed sequence set, Heisenberg scaling is reliably observed for different gate sets in numerical simulations.

In principle pyGSTi is flexible towards different gate parametrizations, but as per default physicality constraints are preserved by parametrizing the gate models via the Lindblad parametrization (Eq. (2.70)). Since long sequences lead to cost functions which are high degree polynomials in the gate parameters, simply optimizing the log-likelihood on all long sequence GST data is daunting. Therefore, pyGSTi uses an optimization pipeline where a local approximation of the log-likelihood function is first optimized on short sequences, and progressively longer sequences are added with the previous estimate as a starting point. The initial starting point is naturally given by the CPT projected linear GST estimate. Only in the last step is the log-likelihood maximized on the full data (see illustration in [171]). Finally, gauge optimization to the target gate set is performed in order to compare gate set estimates to the experimental implementation.

In a well calibrated experiment, one expects the implemented gates to be close to the target gates, and the latter can therefore be used as an initialization. This circumvents the computation steps required for the initial linear GST estimate: Inversion of the Gram matrix and projection onto the model class.

Alternative self-consistent tomography protocols

We have thus far described the most widespread protocol, pyGSTi, as well as regularized self-consistent tomography [177] and the broader scope of non-Markovian process tomography [167].

In Gu *et al.* [181], the GST estimation problem Eq. (2.129) is considerably simplified by only taking the linear approximation of outcomes probabilities in terms of error terms into account, a framework that builds on the earlier work by Merkel *et al.* [28]. To make this more precise, consider a noisy initial state and a noisy POVM given by $|\tilde{\rho}_0\rangle = (\text{id} + \mathcal{E}_{\text{in}})|\rho_0\rangle$ and $\langle\tilde{E}_0| = \langle E_0|(\text{id} + \mathcal{E}_{\text{out}})$, as well as noisy gates $\tilde{\mathcal{G}}_i = (\text{id} + \mathcal{E}_i)\mathcal{G}_i$ and let

$$\max\{\|\mathcal{E}_{\text{in}}\|, \|\mathcal{E}_1\|, \dots, \|\mathcal{E}_k\|, \|\mathcal{E}_{\text{out}}\|\} < \epsilon. \quad (2.149)$$

The linear approximation of an experiment with the gate sequence s is given by the first two terms in

$$\begin{aligned} \langle\tilde{E}_0|\tilde{\mathcal{G}}_{s_L}\cdots\tilde{\mathcal{G}}_{s_1}|\tilde{\rho}_0\rangle &= \langle E_0|\mathcal{G}_{s_L}\cdots\mathcal{G}_{s_1}|\rho_0\rangle \\ &+ \langle E_0|\left(\mathcal{E}_{\text{out}} + \sum_{i=1}^L \mathcal{G}_{s_L}\cdots\mathcal{E}_{s_i}\mathcal{G}_{s_i}\cdots\mathcal{G}_{s_1} + \mathcal{E}_{\text{in}}\right)|\rho_0\rangle \\ &+ \mathcal{O}(\epsilon^2). \end{aligned} \quad (2.150)$$

Let again $p_s = \langle E_0|\mathcal{G}_{s_L}\cdots\mathcal{G}_{s_1}|\rho_0\rangle$ and $\tilde{p}_s = \langle\tilde{E}_0|\tilde{\mathcal{G}}_{s_L}\cdots\tilde{\mathcal{G}}_{s_1}|\tilde{\rho}_0\rangle$. Given a vector of measured probabilities $\tilde{p}$, we see that

$$\tilde{p} - p = \mathbf{M}(\mathcal{E}_{\text{out}}, \{\mathcal{E}_i\}, \mathcal{E}_{\text{in}}) + \mathcal{O}(\epsilon^2), \quad (2.151)$$

where $\mathbf{M}$ is a linear map that depends on the set I and the target gates. The method in Gu *et al.* [181] then boils down to drawing a random set of sequences that is informationally complete, determining $\mathbf{M}_I$ and measuring $\tilde{p}$. The parametrizations of the error channels can then in linear approximation be given by $\mathbf{M}_I^{-1}(\tilde{p})$ and the resulting method is termed *randomized linear GST*. Based on results of different 5 qubit superconducting devices, where single- and two-qubit noise channels were estimated, randomized linear GST performs comparable to pyGSTi, while using a simpler experimental design and reduced post-processing time. However, as expected from a linear approximation, prediction errors on long circuits are shown to be larger compared to results from pyGSTi.

In Evans *et al.* [182] a related model is used, where the noise is not linearized around the identity, but around a previous noise model given by channels $\{\Lambda_{\text{out}}, \{\Lambda_i\}, \Lambda_{\text{in}}\}$. The noise channels are further treated as random variables with means $\{\bar{\Lambda}_{\text{out}}, \{\bar{\Lambda}_i\}, \bar{\Lambda}_{\text{in}}\}$, where the real noise process is assumed to be close to the mean. For gate noise, this results in $\tilde{\mathcal{G}}_i = \Lambda_i\mathcal{G}_i =$

$(\bar{\Lambda}_i + \mathcal{E}_i)\mathcal{G}_i$ where $\|\mathcal{E}_i\|$ is again assumed to be small. This leads to an analogous first order model to the one described in Eq. (2.150), where ideal gates $\mathcal{G}_i$ are replaced by their average noise implementation $\bar{\Lambda}_i\mathcal{G}_i$. Let θ be the vector comprising all parameters of the noise channels $\{\mathcal{E}_{\text{out}}, \{\mathcal{E}_i\}, \mathcal{E}_{\text{in}}\}$. The authors use the linear approximation together with Bayesian estimation to update a prior noise distribution after each sequence probability p_s is measured. The updates are performed according to the Bayes rule

$$\mathbb{P}(p_s|\theta) \propto \mathbb{P}(\theta|p_s)\mathbb{P}(\theta), \quad (2.152)$$

where the noise parameters θ are drawn from a multivariate Gaussian distribution. The ability to update the noise distribution after each measured sequence allows for an online approach, in which progressively more accurate gate set estimates are given during the gathering of sequence data. An interesting way of obtaining the prior distribution is also demonstrated in Evans *et al.* [182], where data from a previous RB experiment is used to get a prior that is consistent with the average gate fidelities obtained as a result of RB. Furthermore, individual sequences measured during RB can be used for individual updates according to Eq. (2.152). In an experimental demonstration, the Bayesian tomography method was then shown to be sufficiently fast for the characterization of two-qubit noise in a spin qubit system, with the approximation errors due to the linearization being generally lower than the shot noise error. However, since each update step for two qubit was reported to take about 1 – 4 seconds and a much larger set of sequences than needed for tomographic completeness was used, the proposed Bayesian tomography protocol appears to be more costly than pyGSTi or randomized linear GST.

A third recent work which is very relevant for GST is the one by Huang *et al.* [183], where a comprehensive framework for self-consistent tomography is given from the perspective of learning theory. Most notably the first theoretical guarantee for a self-consistent algorithm was given. Let the system parameters again be given by the set $\{E_i, \mathcal{G}_j, \rho_k\}$, where the initial states and POVM elements are prepared without using the unknown gates $\mathcal{G}_i$. The authors make a distinction between learning an *intrinsic* description and learning *extrinsic* behavior of a device. An intrinsic description means that there exists a unitary or anti-unitary transformation $\mathcal{U}$ such that the estimates $\{\hat{E}_i, \hat{\mathcal{G}}_j, \hat{\rho}_k\}$ satisfy

$$\begin{aligned} \|\hat{\rho}_k - \mathcal{U}(\rho_k)\|_1 &\leq \epsilon \\ \|\hat{E}_i - \mathcal{U}^\dagger(E_i)\|_1 &\leq \epsilon \\ \|\hat{\mathcal{G}}_j - \mathcal{U}\mathcal{G}_j\mathcal{U}^\dagger\|_\diamond &\leq \epsilon \end{aligned} \quad (2.153)$$

for all i, j, k and for some $\epsilon > 0$. In contrast, an extrinsic description must simply satisfy that it can predict any circuit outcome probabilities from the gate set up to a given sequence length L to error ϵ .

It is then shown in Huang *et al.* [183] that for a universal gate set, i.e. a gate set that contains a pure state, a universal set of unitaries and a non-trivial POVM, an intrinsic description can indeed be learned up to arbitrary precision. However, the specific algorithm is by no means efficient and involves a lengthy procedure to (i) identify all gates in the gate set which are perfect unitaries, (ii) composing these unitaries to sample from an approximate unitary 2-designs, (iii) finding a pure state in the set $\{\rho_i\}$ via the use of the unitary 2-design, (iv) generating a particular basis on $L(\mathcal{H})$ from the pure state and the 2-design and finally, (v) measuring the POVM elements and gates in this basis.

To learn extrinsic behavior of a gate set, a more practical constructive proof is given. It starts with the assumption which is natural for GST, that we have identified a set of sequences such that we can generate frames on $L(\mathcal{H})$ from initial states and POVMs, in complete analogy to the formalism of fiducials which we have seen previously. The experimental design is thus similar to standard linear GST. By meticulously keeping track of all the errors incurred during the estimation of individual gate set elements and the use of appropriate inequalities, a sample

complexity was given for learning the extrinsic behavior. The result is summarized in the following theorem, where we define k as the number gates, r as the number of states and M as the number of POVM elements in the gate set.

Theorem 9 (Theorem 4 in [183]). *Assuming knowledge of sequences that prepare a frame on the set of states and POVM elements, the algorithm in [183] can predict $\text{Tr}[E_i \mathcal{G}_s(\rho_j)]$ for all E_i, ρ_j and all gate sequences s up to length L with a sample complexity of $\mathcal{O}((L^2 k + r + M)/\epsilon^2)$ up to logarithmic factors with high probability.*

The algorithm used in the proof is more of theoretical nature and has thus far not been implemented or numerically tested. Furthermore, the total sample complexity with all constants and logarithmic terms involved is likely far higher than the number of samples for which an ϵ -error in predictions is observed using existing gate set tomography protocols. It thus remains for future work to combine the proof techniques of Theorem 9 with an efficient algorithm to obtain a practically usable GST protocol with theoretical guarantees.

Another interesting result in Huang *et al.* [183] is that even if two gate sets are not related by a unitary or anti-unitary gauge transformation and if their ideal versions are tomographically complete, simple noise models can render them indistinguishable via measurements. The example given was a single qubit gate set consisting of the H, S and T -gate, acting on a perfect 0-state with a perfect computational basis measurement. By writing all possible circuit outcomes comprised of the $\{H, S, T\}$ gate set under Pauli noise as polynomials of the Pauli-noise parameters and the ideal gate actions, it was shown that bit flip noise and dephasing noise on the Hadamard gate leads to the exact same polynomials. Thus, those two noise models can not be distinguished in a GST experiment. This provides a realistic example of a relevant gauge freedom which lies outside the set of unitary or anti-unitary gauge class.

It can further be argued [168, 183] that GST does not reconstruct the true physical description of the device, even if its assumptions are satisfied. The first reason is that due to sampling error, only a gate set estimate that explains the outcomes with respect to shot noise is given, not the true outcomes. Second and more importantly, due to gauge freedom GST only extracts a joint physical description of the device that does not have to correspond to the true description. Thus, GST can only be seen as a model to accurately predict outcomes for arbitrary circuits and consequently the quality of GST estimates should only be judged by their predictive power.

2.3 Estimating quantum state properties with randomized measurements

The need to efficiently estimate properties of a quantum state is central for many applications of quantum computing as well as for the understanding of the system's physics. Properties of interest include for instance correlation functions, entanglement entropies, the energy with respect to a Hamiltonian, or the fidelity to a target state. If the quantum state whose properties are to be studied can be repeatedly prepared, quantum state tomography would give full information about the system, but the required number of measurement settings and the space to store the result increase exponentially in the system size. Very recently [36, 38], a formalism was developed that can simultaneously estimate exponentially many properties of a quantum state, while the number of required measurement samples only scales linearly in the number of observables and is independent of the system size. This is known as the *classical shadows* or *shadow estimation* formalism after the works by Aaronson [184] and Huang *et al.* [36] or alternatively as the *randomized measurement* formalism after Elben *et al.* [37]. The experimental procedure is given by the simple loop:

- (i) prepare a copy of the state,
- (ii) select a random unitary,

(iii) measure in the computational basis.

These steps are then repeated until sufficiently many samples for a desired target accuracy are obtained. The measurement loop describes a single shot scenario, where each setting defined by the random unitary is only measured once (unless the same unitary is drawn again by chance). What turns this simple measurement protocol into a highly successful technique is the construction of an efficient estimator for different state properties, with theoretical guarantees on the sample complexity that match fundamental lower bounds [36].

The utility of shadow estimation has immediately been demonstrated in experiments [39, 40], and a large amount of applications and generalizations were developed in short order. Among these are improved Monte Carlo simulations [42], machine learning [41] and quantum process tomography [43, 142–144]. Another promising application lies in variational quantum algorithms, which are hybrid classical algorithms that require the preparation of a quantum state and the estimation of its energy with respect to a given Hamiltonian at each iteration. Via the use of classical shadows, there now exist highly optimized methods for this task [185, 186].

In this section we give an introduction to the formalism, following the language of classical shadows in the work by Huang *et al.* [36].

We write the computational basis projectors as $|x\rangle\langle x| = E_x \equiv |E_x\rangle$ for $x \in \mathbb{F}_2^n$. After a computational basis measurement is applied, the post-measurement state $|E_x\rangle$ described by n bits can be stored efficiently. For a given state ρ , we define $p_x(\rho) = (E_x|\rho)$ to be the outcome probability of obtaining x . The expectation value over the post-measurement state is consequently given by

$$\mathbb{E}_x |E_x\rangle = \sum_{x \in \mathbb{F}_2^n} |E_x\rangle p_x(\rho) = \sum_{x \in \mathbb{F}_2^n} |E_x\rangle (E_x|\rho) =: M|\rho\rangle. \quad (2.154)$$

The linear map M is called the *Z-basis measurement operator*. In the shadow estimation protocol, a random unitary operation $\omega(g) = g(\cdot)g^\dagger$ from a group G is applied before each Z-basis measurement. This is equivalent to measuring in the basis $\{\omega(g)^\dagger |E_x\rangle\}_{x \in \mathbb{F}_2^n}$ and storing the post measurement state as a tuple (g, x) . For the storage to be efficient, we require g to have an efficient classical representation, which is the case for the most widely studied instances, where G is either the local or the global Clifford group. Let g be drawn from G according to the measure μ . The expectation value over the post measurement state is then given by

$$\mathbb{E}_{g \sim \mu} \sum_{x \in \mathbb{F}_2^n} \omega(g)^\dagger |E_x\rangle (E_x|\omega(g)|\rho) =: S|\rho\rangle, \quad (2.155)$$

where we call S the frame operator for the given measurement process. The reason for the nomenclature is that we can regard all measurement settings labeled by (g, x) as a single POVM with elements $\{\mu(g)g^\dagger |x\rangle\langle x|g\}_{g \in G, x \in \mathbb{F}_2^n}$, which for suitable μ forms a frame on $L(\mathcal{H})$.

Assuming in round i of the protocol, gate g_i was drawn and the outcome $x_i \in \mathbb{F}_2^n$ was recorded. The *classical* shadow of the state ρ for a protocol with N measurement rounds is defined to be the set of states

$$\{S^{-1}\omega(g_i)^\dagger |E_{x_i}\rangle\}_{i \in [N]}. \quad (2.156)$$

The average over this set of states gives an unbiased estimator of ρ , since

$$\mathbb{E}_{g,x} \frac{1}{N} \sum_{i=1}^N S^{-1}\omega(g_i)^\dagger |E_{x_i}\rangle = \frac{1}{N} \sum_{i=1}^N S^{-1} \mathbb{E}_g \sum_{x \in \mathbb{F}_2^n} \omega(g)^\dagger |E_x\rangle (E_x|\omega(g)|\rho) = S^{-1}S|\rho\rangle = |\rho\rangle. \quad (2.157)$$

We can hence obtain an unbiased estimator for any observable O via the average

$$\hat{o}_N = \frac{1}{N} \sum_{i=1}^N (O|S^{-1}\omega(g_i)^\dagger |E_{x_i}\rangle). \quad (2.158)$$

In Huang *et al.* [36], the *median of means* estimator was used instead of the mean, which provides better concentration for heavy tailed distributions [187]. It is computed by splitting the data into K batches of size N/K (assuming N is divisible by K), calculating the mean of each batch, and finally forming the median over all batch means.

In Helsen and Walter [188] it was shown that for $\mu(g)$ being the Haar measure on the unitary group, the mean can be used instead of the median of means without sacrificing efficiency. The same holds true for uniform sampling from the local Clifford group $\text{Cl}_1^{\times n}$ [189], while for uniform sampling from the n -qubit Clifford group Cl_n , Helsen and Walter [188] show that the distribution of the random variables $(O|S^{-1}\omega(g_i)^\dagger|E_{x_i})$ is heavy tailed and hence using the median of means is warranted.

In order for the estimator in Eq. (2.158) to be practical, we need an efficient method to compute the inverse of the frame operator $S \in \mathbb{C}^{d^2 \times d^2}$. If the gates g are drawn from a group G , we see that the frame operator is just the group twirl given by

$$S = \int \omega(g)^\dagger M \omega(g) d\mu(g). \quad (2.159)$$

To derive an analytical expression for S we can use Schur's Lemma as shown in Section 2.1.2, where the resulting expression is given in terms of projectors onto the invariant subspaces of the representation ω (see Eq. (2.7)). The invariant subspaces of the representation $\omega(g)(X) = gXg^\dagger$ are known to be $\text{span}(\{\sigma_0\})$ and $\text{span}(\{\sigma_a\}_{a \in \mathbb{F}_2^{2n} \setminus 0})$, both with multiplicity 1. The projectors onto these subspaces are given by $\Pi_0 = |\check{\sigma}_0\rangle\langle\check{\sigma}_0|$ and its complement

$$\Pi_{\text{ad}} = \sum_{a \in \mathbb{F}_2^{2n} \setminus 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a|, \quad (2.160)$$

where we wrote normalized Pauli matrices as $\check{\sigma}_a$ (see Section 2.1.1). Using

$$M = \sum_{x \in \mathbb{F}_2^n} |x\rangle\langle x| = (|\check{\sigma}_0\rangle\langle\check{\sigma}_0| + |\check{\sigma}_z\rangle\langle\check{\sigma}_z|)^{\otimes n} \quad (2.161)$$

it is straightforward to verify that the frame operator for sampling from the Haar measure over $U(2^n)$ is given by

$$\begin{aligned} S &= \Pi_0 \text{Tr}[\Pi_0 M] + \frac{1}{d^2 - 1} \Pi_{\text{ad}} \text{Tr}[\Pi_{\text{ad}} M] \\ &= |\check{\sigma}_0\rangle\langle\check{\sigma}_0| + \frac{1}{d+1} \sum_{a \in \mathbb{F}_2^{2n} \setminus 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \\ &= |\check{\sigma}_0\rangle\langle\check{\sigma}_0| + (d+1)^{-1} (\text{id} - |\check{\sigma}_0\rangle\langle\check{\sigma}_0|), \end{aligned} \quad (2.162)$$

where id is the identity channel. The same holds true for uniform sampling over the Clifford group Cl_n . Since the frame operator is diagonal in the Pauli basis, its inverse is easily computed as $S^{-1} = |\check{\sigma}_0\rangle\langle\check{\sigma}_0| + (d+1)(\text{id} - |\check{\sigma}_0\rangle\langle\check{\sigma}_0|) = d|\check{\sigma}_0\rangle\langle\check{\sigma}_0| + (d+1)\text{id} = (d+1)\text{id} - |\sigma_0\rangle\langle\sigma_0|$. As a result, the classical shadows (Eq. (2.156)) are just given as

$$\{(d+1)\omega(g_i)^\dagger|E_{x_i}\rangle - |\sigma_0\rangle\}_{i \in [N]}. \quad (2.163)$$

For a protocol that samples uniformly over the local Clifford group, the frame operator factorizes in the same way as the Z-basis measurement map in Eq. (2.161), and the classical shadow consists of local shadows, where each is given by Eq. (2.156) with $d = 2$. Apart from the local or global Clifford group, random shallow circuits have been proposed for shadow estimation, where the frame operator is either calculated analytically or numerically [190–193].

The sample complexity bounds for classical shadows are reviewed in the following. We call a measure ν over $U(d)$ informationally complete if for all states $\rho \neq \tilde{\rho}$ there exists a $g \in U(d)$ and a $x \in \mathbb{F}_2^n$ for which $\nu(g) \neq 0$ and $(x|\omega(g)|\rho) \neq (x|\omega(g)|\tilde{\rho})$. We further define the traceless part of an observable as $O_0 = O - \text{Tr}(O)\mathbb{1}/d$.

Theorem 10 (Theorem 3 in Huang *et al.* [36]). *Consider a list of observables $O_i, \dots, O_M$, and let $\epsilon > 0$, $\delta \in [0, 1]$, as well as*

$$K = 2 \log(2M/\delta) \text{ and } N = \frac{34}{\epsilon^2} \max_{i \in [M]} \|O_{0,i}\|_{\text{shadow}, \nu}^2. \quad (2.164)$$

Then the median of means estimator $\hat{o}_i$ of NK classical shadows obtained via an informationally complete measure satisfies

$$|\hat{o}_i - \text{Tr}[O\rho]| \leq \epsilon \quad (2.165)$$

for all $i \in [M]$ with probability $1 - \delta$.

The sample complexity depends on the observable, as well as the measure ν , according to which unitaries are drawn in each round. This dependence is governed by the so-called *stabilizer norm*, defined as

$$\|O\|_{\text{shadow}, \nu} := \max_{\rho \in \mathcal{S}(H)} \left(\mathbb{E}_{g \sim \nu} \sum_{x \in \mathbb{F}_2^n} (E_x |\omega(g)|\rho)(E_x |\omega(g)^\dagger S^{-1}|\rho)^2 \right)^{1/2}. \quad (2.166)$$

The Theorem further shows that the number of samples only scales logarithmically in the number of observables to be estimated. This is to be understood in terms of the failure probability δ , in the sense that achieving error ϵ on M observables with probability $1 - \delta$ requires the same number of samples as the estimation of a single observable to error ϵ with probability $1 - \delta/M$. Since the error for the median of means estimator scales logarithmically in δ , we also get a logarithmic scaling in M . The theorem is furthermore agnostic to the list of observables, provided they have low shadow norm, which implies they can be chosen after the data acquisition phase. It was further shown that the sample complexity in Theorem 10 matches fundamental lower bounds up to constant factors [36].

For efficient estimators, we require the squared shadow norm to be independent of the system dimension. Two main examples are given in the original work [36]. The first considers uniform sampling from the global Clifford group, where it is shown that $\|O_{0,i}\|_{\text{shadow}, \nu}^2 \leq 3\|O\|_2^2$. observables with constant 2-norm include all pure states, resulting in the efficient estimation of pure state fidelities $\langle \psi | O | \psi \rangle$ and entanglement witnesses [194]. The second example considers uniform sampling from local Clifford group, for which it is shown that $\|O_{0,i}\|_{\text{shadow}, \nu}^2 \leq 4^{|\text{supp}(O)|} \|O\|_\infty^2$, where $|\text{supp}(O)|$ is the number of qubits on which O acts non-trivially. Shadows obtained from the local Clifford group can thus be used to estimate local observables efficiently. An example of a relevant observable where the sample complexity of both the local and global Clifford protocol becomes exponential is given by a Pauli observable σ_a with $|\text{supp}(a)| = n$. The shadow norm bounds in this case suggest a scaling of $\|O\|_2^2 = 2^n$ for global Cliffords and a scaling of 4^n for local Cliffords.¹

An obstacle that was not taken into account in the original works is the presence of noise, which can make a crucial difference for current and near term quantum devices. In Chen *et al.* [44], a modified protocol called *robust shadow estimation* was proposed, which makes use of an additional calibration experiment to mitigate the noise-induced estimation bias. In the noise model of robust shadow estimation, each gate is only affected by a constant (left-sided) noise channel Λ , such that $\phi(g) = \Lambda\omega(g)$ for all $g \in G$. The resulting noisy frame operator takes the form

$$\tilde{S} = \mathbb{E}_{g \sim \mu} \omega(g)^\dagger \sum_{x \in \mathbb{F}_2^n} |E_x\rangle \langle E_x| \Lambda \omega(g), \quad (2.167)$$

where it can be seen that Λ effectively incorporates measurement noise. This noise model can be easily analyzed using Schur's Lemma, since the noisy frame operator is just a twirl of the channel $M\Lambda$ over the group G . In analogy to Eq. (2.162) we get

¹The scaling for tensor products of 1-local observables can actually be improved to 3^k [36].

$$\tilde{S} = \Pi_0 \text{Tr}[\Pi_0 M \Lambda] + \frac{1}{d^2 - 1} \Pi_{\text{ad}} \text{Tr}[\Pi_{\text{ad}} M \Lambda] = |\check{\sigma}_0\rangle\langle\check{\sigma}_0| + \frac{\text{Tr}[\Pi_{\text{ad}} M \Lambda]}{d^2 - 1} (\text{id} - |\check{\sigma}_0\rangle\langle\check{\sigma}_0|), \quad (2.168)$$

where Λ is assumed to be CPT. The noisy frame operator thus only depends on the single parameter $f := \text{Tr}[\Pi_{\text{ad}} M \Lambda]/(d^2 - 1)$, which is estimated in Chen *et al.* [44] via the following procedure. In each round, the state $|0\rangle$ is prepared, a random gate $\omega(g)$ is applied, and the resulting state is measured in the computational basis, yielding an outcome $x \in \mathbb{F}_2^n$. For each round $i \in [N]$, a single shot estimator of f given by

$$\hat{f}^{(i)} = \frac{d(E_{x_i} |\omega(g_i)|0\rangle) - 1}{d^2 - 1} \quad (2.169)$$

is stored. It can then be shown via familiar representation theory results that $\mathbb{E}\hat{f}^{(i)} = f$. After the calibration experiment, the noisy frame operator is known and an improved shadow estimation protocol where S^{-1} is replaced by $\tilde{S}^{-1} = |\check{\sigma}_0\rangle\langle\check{\sigma}_0| + \hat{f}^{-1}(\text{id} - |\check{\sigma}_0\rangle\langle\check{\sigma}_0|)$ can be performed. A modified procedure is also given to estimate the noisy frame operator for shadow estimation with the local Clifford group, where the two irreducible representations of each local Clifford group lead to 2^n irreducible representation of the n -fold tensor product. As a consequence, the noisy frame operator depends on 2^n parameters in this case, each of which can be estimated efficiently. The robust shadow estimation protocol comes with concrete theoretical guarantees. For the global Clifford protocol, the bias of the robust shadow estimator $\hat{o}_{\text{RS}}$ is bounded as

$$|\mathbb{E}\hat{o}_{\text{RS}} - \text{Tr}[O\rho]| \leq (\epsilon + r)\|O\|_\infty, \quad (2.170)$$

provided $\mathcal{O}(\epsilon^{-2}(\text{Tr}[M\Lambda]/d)^{-2})$ samples are used in the calibration procedure and the state $|0\rangle$ can be prepared with infidelity at most r . Note that $(\text{Tr}[M\Lambda]/d)^{-2} = \mathcal{O}(1)$. An analogous bound is given for the local Clifford protocol [44].

In Chapter 4 and [47], we consider a less restricted noise model and show that this can lead to problems for both the vanilla shadow estimation protocol and for robust shadows.

Chapter 3

Compressive gate set tomography

As motivated in Section 2.2.4, GST allows for the self-consistent characterization of a set of quantum gates, initial state preparation and measurement operators and has become a standard tool for the improvement of current NISQ devices. Previous works on GST however each have serious shortcomings. The standard implementation pyGSTi [168] uses a large amount of measurement settings and post-processing time. This makes GST difficult to implement on platforms with long single shot measurements times, while the classical post-processing time limits the use of GST to at most two-qubit gates. Other recent proposals [181, 182] simplify the post-processing problem by assuming only small deviations between a prior theoretical model and the experimental implementation, thus limiting themselves to near perfect implementations and relaxing the notion of self-consistency. With compressive GST [45] we take a different approach: Instead of assuming prior knowledge about the concrete gate implementation, we offer the possibility to reconstruct a low complexity approximation to the gates implemented in an experiment, where complexity is quantified by the Kraus rank. This is physically well-motivated and allows us to significantly reduce the data acquisition and post-processing time.

In this chapter, we give an overview of our work done on compressive GST. The material presented here mostly follows the publication [45], which is also included in Appendix A. Characterization results presented in this chapter are obtained through our manifold optimization algorithm for GST called *mGST*, which is derived in [45] and implemented in our Python package [46]. We start by explaining the framework of the compressive GST approach in Section 3.1 where we also showcase its performance via numerical simulations. Thereafter, in Section 3.2 we present recent and unpublished results on the use of *mGST* to characterize single- and two-qubit gates in an ion trap experiment. Finally, in Section 3.3, we present an error mitigation scheme for shadow estimation that uses GST results, which was also part of [45].

3.1 The compressive GST framework and algorithm

Our starting point is a gate set as defined in Eq. (2.128), where w.l.o.g. we assume that there is a fixed initial state ρ that is prepared in each experiment. The gate set is thus given by

$$\mathcal{X} = ((E_1, \dots, E_M), (\mathcal{G}_1, \dots, \mathcal{G}_k), \rho). \quad (3.1)$$

As discussed in Section 2.2.4, experimental access to gate set parameters is gained through outcome probabilities, which are, up to shot noise, entries of the probability tensor given by

$$p_{i,s_1,\dots,s_L} = (E_i | \mathcal{G}_{s_L} \cdots \mathcal{G}_{s_1} | \rho). \quad (3.2)$$

In this formulation GST is a tensor completion problem, as it solves the question of how many entries are required to reconstruct the full tensor. Tensors with low rank can be provably reconstructed from a number of entries that is far lower than the total dimension of the tensor [195, 196]. In our case low rank is to be understood in the spirit of matrix product states: Each

tensor can be decomposed into a product of matrices with fixed dimension, which can be seen in the tensor network picture of a gate sequence on the left side of Figure 3.1. However, in our case the matrix dimensions are still of size d^2 and decreasing them would amount to reducing the size of the Hilbert space, which is not expected to yield a meaningful approximation. As we have seen in Section 2.2.3, the natural low rank structure of quantum gates is given through the Kraus decomposition, since the ideal implementation is of rank 1 and typical noise models are of low rank or at least approximately of low rank. In our formulation we therefore treat GST as a completion problem of a structured tensor with low Kraus bond dimensions.

We regard the set of gates $(\mathcal{G}_1, \dots, \mathcal{G}_k)$ as a single tensor of size $k \times d^2 \times d^2$, and we write its Kraus decomposition as

$$\mathcal{G}_i(\cdot) = \sum_{l=1}^{r_K} \mathcal{K}_{il}(\cdot) \mathcal{K}_{il}^\dagger, \quad (3.3)$$

with $\mathcal{K}_i \in \mathbb{C}^{r_K \times d \times d}$. We can similarly approximate initial state and POVM elements to be almost pure, i.e.

of low matrix rank: $E_j = A_j^\dagger A_j$ and $\rho = BB^\dagger$ with $A_j \in \mathbb{C}^{r_E \times d}$ and $B \in \mathbb{C}^{d \times r_\rho}$. The resulting tensor $p_{j,s_1,\dots,s_L}$ is depicted on the right side of Figure 3.1.

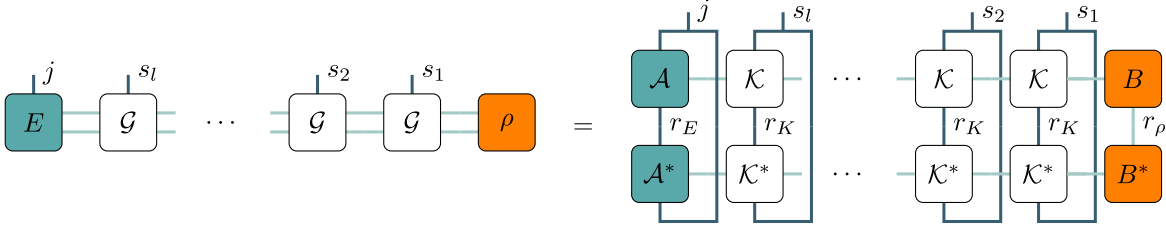


Figure 3.1: Low rank factorization of the sequence tensor to be estimated via GST.

This factorization is possible since gates, initial state and POVM are assumed to be positive. We hence naturally incorporate all positivity constraints. What is still left are the trace constraints, i.e.

unit trace of the state, trace preservation of the gates and the POVM condition $\sum_{j=1}^M E_j = 1$. Trace preservation for all gates is satisfied if $\sum_{l=1}^{r_K} \mathcal{K}_{il}^\dagger \mathcal{K}_{il} = \mathbb{1}$ for all $i \in [k]$. We can rewrite these equations by stacking the matrices $\mathcal{K}_{il}$ together along the l -index, for example $\{\mathcal{K}_{il} \in \mathbb{C}^{d \times d}\}_{l=1}^{r_K}$ turns into $\mathcal{K}_i \in \mathbb{C}^{r_K d \times d}$. The trace preservation constraints are thus simply given by the equations $\mathcal{K}_i^\dagger \mathcal{K}_i = \mathbb{1}$, meaning each Kraus tensor $\mathcal{K}$ lives on a complex Stiefel manifold (see Section 2.1.3). An analogous argument can be made for the state and for POVM elements, and the resulting constraints are summarized in Table 3.1. This realization allows us to combine the manifold

	POVM	State	Gates
Constraint	$\sum_j E_j = \mathbb{1}$	$\text{Tr}(\rho) = 1$	$\mathcal{G}_i$ CPT
Factorization	$\sum_j A_j^\dagger A_j = \mathbb{1}$	$\text{Tr}(BB^\dagger) = 1$	$\sum_{l=1}^{r_K} \mathcal{K}_{il}^\dagger \mathcal{K}_{il} = \mathbb{1}$
Stiefel element	$A^\dagger A = \mathbb{1}$	$\text{vec}(B)^\dagger \text{vec}(B) = 1$	$\mathcal{K}_i^\dagger \mathcal{K}_i = \mathbb{1}$
Manifold	$\text{St}(r_E M, d)$	$\text{St}(r_\rho d, 1)$	$\text{St}(r_K d, d)$

Table 3.1: Physicality constraints reformulated as constraints to complex Stiefel manifolds.

optimization techniques reviewed in Section 2.1.3 with the compressed sensing model of a low rank gate set.

To get a rough estimate of the number of unique sequences that need to be measured, we can count the parameters of a gate set with a single initial state, a POVM with M elements and a set of k gates. First, note that a $d \times d$ complex Hermitian matrix is given by d^2 real numbers.

The same holds for a $d \times d$ unitary matrix, since it can be defined via a hermitian generator. Subtracting the unit trace constraint, we therefore have $d^2 - 1$ free parameters for the initial state without any low rank approximation. In the following we assume all gates have the same Kraus rank r_K . The matrix of Kraus operators $\mathcal{K}_i$ corresponding to gate i has $2d^2 r_K$ real parameters given by the real and imaginary part of each entry. The Stiefel manifold constraint $\mathcal{K}^\dagger \mathcal{K} = \mathbb{1}$ gives us d real constraints for the diagonal of $\mathbb{1}$ and $d(d-1)/2$ complex constraints for the off-diagonal elements. Moreover, we can apply an $r_K \times r_K$ unitary to the Kraus index of $\mathcal{K}$ without changing the measurement outcomes. This corresponds to applying a unitary and its inverse along the vertical indices in Figure 3.1. This unitary freedom removes another r_K^2 real parameters from $\mathcal{K}$ so that we are left with

$$2d^2 r_K - 2d(d-1)/2 - d - r_K^2 = d^2(2r_K - 1) - r_K^2 \quad (3.4)$$

parameters. For the special case of full Kraus rank $r_K = d^2$ this leads to $d^4 - d^2$ real parameters, which coincides with the number of real parameters used to describe the Pauli transfer matrix of a gate. For the POVM-tensor $\mathcal{A} \in \mathbb{C}^{Mr_E \times d}$ the condition $\mathcal{A}^\dagger \mathcal{A} = \mathbb{1}$ gives us d^2 constraints. Additionally, since each POVM element is factorized via $E_i = A_i^\dagger A_i$, there is a unitary freedom for each $i \in [M]$, removing Mr_E^2 parameters. Therefore, the POVM with rank r_E elements is described by $2dMr_E - d^2 - Mr_E^2$ real parameters, which reduces to $Md^2 - d^2$ for full rank POVMs. Since the rank of $\sum_{i=1}^M E_i$ with each E_i of rank r_E is at most Mr_E , while $\mathbb{1}$ is of rank d , a valid POVM needs to satisfy $Mr_E \geq d$. For the initial state parametrization $\rho = BB^\dagger$ with $B \in \mathbb{C}^{d \times r_\rho}$, we need $2dr_\rho - r_\rho^2$ parameters. Finally, when we consider a unitary gauge freedom we remove another d^2 real parameters. In summary, for a gate set with k gates this leads us to a total of

$$\underbrace{k(d^2(2r_K - 1) - r_K^2)}_{\text{gates}} + \underbrace{d(2Mr_E - d) - Mr_E^2}_{\text{POVM}} + \underbrace{2dr_\rho - r_\rho^2}_{\text{state}} - \underbrace{d^2}_{\text{gauge}} \quad (3.5)$$

free parameters which can be probed through experiments. As an example, for a two qubit system with $d = 4, k = 6, M = 4, r_K = r_E = r_\rho = 2$ this amounts to 292 parameters, while the same system with full ranks ($r_K = 16, r_E = r_\rho = 4$) is described by 1488 parameters.

With an understanding of our chosen parametrization, we now turn to the question of how these parameters can be learned from measurement data. Let $\hat{p}_{j,s}$ be the outcome probabilities estimated from repeated measurements of the circuits $s \in I$ and let $f(\mathcal{A}, \mathcal{K}, B|\hat{p})$ be a cost function that quantifies the discrepancy between model predictions $p_{j,s}(\mathcal{A}, \mathcal{K}, B)$ and estimated probabilities $\hat{p}_{j,s}$. We choose this to be the least-squares error (Eq. (2.99)) for our main GST algorithm, with the option to improve the final least-squares estimate further by running a subsequent likelihood error (Eq. (2.75)) optimization. The estimation problem for compressive GST is thus given by

$$\begin{aligned} & \underset{\mathcal{A}, \mathcal{K}, B}{\text{minimize}} && f(\mathcal{A}, \mathcal{K}, B|\hat{p}) \\ & \text{subject to} && \mathcal{A} \in \text{St}(r_E M, d), \quad \mathcal{K}_i \in \text{St}(dr_K, d), \quad B \in \text{St}(dr_\rho, 1). \end{aligned} \quad (3.6)$$

Here the choice also arises as to whether we want to jointly optimize on the full product manifold $\text{St}(r_E M, d) \times \text{St}(r_K d, d)^{\times k} \times \text{St}(r_\rho d, 1)$ or take an alternating minimization approach. Through trial and error, we found that jointly optimizing over the product manifold of gates and alternating between POVM, gates and initial state leads to fast convergence on single- and two qubit examples. The resulting main routine of the optimization algorithm mGST is given in Algorithm 1. Since the literature on optimization algorithms for the complex Stiefel manifold is sparse and the GST estimation problem involves a non-convex cost function given by the higher order polynomial $f(\mathcal{A}, \mathcal{K}, B|\hat{p})$, we formulate our own tailored optimization procedure to find update directions. It combines the saddle free Newton method of Dauphin *et al.* [197] with second order optimization on the complex Stiefel manifold. To formulate this method, we first derive the geodesic equation and prove that a geodesic ansatz formed by generalizing the real

Algorithm 1: mGST - main routine

input: Data $\{\hat{p}_{j,s}\}_{j \in [M], s \in I}$, Ranks r_K, r_E, r_ρ , batch size κ , initialization $(\mathcal{A}^0, \mathcal{K}^0, B^0)$, stopping criterion

```

1  $i \leftarrow 0$ 
2 repeat
3   Select batch  $J \subset I$  of size  $|J| = \kappa$  at random
4    $\mathcal{A}^{i+1} \leftarrow$  update  $\mathcal{A}^i$  with objective  $f(\cdot, \mathcal{K}^i, B^i | \hat{p}_{j,s})$  along geodesic on  $\text{St}(Mr_E, d)$ 
5    $\mathcal{K}^{i+1} \leftarrow$  update  $\mathcal{K}^i$  with objective  $f(\mathcal{A}^{i+1}, \cdot, B^i | \hat{p}_{j,s})$  along geodesic on  $\text{St}(r_K d, d)^{\times k}$ 
6    $B^{i+1} \leftarrow$  update  $B^i$  with objective  $f(\mathcal{A}^{i+1}, \mathcal{K}^{i+1}, \cdot | \hat{p}_{j,s})$  along geodesic on  $\text{St}(dr_\rho, 1)$ 
7    $i \leftarrow i + 1$ 
8 until stopping criterion is met;
9 return  $(\mathcal{A}^i, \mathcal{K}^i, B^i)$ 

```

valued case satisfies this equation. Equipped with an expression for the geodesic, we derive the Riemannian Hessian via the use of Proposition 8. The full calculations can be found in Appendix A.

By the standard assumption in GST and QPT that the underlying dynamics change negligibly between different shots, we know that the measurement outcomes for a given sequence follow a multinomial distribution. Hence, the estimation error for $\hat{p}_{i,s}$ given a fixed number of shots is known. We use this information to define a stopping criterion that lets the algorithm terminate if the least-squares error approaches the expected least-squares error given the shot noise. If the algorithm converges to a local minimum that does not satisfy the convergence criterion we restart it with a new random initialization. For the random gate initialization, we first form a random Hermitian matrix $H \in \mathbb{C}^{r_K d \times r_K d}$ with real and imaginary part of each independent entry drawn from the normal distribution with zero mean and unit variance. Taking only the first d columns of e^{iH} gives us a random element of $\text{St}(r_K d \times d)$. Initial state and POVM are handled analogously.

We implement Algorithm 1 in Python, using the package Numba for compiling performance-critical functions into optimized machine code. The full Python package [46] contains an example notebook included in Appendix D, where the principal use of the algorithm for a qubit and a qutrit example is showcased.

The performance of mGST is studied in numerical simulations in our article [45], where the plots in Figure 3.3 originate from. In these plots the promise of our compressed sensing approach can be observed: For a given generalization error accuracy (measured by the MVE), a low rank model requires fewer measurement settings. This behavior is clearly visible for the single qubit case, while for two qubits we observe that the rank 1 optimization runs into local minima significantly more often for a lower sequence count. This leads to a competing effect to the compressed sensing advantage of low rank: It is possible to obtain a low error estimate for fewer sequences, but the post-processing time increases. For the plots in Figure 3.2 we limit the maximal algorithm runtime, and hence the results indicate what error can be achieved with both a limited measurement and a limited classical computation effort.

The question also arises as to what happens when there is a model mismatch, in the sense that the true gate is not of low rank. In Figure 3.3 we analyze the case where there is a large amount of depolarizing noise, which is of full rank. We see that using the unitary ($r_K = 1$) approximation, the algorithm terminates at a high cost function value, indicating a poor fit to the data. Up to 10^3 shots per sequence, the shot noise error dominates the model mismatch error, and ranks 2 – 4 fit the data equally well. The benefit of using the full rank $r_K = 4$ only becomes apparent for a large number of shots, which speaks for the adequacy of lower rank models in realistic scenarios where the number of shots is limited.

In [45] we further extensively compare mGST to pyGSTi for different noisy single qubit gate

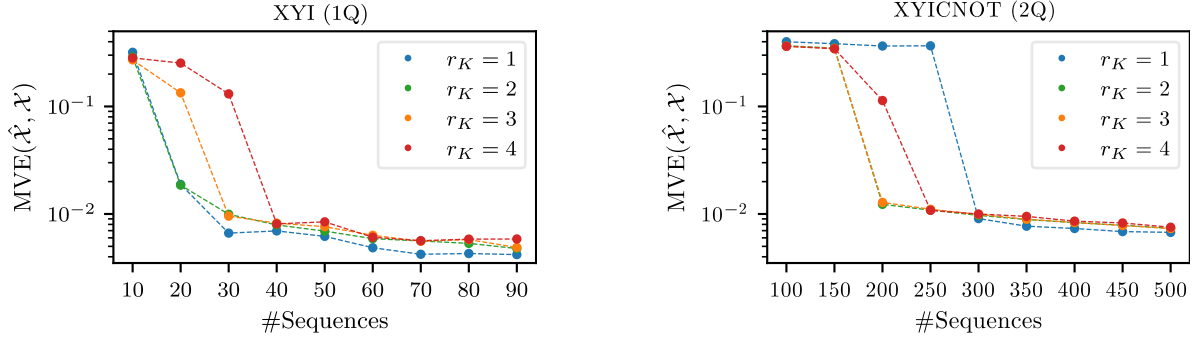


Figure 3.2: Mean variation error (MVE) between the GST estimate $\hat{\mathcal{X}}$ and the true gate set $\mathcal{X}$, evaluated on random sequences which are not in the data set. The XYI model is comprised of $\pi/2$ Bloch rotations around the X and Y axes on the Bloch sphere, as well as the identity gate. The XYICNOT model consists of the same X and Y rotations on two qubits, as well as the CNOT gate. Sequence probabilities $\hat{p}_{i,s}$ were estimated from 1000 shots for each sequence. Data points are the median over 10 repetitions, where each repetition uses a new random sequence set.

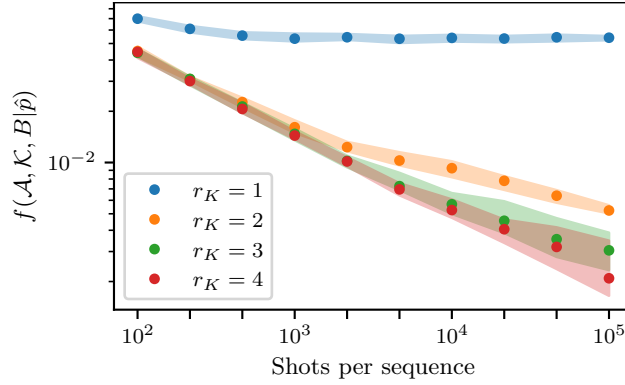


Figure 3.3: Final cost function over the number of samples. The gates are each subject to depolarizing noise with magnitude $p = 0.1$ and small random overrotations, while the measurement is subject to depolarizing noise with $p = 0.001$. Bands indicate the limits of the 25th and 75th percentile out of all data for a fixed sample count, while markers indicate the median.

sets in the short sequence regime. We find that the achieved accuracy is identical for standard settings, while mGST fares better for high amplitude damping noise and generally for random gate sets. Furthermore, mGST uses a low number of random sequences and uses less classical post-processing resources due to the low rank approximation. This allows us to perform GST in numerical simulations for a 3 qubit example. For the same example and the same measurements, pyGSTi did not produce any outcome within 4 1/2 days of computing time, and we therefore deem pyGSTi unsuccessful for this task.

Throughout the numerical simulations in [45] and the application of mGST for the characterization of gates in ion traps which will be presented in Section 3.2, compressive GST emerges as a very practical alternative to previous GST implementations. Additionally, there remain a few opportunities for improvement: First, long error amplification sequences as described by Nielsen *et al.* [168] could be incorporated in the optimization to achieve a $1/L$ scaling in the final error. Second, the proof methods for recovery guarantees in Huang *et al.* [183], which came out after our work, could also inspire recovery guarantees for compressive GST in the random sequence setting. And third, a reduction in post-processing time could be achieved by avoiding the computation of the Hessian for the Newton method, but it remains to be shown that this

can be done without sacrificing the convergence properties of the algorithm. Improvements are currently also underway to the python package, where we aim to increase the user-friendliness by fully automating the running of mGST and the generation of report files such as the ones shown in Appendix C. Thus far, the package is to be used via a python script or in a Jupyter notebook, and in Appendix D we include the tutorial notebook from the mGST package, which comes with running instructions and basic examples.

3.2 Characterizing quantum gates on a trapped ion system using mGST

In this section we summarize unpublished results of compressive GST performed on the linear Paul trap in the group of Prof. Wunderlich at the University of Siegen. GST was performed in the context of a joint research project (MIQRO) with the aim of developing a universal ion trap-based quantum computer. The experiment uses $^{171}\text{Yb}^+$ ions which are trapped with an axial frequency of 89 kHz and are Doppler cooled to form a Coulomb crystal. Qubits are encoded in the hyperfine states of the ground state $^2\text{S}_{1/2}$, forming a Spin 1/2 system. Along the trapping axis a static magnetic field gradient of 17 T/m is applied, resulting in different addressing frequencies of each qubit near 12.64 GHz. This allows for individual single qubit operations on each qubit. Entangling operations use pairwise spin interactions according to the Hamiltonian

$$H = \sum_{i \neq j} J_{ij} \sigma_z^{(i)} \otimes \sigma_z^{(j)}, \quad (3.7)$$

where J_{ij} is the coupling strength between qubits at sites i and j . This interaction is present due to the magnetic field gradient, and is therefore termed magnetic gradient induced coupling (MAGIC). By encoding the qubits in different hyperfine levels, the sign of the coupling strength can be changed or spin-spin couplings can be entirely suppressed, leading to modifiable interaction graphs. For a more detailed description of the experimental setup and its prospect for quantum computing see Piltz *et al.* [198].

In order to reliably assess the performance of the system and characterize error sources, we perform GST with varying Kraus ranks on single- and a two-qubit subsystem. Since the overall number of shots is very limited, mostly due to the cooling time required to prepare the initial state, a method that can reconstruct a gate set from few sequences is needed in this scenario.

To quantify the effect of shot noise, we use non-parametric bootstrapping to get error bars on our estimates. The idea behind bootstrapping is as follows. Assume we are given outcome probabilities $\hat{p}_{j,s}^\#$ for POVM elements $j \in [M]$ after sequences s are applied, where each sequence is measured m times to estimate $\hat{p}_{j,s}^\#$. Since we assume that the underlying experiment does not change between shots and outcomes are given by the Born rule, each shot represents an independent draw from the multinomial distribution $(p_{1,s}, \dots, p_{M,s})$ with replacement. In non-parametric bootstrapping we take $(\hat{p}_{1,s}^\#, \dots, \hat{p}_{M,s}^\#)$ obtained from the experiment and draw m new samples from it. Using these samples we get a new empirical distribution $\hat{p}_{1,s}^1, \dots, \hat{p}_{M,s}^1$. This process is then repeated N times such that we have a list of artificially created datasets, on each of which we run the full GST procedure. From all N GST estimates obtained in this way, a confidence interval in, say, the average gate fidelity, is given by the interval in which 95% of average gate fidelities computed from the N artificial datasets lie.

The advantages of bootstrapping are its simplicity and the fact that all potential inaccuracies incurred during GST, i.e. from not fully converged fits or suboptimal results of the gauge optimization, are accounted for. The main disadvantages of bootstrapping is that it is time-consuming. Furthermore, we observe that for a low number of shots per sequence, artificial datasets $(\hat{p}_{1,s}^i, \dots, \hat{p}_{M,s}^i)$ do not admit as good a fit with a physical gate set model compared to the original dataset $(\hat{p}_{1,s}^\#, \dots, \hat{p}_{M,s}^\#)$.

For all results presented in this section, we obtain physical estimates of the gate sets using mGST [46] and a subsequent gauge optimization is done via the gauge optimizer in the pyGSTi [172] package. As discussed in Section 2.2.4, the set of allowed gauge operations is different for each point in parameter space and as of yet there is no gauge optimization algorithm that takes the full gauge set into account. We therefore take what appears to be the prudent approach and only gauge-optimize over the unitary group, which is expected to cover most gauge operations for near-unitary gates and leads at worst to a suboptimal gauge choice.

We represent gate estimates via Pauli transfer matrices (Eq. (2.61)), for which common single qubit noise models are defined in Table 2.1. For rank 1 (unitary) gate estimates $\hat{U}_i$ we compute the Hamiltonians as $\hat{H}_i = i \log(e^{i\varphi} \hat{U}_i)$ where the principal logarithm is chosen. The Hamiltonians are then represented in the Pauli basis as

$$H_i = \sum_{a \in \mathbb{F}_2^{2n}} c_{i,a} \sigma_a = \varphi \mathbb{1} + \frac{\alpha_i}{2} \sum_{a \neq 0} n_{i,a} \sigma_a, \quad (3.8)$$

where we set $n_{i,a} = c_{i,a} / \|c_i\|_{\ell_2}$ and $\alpha_i = \|c_i\|_{\ell_2}$. For a single qubit, this allows us to interpret the vector n_i as the rotation axis on the Bloch sphere and α_i as the rotation angle. The additional phase φ is just a global phase and can thus be ignored. For more than one qubit where the Bloch sphere picture does not apply, we can interpret $\alpha_i n_{i,a}$ as the total rotation angle of the σ_a interaction and compare it to the target angle.

To analyze noise contributions for arbitrary Kraus ranks, we use the left-error model $\mathcal{G}_i = \Lambda_i \mathcal{U}_i^{\text{target}}$. Note that there always exists a left noise channel that satisfies this condition, but the representation is not unique, as one could analogously define right-sided or two-sided noise. In reality, different interactions happen simultaneously, and the resulting process does not generally factorize into a target unitary interaction and a noise map. The left sided noise model nonetheless provides valuable information and is commonly used in the literature. Now let $\mathcal{E}(p)$ be a given noise model with parameter p such as for instance dephasing noise. To assess how dominant the contribution of $\mathcal{E}(p)$ is, we find the noise parameter p that best describes Λ_i via the minimization problem

$$\min_p \|\mathcal{E}(p) - \Lambda_i\|_F. \quad (3.9)$$

For the purpose of the trapped ion system analyzed here, local Z-dephasing noise given by $\mathcal{E}(p)(\rho) = (1-p)\mathbb{1} + p\sigma_Z\rho\sigma_Z$ (see also Table 2.1) is the most relevant.

3.2.1 Single qubit GST

The gate set used for single qubit GST is given by

$$\mathfrak{G} = \left(\text{Idle}(T), \text{Idle}(2T), e^{-i\frac{\pi}{2}\sigma_x}, e^{-i\frac{\pi}{2}\sigma_y}, e^{-i\frac{\pi}{4}\sigma_x}, e^{-i\frac{\pi}{4}\sigma_y} \right), \quad (3.10)$$

together with the initial state $\rho = |0\rangle\langle 0|$ and the computational basis measurement $E_0 = |0\rangle\langle 0|, E_1 = |1\rangle\langle 1|$. The idle gates of different lengths were added to characterize the background noise present when no control operation is applied. Here the gate time T is the same overall time to apply the $\pi/2$ Bloch rotation given by $e^{-i\frac{\pi}{4}\sigma_x}$. To perform GST on this gate set, we use 200 random sequences of lengths $L = 0, \dots, 24$, where each gate in each sequence is drawn uniformly at random from $\mathfrak{G}$. In Table 3.2, key gate quality measures which were computed from $r_K = 4$ mGST estimates are shown.

During one day of measurement acquisition, an average of 190 shots per sequences were used to estimate the outcome probabilities $(\hat{p}_{1,s}^\#, \dots, \hat{p}_{M,s}^\#)$. We generally find average gate fidelities between 0.9957 and 0.9998. The bootstrapping error bars indicate that due to the relatively low number of shots, average gate fidelities can only be estimated up to an error of 0.005 and a higher number of shots is required to reliably certify average gate fidelities of 0.999 or higher. We further compute the diamond distance to the target gates, which is generally more sensitive to

	Average gate fidelity $\mathcal{F}_{\text{avg}}(\mathcal{U}_i, \hat{\mathcal{G}}_i)$	Diamond distance $\frac{1}{2} \ \mathcal{U}_i - \hat{\mathcal{G}}_i\ _{\diamond}$	Unitarity $u(\hat{\mathcal{G}}_i)$
Idle-short	0.9992 [0.9971,0.9999]	0.0165 [0.0116,0.0265]	0.9972 [0.9890,1.0000]
Idle-long	0.9950 [0.9924,0.9978]	0.0197 [0.0136,0.0309]	0.9807 [0.9716,0.9923]
Rx(pi)	0.9998 [0.9968,0.9999]	0.0081 [0.0045,0.0200]	0.9995 [0.9877,1.0000]
Ry(pi)	0.9998 [0.9981,0.9999]	0.0180 [0.0108,0.0288]	1.0000 [0.9933,1.0000]
Rx(pi/2)	0.9988 [0.9970,0.9997]	0.0209 [0.0116,0.0284]	0.9964 [0.9891,1.0000]
Ry(pi/2)	0.9957 [0.9934,0.9988]	0.0174 [0.0147,0.0276]	0.9834 [0.9744,0.9960]

Table 3.2: Gate quality measures for the standard experimental setup tested on the 15th of January 2024. Confidence intervals encompass 95th percent of results over 50 bootstrapping runs.

coherent errors [71]. If we use the diamond distance to compare the results presented in Table 3.2 to results of an improved experimental setup Table 3.3, we can more reliably certify an across the board improvement for in the implementation of the X- and Y-rotations.

	Average gate fidelity $\mathcal{F}_{\text{avg}}(\mathcal{U}_i, \hat{\mathcal{G}}_i)$	Diamond distance $\frac{1}{2} \ \mathcal{U}_i - \hat{\mathcal{G}}_i\ _{\diamond}$	Unitarity $u(\hat{\mathcal{G}}_i)$
Idle-short	0.9977 [0.9961,0.9990]	0.0207 [0.0150,0.0281]	0.9918 [0.9859,0.9972]
Idle-long	0.9972 [0.9948,0.9979]	0.0221 [0.0185,0.0287]	0.9897 [0.9810,0.9931]
Rx(pi)	0.9989 [0.9969,0.9999]	0.0048 [0.0029,0.0123]	0.9958 [0.9880,0.9999]
Ry(pi)	0.9991 [0.9977,1.0000]	0.0060 [0.0039,0.0129]	0.9965 [0.9911,1.0000]
Rx(pi/2)	0.9994 [0.9980,1.0000]	0.0069 [0.0040,0.0155]	0.9977 [0.9919,1.0000]
Ry(pi/2)	0.9986 [0.9960,0.9999]	0.0079 [0.0053,0.0158]	0.9944 [0.9843,0.9997]

Table 3.3: Gate quality measures for the improved experimental setup using a digital up-converter, tested on the 11th and 12th of January 2024. Confidence intervals encompass 95th percent of results over 50 bootstrapping runs.

For the error measures shown in Table 3.2 and Table 3.3, we require a full rank estimate of the gates. To learn about coherent errors we can run mGST for $r_k = 1$ to obtain the best unitary approximations to each gate, while saving on post-processing time.

	Rotation angle $/\pi$	Axes tilt vs. target (in $^{\circ}$)	Axes estimation error (in $^{\circ}$)
Idle-short	0.0105 [0.0067,1.9903]	–	27.2264
Idle-long	0.0110 [0.0095,1.9896]	–	7.3257
Rx(pi)	0.9994 [0.9949,0.9999]	0.1603	0.0498
Ry(pi)	0.9990 [0.9965,0.9999]	0.2884	0.0961
Rx(pi/2)	0.5051 [0.4989,0.5089]	0.2031	0.5884
Ry(pi/2)	0.5005 [0.4966,0.5024]	0.4722	0.6220

Table 3.4: Single qubit rotation angle and axes tilt for the digital up-converter setup, with errors computed from the 95th percentile over 50 bootstrapping runs.

In Table 3.4 we summarize the results using the single qubit gate parametrization via the Bloch rotations outlined in Eq. (3.8). We observe that the rotation angles and rotation axes of the $\pi/2$ rotations agree with the ideal angles within error bars. For the π rotations, a slight underrotation and a slight axis tilt is observed.

The corresponding full reports, which include Hinton diagrams of all gates in the Pauli basis and state preparation/measurement in the standard basis as well as additional quality measures,

can be found in Appendix C.

3.2.2 Two qubit GST

For the two qubit GST experiment we choose the gate set

$$\mathfrak{G} = \left(\mathbb{1}_4, e^{-i\frac{\pi}{4}\sigma_x} \otimes \mathbb{1}_2, e^{-i\frac{\pi}{4}\sigma_y} \otimes \mathbb{1}_2, \mathbb{1}_2 \otimes e^{-i\frac{\pi}{4}\sigma_x}, \mathbb{1}_2 \otimes e^{-i\frac{\pi}{4}\sigma_y}, e^{-i\frac{\pi}{4}\sigma_z \otimes \sigma_z} \right), \quad (3.11)$$

again in combination with the initial state $|00\rangle$ and the POVM given by projectors onto the computational basis states. The $\sigma_z \otimes \sigma_z$ rotation used as the entangling gate has a gate time of 3.0147ms, which is about two orders of magnitude longer than gate times for single qubit rotations. In Table 3.5, gate quality measures resulting from a full rank mGST estimate are shown. The dataset used for mGST consists of 360 measured sequences with an average of 118 shots per sequence, where sequences are drawn uniformly at random up to a maximum length of 16. Additionally, effective local dephasing parameters computed according to Eq. (3.9) are

	$\mathcal{F}_{\text{avg}}(\mathcal{U}_i, \hat{\mathcal{G}}_i)$	$\frac{1}{2} \ \mathcal{U}_i - \hat{\mathcal{G}}_i\ _{\diamond}$	$u(\hat{\mathcal{G}}_i)$	Dephasing probability of Λ_i for qubit 1	Dephasing probability of Λ_i for qubit 2
Idle	0.9649	0.1270	0.9140	0.0075	0.0301
Rx(pi/2):0	0.9927	0.0308	0.9815	0.0024	0.0050
Ry(pi/2):0	0.9786	0.0766	0.9466	0.0105	0.0114
Rx(pi/2):1	0.9967	0.0442	0.9934	0.0018	0.0012
Ry(pi/2):1	0.9746	0.0788	0.9371	0.0143	0.0115
$e^{-i\frac{\pi}{4}Z \otimes Z}$	0.6852	0.5056	0.4164	0.1481	0.2597

Table 3.5: Gate quality measures for the two qubit full rank reconstruction.

added, since we find local dephasing noise to be the principal noise source in the experiment. The single qubit gates are modeled as two qubit gates in the optimization to also detect potential cross talk errors, and average gate fidelities are thus expected to be higher than for the single qubit gates studied in Table 3.3. We find that the local σ_x rotations have notably higher fidelities and lower diamond distances than the σ_y rotations. The observation that σ_y rotations have a lower unitarity hints at the fact that the error discrepancy is due to incoherent noise. And indeed, the local dephasing probabilities are noticeably higher on the σ_y rotations.

The Pauli transfer matrix of the entangling gate is shown in Figure 3.4. The associated left noise channel reveals a pattern on its diagonal which is well explained by the tensor product of local dephasing channels. Looking at the dephasing probabilities of the entangling gate shown in Table 3.5, we see that they are more than a magnitude higher than the dephasing probabilities of single qubit gates. Thus, the comparatively low average gate fidelity of 0.685 is mostly due to dephasing noise. Additionally, there is an overrotation present, which can be seen by considering the anti-diagonal of the matrix $\hat{\mathcal{G}}\mathcal{U}^{-1}$ and noting that the pattern corresponds to the ideal $\sigma_z \otimes \sigma_z$ rotation. Through the computation of the Hamiltonian for the $r_K = 1$ fit, we also find that the angle of the $\sigma_z \otimes \sigma_z$ is 0.592π and thus higher than the target of $\pi/2$. In terms of other unwanted rotations, the most dominant is a local Z rotation on the first qubit with a rotation angle of -0.08π . The error pattern produced by the local Z rotation is also consistent with the pattern observed in the top right and bottom left of the plot for $\hat{\mathcal{G}}\mathcal{U}^{-1}$.

Since the number of free parameters for an $r_K = 16$ two qubit gate set is much larger than the number of sequences measured, the question arises whether our estimates give an accurate model of the true gates. Two arguments can be made in favor. First, we observe that no overfitting occurs, i.e. the cost function does not converge to a lower error than what is expected for the given shot noise. Second, we expect the true gates to be of low Kraus rank, especially the high fidelity single qubit gates. Therefore, the data should be consistent with a sparse, low parameter count model. In Table 3.6 we include the largest eigenvalues $\{\lambda_i\}$ of the Choi matrices

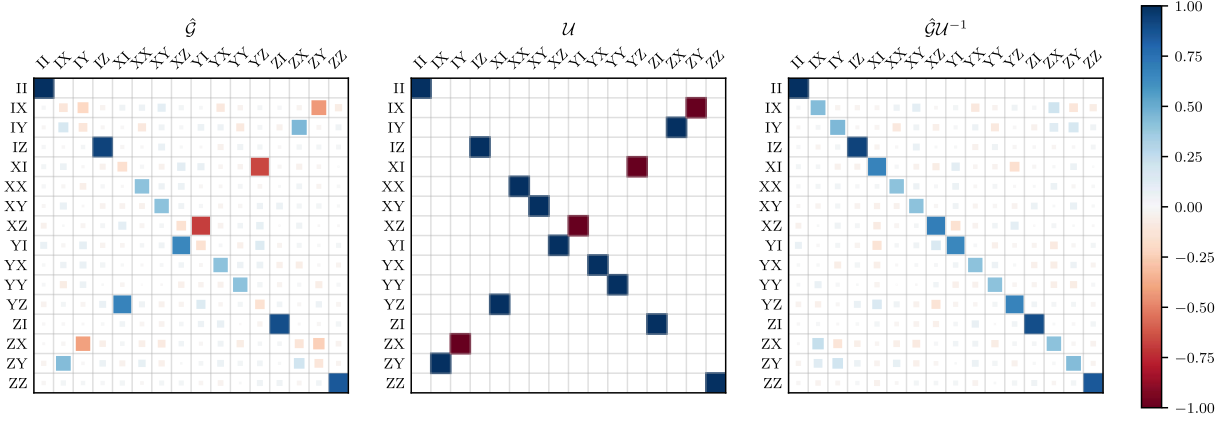


Figure 3.4: Pauli transfer matrix plots for the estimate $\hat{\mathcal{G}}$, whose target implementation is $U = e^{-i\frac{\pi}{4}\sigma_z \otimes \sigma_z}$. On the left is the full rank mGST reconstruction, in the center the ideal unitary gate and on the right the left noisy gate $\hat{\mathcal{G}}U^{-1}$. The area of each square tile represents the magnitude of the matrix entry, while the color coding represents both the magnitude and the sign.

of our estimates $\hat{\mathcal{G}}$, sorted from highest to lowest. The eigenvalues provide information about the approximate Kraus rank. More precisely, if we keep only the k largest Kraus operators $K_1, K_2, \dots, K_k$, the Frobenius norm error incurred by approximating the Choi matrix is given by $\|\mathcal{J}(\hat{\mathcal{G}})|_{r_K=k} - \mathcal{J}(\hat{\mathcal{G}})\|_F = \sum_{i>k} \lambda_i^2$.

	λ_1	λ_2	λ_3	λ_4	λ_5	λ_6	λ_7
Idle	0.95819	0.03885	0.00288	0.00006	0.00001	0.00000	0.00000
Rx(pi/2):0	0.99130	0.00679	0.00153	0.00035	0.00003	0.00000	0.00000
Ry(pi/2):0	0.97445	0.02083	0.00458	0.00010	0.00004	0.00000	0.00000
Rx(pi/2):1	0.99688	0.00250	0.00059	0.00003	0.00000	0.00000	0.00000
Ry(pi/2):1	0.96982	0.02335	0.00486	0.00161	0.00031	0.00004	0.00000
$e^{-i\frac{\pi}{4}}Z \otimes Z$	0.62867	0.21909	0.09275	0.03678	0.01879	0.00364	0.00024

Table 3.6: Eigenvalues of the Choi state for the two qubit full rank reconstruction.

In Table 3.6 we see that for single qubit gates the magnitude of the eigenvalues rapidly decays. For the entangling gate we find a slower decay, while nonetheless the contribution after the 4–5 largest eigenvalues becomes negligible. Thus far mGST uses the same predefined Kraus rank for all gates, but the above results suggest that choosing the rank on a per-gate basis or even adaptively could yield accurate estimates while further saving on computation time. More information about the mGST estimates and quality measures can again be found in the full report (Appendix C).

3.3 Improving shadow estimation with compressive GST

The classical shadow estimation protocol which we reviewed in Section 2.3 promises the *scalable* estimation of many observables at once and is thus very promising for current and near future quantum device consisting of hundreds of qubits. In addition to the theoretical error bounds under noise which we derive in Chapter 4, we develop a simple method in [45] to mitigate the effects of noise using compressive GST. Let $\phi(g)$ be the (noisy) implementation of a gate g and let $\omega(g)$ be the unitary target implementation. Then according to Eq. (2.158), the expectation value for the shadow estimate $\hat{o}$ of a given observable O on the state ρ is given by

$$\mathbb{E}[\hat{o}] = (O|S^{-1}\mathbb{E}[\phi(g)^\dagger|E_x]) = (O|S^{-1}\tilde{S}|\rho), \quad (3.12)$$

where S is the ideal frame operator and $\tilde{S}$ its noisy counterpart. In the protocol, $\tilde{S}$ directly depends on the experiment through $\phi(g)$, while S^{-1} is known beforehand and applied in post-processing. It is immediately clear that if we have information about the noise, we can construct an estimate $\hat{\tilde{S}}$ and use this operator instead for the post-processing step, resulting in a lower error incurred by the inversion: $\|\hat{\tilde{S}}^{-1}\tilde{S} - \text{id}\| < \|S^{-1}\tilde{S} - \text{id}\|$.

The specific shadow estimation protocol we consider here is the easiest to implement experimentally, where random Pauli bases measurements are used. We assume they are facilitated with the native computational basis measurement and the gate set $\mathfrak{G} = (\text{Id}, H, HS)$ where H is the Hadamard gate and S is the phase gate. The action of H enables a measurement in the eigenbasis of σ_x , while the action of HS enables a measurement in the eigenbasis of σ_y . The resulting noisy frame operator is defined as

$$\tilde{S} := \frac{1}{3^n} \sum_{g \in \mathfrak{G}^n} \sum_{x \in \{0,1\}^n} \omega(g) |E_x\rangle \langle E_x| \phi(g), \quad (3.13)$$

which for arbitrary global noise is a matrix in $\mathbb{C}^{d^2 \times d^2}$. In order for $\tilde{S}$ to be efficiently invertible, we assume that noisy implementations $\phi(g) = \omega(g)\Lambda(g)$ act only locally on two qubit pairs, which means we can factorize the frame operator as

$$\tilde{S} = \bigotimes_{i=1}^{n/2} \tilde{S}_{i,i+1}, \quad (3.14)$$

where we have w.l.o.g. taken n to be even. In Figure 3.5, circuits used in the shadow estimation protocol with the above noise assumption are illustrated.

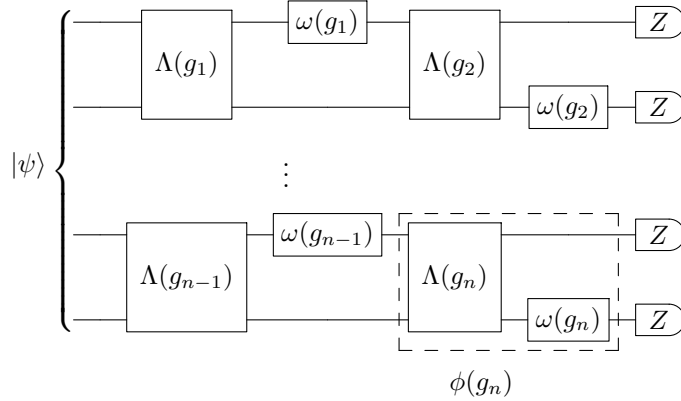


Figure 3.5: Circuit diagram for the noise model consisting of two qubit noise channels Λ_i acting on isolated two-qubit pairs.

Using compressive GST on these two-qubit pairs, the 2-local noise channels $\Lambda(g_i)$ can be estimated with $n/2$ additional GST experiments. In an experimental setting this would be done once, before any data for shadow estimation is taken.

In the following, an overview of the numerical simulations done in our publication [45] is given. As an observable we take the 10-qubit Heisenberg Hamiltonian

$$\mathcal{H} = \frac{1}{2} \sum_{j=1}^{10} (\sigma_x^j \sigma_x^{j+1} + \sigma_y^j \sigma_y^{j+1} + \sigma_z^j \sigma_z^{j+1} - \sigma_z^j) \quad (3.15)$$

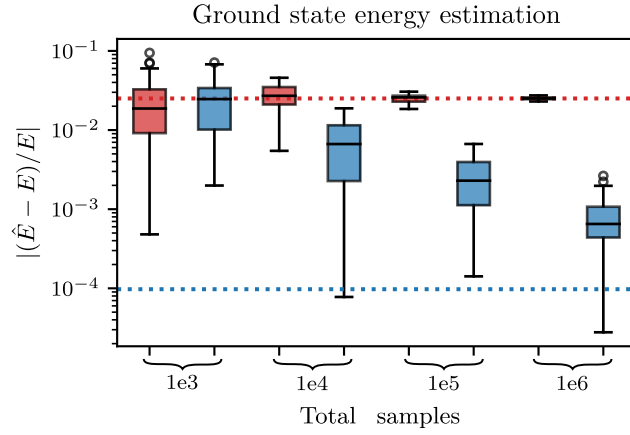


Figure 3.6: Relative error for the ground state energy estimate produced by standard shadow estimation (**red**) and error mitigated shadow estimation using mGST (**blue**). The boxes extend from the 25th to the 75th percentile of data, while the whiskers extend from the 5th to the 95th percentile. Along the x-axis the total number of samples for the shadow estimation protocol are varied. The dotted lines represent the infinite-sample limit.

with periodic boundary conditions. For the noise model we consider random two qubit unitaries given by $e^{i\gamma K}$, where K is drawn from the Gaussian unitary ensemble. The noise parameter γ is selected for to get average gate fidelities of $\mathcal{F}_{\text{avg}}(\omega(g_i), \phi(g_i)) = 0.99 \pm 0.001$. In Figure 3.6, standard shadow estimation results are compared to error mitigated results which use compressive GST estimates of Kraus rank $r_K = 2$. We see that for a low number of samples (10^{-3}), shot noise is still dominant and no reduction in relative error can be observed. Already for 10^4 samples, we observe considerable improvements in relative error using GST estimates. It is also important to note that for the given noisy model, the relative error for standard shadow estimation does not decrease when adding more samples, since the error is entirely dominated by the bias resulting from imperfect gates.

We further observe in Figure 3.6 that the relative error for infinitely many samples in the GST-assisted setting is limited to about 10^{-4} . This is due to the fact that the GST-estimates are finite sample estimates themselves, where we use 400 random sequences with 10^4 shots per sequence in the simulations. In the full article (Appendix A), the bias reduction is further studied not only for the ground state, but also for random states and all eigenstates of the Hamiltonian. We find a bias reduction between half an order of magnitude (for random states) and one order of magnitude (for Hamiltonian eigenstates) on average.

In summary, we presented a new error mitigation scheme which makes use of GST estimates and demonstrated in numerical simulations that significant error mitigation can be achieved in a practically relevant scenario. Implementing the mitigation protocol for real quantum hardware, for instance to estimate the energy of experimentally prepared states with respect to a quantum chemistry Hamiltonian, it thus a promising next step. The protocol can also be generalized to include different locality structures of noise channels, and in contrast to mitigation via robust shadows [44], the local noise channels are allowed to be fully general.

Chapter 4

Shadow estimation under gate-dependent noise

The ability to efficiently extract information about the final state of a quantum system is central for quantum computing as well as for gaining an understanding about the physical properties of the system. In Section 2.3 we reviewed a very promising recent development known as shadow estimation, which uses randomized measurements to efficiently obtain a classical model of the quantum state, from which many properties of the state can be estimated simultaneously. Even though the model is 'classical', in the precise sense that we can efficiently store it on a classical computer, we can use it to accurately predict uniquely quantum properties of the quantum state in the experiment, for instance the entanglement between subsystems. Due to its practical promise, shadow estimation has seen huge interest since the original works in 2019. Surprisingly though, there was very little theoretical understanding about the behavior of shadow estimation under noise. The original work [36] did not consider any type of noise, while later works [44, 199, 200] limit themselves to a simple gate-independent noise model. There are also noise mitigation methods which go beyond gate-independent or readout noise like our own work [45] or the work by Jnane *et al.* [201], which uses a probabilistic error cancellation method. Yet those methods require a model of the noise occurring in the experiment, which is either simplified, or costly to obtain.

This chapter gives an overview of our results on shadow estimation under general gate-dependent noise, where we follow the preprint [47], which is also included in Appendix B. We are mostly concerned with the estimation bias. Let $\hat{o}$ be the estimator of the expectation value of an observable O for the state ρ , then the bias is defined as $|\mathbb{E}[\hat{o}] - \text{Tr}(O\rho)|$. We first present two propositions, which show how a theoretical understanding of gate-dependent noise is fundamentally required to build trust in the shadow estimator.

Proposition 11. *Gate-dependent noise can lead to a bias of $\Omega(2^{n/4})$ on the shadow estimate of a pure state fidelity on an n -qubit state.*

An explicit example realizing this error scaling is given in [47]. We see that the bias resulting from noisy operations in the implementation of randomized measurements can amplify exponentially in the systems size in the worst case. What is also worrying is that 'robust shadows' [44], a method that provably mitigates noise for gate-independent errors, can itself introduce a bias which scales exponentially in the system size under gate-dependent noise.

Proposition 12. *'Robust classical shadows' can introduce a bias of up to $e^{\Omega(n)}$ while the standard classical shadows yield unbiased estimators.*

These propositions are based on explicit examples, and there is still the hope that they represent a worst case scaling, which can be avoided in everyday scenarios. To see if this is indeed the case, we derive general bounds on the bias.

Recall from Section 2.3 that a single shot shadow estimator is given by

$$\hat{o} := (O | S^{-1} \omega(g)^\dagger | E_x), \quad (4.1)$$

where $x \in \{0, 1\}^n$ is the computational basis measurement outcome, and $\omega(g)$ is a random basis change operation. We assume that g is sampled randomly from the Clifford group Cl_n according to a probability distribution p . This definition covers most random measurement protocols of classical shadows, such as uniform sampling from the global Clifford group Cl_n , uniform sampling from the local Clifford group $\text{Cl}_1^{\times n}$, shallow Clifford circuits, random Pauli basis measurements and matchgate 3-designs. Such a protocol is called informationally complete if the frame operator

$$S = \mathbb{E}_{g \sim p} [\omega(g)^\dagger \sum_x |E_x\rangle \langle E_x| \omega(g)], \quad (4.2)$$

is invertible (see Definition in Eq. (2.155)). The noise model we consider is time-stationary and context independent, and we write the noisy implementation of a gate g as $\phi(g) = \omega(g)\Lambda(g)$, where $\omega(g)$ is the ideal implementation and $\Lambda(g)$ is a gate-dependent CPT map. Note that any quantum channel $\phi(g)$ can be written in this form. To derive our main results we first give the following Lemma.

Lemma 13. *Consider an informationally complete Clifford-based shadow estimation protocol. Then, the shadow estimator's expected value is given by*

$$\mathbb{E}[\hat{o}] = \frac{1}{d} (O | \sum_{a \in \mathcal{P}_n} |\sigma_a\rangle \langle \sigma_a| \bar{\Lambda}_a | \rho), \quad (4.3)$$

where the quantum channels $\bar{\Lambda}_a$ are averages over noise channels $\Lambda(g)$ applied in the protocol.

We recall that $d^{-1} \sum_{a \in \mathcal{P}_n} |\sigma_a\rangle \langle \sigma_a|$ is just the identity channel. In effect, the protocol does not see the correct Pauli basis coefficients $(\sigma_a | \rho)$, but only the basis coefficients of perturbed states: $(\sigma_a | \bar{\Lambda}_a(\rho))$.

For the formulation of our main theorem, we also need what is called the *stabilizer norm*

$$\|O\|_{\text{st}} := \frac{1}{d} \sum_{a \in \mathcal{P}_n} |(\sigma_a | O)|, \quad (4.4)$$

which interestingly also turns up in the resource theory of magic, where low stabilizer norm observables and states lead to lower upper bounds on the computing time for a classical simulation of the system [202].

Theorem 14. *Consider an informationally complete Clifford-based shadow estimation protocol and let O be any observable. Then the estimation bias is bounded as*

$$|\mathbb{E}[\hat{o}] - \text{Tr}(O\rho)| \leq \begin{cases} \|O\|_{\text{st}} \cdot \epsilon_{\max} & \text{for arbitrary noise,} \\ \min\{\|O\|_2, \|O\|_{\text{st}}\} \cdot \epsilon_{\max} & \text{for Pauli noise,} \end{cases}$$

where $\epsilon_{\max} \leq 1$ is a worst-case error measure over all gates.

In the research paper [47] we show that $\epsilon_{\max}$ can either be taken as the maximum diamond norm error over unitaries in the protocol, or the maximum diamond norm error over the average noise channels $\bar{\Lambda}_a$. The utility of the bound depends on the stabilizer norm, which ranges from $\|O\|_{\text{st}} = 1$ for stabilizer states and Pauli observables to $\|O\|_{\text{st}} = \mathcal{O}(d)$ for magic states. Thankfully, low stabilizer norm observables encompass important use cases like fidelity estimation with respect to a stabilizer state, correlations functions and the energy estimation of a local Hamiltonian. When dealing with Pauli noise, the bound can be also formulated in terms of the Hilbert-Schmidt norm $\|O\|_2$, which satisfies $\|\psi\rangle\langle\psi\|_2 = 1$ for all $\psi \in \mathcal{H}$ (including magic

states). We know that the bound for arbitrary noise is essentially tight, since the example used to proof Proposition 11 is for a high stabilizer norm observable where $\|O\|_{\text{st}}$ scales as $2^{n/4}$.

Apart from control over the bias, the sample complexity of shadow estimation also has to remain controlled under noise in order for the protocol to be applicable in practice. The sample complexity is fully determined by the variance $\mathbb{V}[\hat{o}]$ of the estimator, and in the following we also present bounds on the variance for informationally complete shadow estimation protocols based on either the local or the global Clifford group.

Theorem 15. *Let O_0 be the traceless part of the observable O . Then, the variance of shadow estimation for uniform sampling from the global Clifford group is bounded by*

$$\mathbb{V}_{\text{global}}[\hat{o}] \leq \frac{2(d+1)}{(d+2)} \|O_0\|_{\text{st}}^2 + \frac{d+1}{d} \|O_0\|_2^2.$$

For uniform sampling from the local Clifford group, the variance is bounded by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 4^k \|O_{\text{loc}}\|_{\infty}^2$ for k -local observables $O = O_{\text{loc}} \otimes \mathbb{1}^{n-k}$ and by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 3^{\text{supp}(\sigma_a)}$ for Pauli observables $O = \sigma_a$.

The proof extensively uses the representation theory tools which we describe in Section 2.1.2. For comparison, the noise-free bound given in [36] is given by $\mathbb{V}_{\text{global}}[\hat{o}] \leq 3\|O\|_2^2$, while the noise-free bound for local Cliffords is exactly the same as in Theorem 15. Consequently, only the variance for the global Clifford group depends on the stabilizer norm, but in contrast to Theorem 14, it is not clear whether this bound is tight. In the paper we also give a variance bound on arbitrary informationally complete Clifford-based protocols in analogy to how the bound on the bias is formulated, but it depends on properties of the frame operator for the given distribution p and has to be evaluated separately for each protocol.

In Proposition 12 we already saw that there exist gate-dependent noise models for which the robust shadow estimation protocol can increase the bias dramatically. However, for certain reasonable gate-dependent noise models that were studied numerically in the literature [44, 203], robust shadows was shown to perform well. We thus wanted to see if a general statement can be made, i.e. if there is class of noise models (larger than gate-independent noise), for which the method still works. This led us to study analytically what happens in robust shadow estimation when the noise is gate-dependent. We found that under Pauli noise, robust shadows effectively corrects for the average Pauli-noise. We then show that if the Pauli-noise is sufficiently isotropic, robust shadow estimation still works with high probability, for any state and observable.

What we learn from the results summarized here is that there are scenarios in which shadow estimation and robust shadow estimation are not well-behaved, and care needs to be taken in an experiment to avoid them. Fortunately, as we show in Theorem 14 and Theorem 15, for most use cases the standard shadow estimation protocol is inherently stable against the effects of gate-dependent noise. Furthermore, our results on robust shadows under gate-dependent Pauli noise provide the basis for future error mitigation protocols that take more general noise models into account. We also expect our proof techniques to be useful for the study of other Clifford-based protocols in terms of their robustness against noise.

Chapter 5

Conclusion

To further reduce errors in current quantum devices and to learn from the quantum states they prepare, we need characterization methods that work even with imperfect device control. In this thesis we have concerned ourselves with two crucial tasks, the characterization of gate sets and the estimation of state properties via shadow estimation. We have identified limits in previous approaches, namely the inefficiency of GST and the lack of noise resilience guarantees for shadow estimation. To be equipped with the right tools to tackle these problems, we have given a theoretical background of several key ideas. First, we have familiarized ourselves with selected results from representation theory and from the study of matrix manifolds, as well as theoretical foundations of characterization tasks: Quantum operations, tensor networks, distance measures and the shadow estimation formalism. Additionally, we have taken a shot at summarizing the most influential ideas from the vast literature of quantum process tomography, of which compressed sensing and self-consistent methods are important for our own research.

In Chapter 3 we have presented our compressive GST algorithm and showed that it can accurately learn low rank approximations of gates in an experiment. Our algorithm uses fewer measurement settings and less classical computing time than previous methods, while only requiring random measurements and thus circumventing the need to design measurement settings on a per-case basis. This is all made possible with manifold optimization techniques as well as several problem-specific heuristic methods, showing that with the right tools, difficult optimization problems can be tamed in practice. In the application of our algorithm to a real world experiment, we have seen how low rank models are really well suited to capture realistic noise, even exceeding our expectations. We have for instance seen how an entangling gate with substantial noise contributions can still be accurately represented with just the 2-3 highest magnitude Kraus operators. These results validate our low rank model and have allowed us to provide helpful advice to experimentalists regarding error sources.

We have further numerically demonstrated how compressive GST can be used to mitigate errors in shadow estimation, leading to an increase in accuracy of an order of magnitude or higher for realistic noise models. This idea of mitigating errors uses full knowledge of local noise models obtained from GST and could be promising for error mitigation in other protocols as well.

To fill a gap in the understanding of shadow estimation under noise, we have presented theoretical guarantees on its bias and variance in Chapter 4. These guarantees show that the protocol is inherently noise-resilient, but only for the estimation of observables with a low stabilizer norm. For a high stabilizer norm, for instance in the task of estimating the fidelity to a magic state, we have shown that gate-dependent errors can accumulate and lead to an exponential bias. This suggests an interesting connection: observables with high magic, which is often regarded as a measure of 'quantumness' due to the difficulty to classically simulate systems with magic states or magic observables, are also the ones which are most susceptible to noise in the estimation protocol.

We have also analyzed previous error mitigation methods and shown that worryingly, they can themselves increase the bias in the presence of gate-dependent noise. To show when these

protocols can still be trusted, we have given more general conditions than in previous works for when this unwanted behavior does not occur.

There are several promising follow-up directions to our works. Regarding the analysis of error mitigation protocols under gate-dependent noise, we have only been able to show that they still work under unstructured Pauli noise. We believe that this can be generalized to larger classes of noise models using our proof techniques. Regarding compressive GST, we are currently working on further decreasing the classical computation time of our algorithm through the elimination of unnecessary gauge parameters and through avoiding the computation of the Hessian matrix when applicable. It would further be of high interest for compressive GST and the GST literature in general if theoretical guarantees for the learnability of gate sets with randomized measurements could be formulated.

Bibliography

- [1] J. Bardeen and W. H. Brattain, *The transistor, a semi-conductor triode*, Phys. Rev. **74**, 230 (1948).
- [2] R. P. Feynman, *Simulating physics with computers*, Int. J. Theor. Phys. **21**, 467 (1982).
- [3] P. Shor, *Algorithms for quantum computation: discrete logarithms and factoring*, Proceedings 35th Annual Symposium on Foundations of Computer Science 10.1109/sfcs.1994.365700 (1994).
- [4] L. K. Grover, *A fast quantum mechanical algorithm for database search*, in *Proceedings of the twenty-eighth annual ACM symposium on Theory of computing* (1996) pp. 212–219, arXiv:quant-ph/9605043 [quant-ph] .
- [5] C. H. Bennett and G. Brassard, *Quantum cryptography: Public key distribution and coin tossing*, Theoretical computer science **560**, 7 (2014).
- [6] S. Pirandola, U. L. Andersen, L. Banchi, M. Berta, D. Bunandar, R. Colbeck, D. Englund, T. Gehring, C. Lupo, C. Ottaviani, J. L. Pereira, M. Razavi, J. Shamsul Shaari, M. Tomamichel, V. C. Usenko, G. Vallone, P. Villoresi, and P. Wallden, *Advances in quantum cryptography*, Advances in Optics and Photonics **12**, 1012 (2020), arXiv:1906.01645 [quant-ph].
- [7] Y. Cao, J. Romero, J. P. Olson, M. Degroote, P. D. Johnson, M. Kieferová, I. D. Kivlichan, T. Menke, B. Peropadre, N. P. D. Sawaya, S. Sim, L. Veis, and A. Aspuru-Guzik, *Quantum chemistry in the age of quantum computing*, Chemical Reviews **119**, 10856 (2019), arXiv:1812.09976 [quant-ph].
- [8] M. Cerezo, A. Arrasmith, R. Babbush, S. C. Benjamin, S. Endo, K. Fujii, J. R. McClean, K. Mitarai, X. Yuan, L. Cincio, and P. J. Coles, *Variational quantum algorithms*, Nat. Rev. Phys. **3**, 625 (2021), arXiv:2012.09265 [quant-ph].
- [9] D. E. Deutsch, *Quantum computational networks*, Proceedings of the royal society of London. A. mathematical and physical sciences **425**, 73 (1989).
- [10] A. Y. Kitaev, *Quantum computations: algorithms and error correction*, Russian Math. Surv. **52**, 1191 (1997).
- [11] A. Barenco, D. Deutsch, A. Ekert, and R. Jozsa, *Conditional Quantum Dynamics and Logic Gates*, Phys. Rev. Lett. **74**, 4083 (1995), arXiv:quant-ph/9503017 [quant-ph].
- [12] J. I. Cirac and P. Zoller, *Quantum computations with cold trapped ions*, Phys. Rev. Lett. **74**, 4091 (1995).
- [13] C. Monroe, D. M. Meekhof, B. E. King, W. M. Itano, and D. J. Wineland, *Demonstration of a fundamental quantum logic gate*, Phys. Rev. Lett. **75**, 4714 (1995).
- [14] F. Schmidt-Kaler, H. Häffner, M. Riebe, S. Gulde, G. P. Lancaster, T. Deuschle, C. Becher, C. F. Roos, J. Eschner, and R. Blatt, *Realization of the cirac–zoller controlled-not quantum gate*, Nature **422**, 408 (2003).
- [15] Google AI Quantum and Collaborators, *Quantum supremacy using a programmable superconducting processor*, Nature **574**, 505 (2019), arXiv:1910.11333 [quant-ph].
- [16] Q. Zhu, S. Cao, F. Chen, M.-C. Chen, X. Chen, T.-H. Chung, H. Deng, Y. Du, D. Fan, M. Gong, C. Guo, C. Guo, S. Guo, L. Han, L. Hong, H.-L. Huang, Y.-H. Huo, L. Li, N. Li, S. Li, Y. Li, F. Liang, C. Lin, J. Lin, H. Qian, D. Qiao, H. Rong, H. Su, L. Sun, L. Wang,

- S. Wang, D. Wu, Y. Wu, Y. Xu, K. Yan, W. Yang, Y. Yang, Y. Ye, J. Yin, C. Ying, J. Yu, C. Zha, C. Zhang, H. Zhang, K. Zhang, Y. Zhang, H. Zhao, Y. Zhao, L. Zhou, C.-Y. Lu, C.-Z. Peng, X. Zhu, and J.-W. Pan, *Quantum computational advantage via 60-qubit 24-cycle random circuit sampling*, Science Bulletin **67**, 240 (2022), arXiv:2109.03494 [quant-ph].
- [17] L. S. Madsen, F. Laudenbach, M. F. Askarani, F. Rortais, T. Vincent, J. F. Bulmer, F. M. Miatto, L. Neuhaus, L. G. Helt, M. J. Collins, *et al.*, *Quantum computational advantage with a programmable photonic processor*, Nature **606**, 75 (2022).
- [18] Y. Kim, A. Eddins, S. Anand, K. X. Wei, E. Van Den Berg, S. Rosenblatt, H. Nayfeh, Y. Wu, M. Zaletel, K. Temme, *et al.*, *Evidence for the utility of quantum computing before fault tolerance*, Nature **618**, 500 (2023).
- [19] F. Pan, K. Chen, and P. Zhang, *Solving the Sampling Problem of the Sycamore Quantum Circuits*, Phys. Rev. Lett. **129**, 090502 (2022), arXiv:2111.03011 [quant-ph].
- [20] T. Begušić and G. Kin-Lic Chan, *Fast classical simulation of evidence for the utility of quantum computing before fault tolerance*, arXiv:2306.16372 [quant-ph] (2023).
- [21] E. T. Campbell, B. M. Terhal, and C. Vuillot, *Roads towards fault-tolerant universal quantum computation*, Nature **549**, 172 (2017), arXiv:1612.07330 [quant-ph].
- [22] D. Aharonov and M. Ben-Or, *Fault-tolerant quantum computation with constant error rate*, in *29th ACM Symp. on Theory of Computing (STOC)* (New York, 1997) pp. 176–188.
- [23] L. Postler, F. Butt, I. Pogorelov, C. D. Marciniak, S. Heußen, R. Blatt, P. Schindler, M. Risppler, M. Müller, and T. Monz, *Demonstration of fault-tolerant Steane quantum error correction*, arXiv:2312.09745 [quant-ph] (2023).
- [24] D. Bluvstein, S. J. Evered, A. A. Geim, S. H. Li, H. Zhou, T. Manovitz, S. Ebadi, M. Cain, M. Kalinowski, D. Hangleiter, J. P. Bonilla Ataides, N. Maskara, I. Cong, X. Gao, P. Sales Rodriguez, T. Karolyshyn, G. Semeghini, M. J. Gullans, M. Greiner, V. Vuletić, and M. D. Lukin, *Logical quantum processor based on reconfigurable atom arrays*, Nature **626**, 58 (2024), arXiv:2312.03982 [quant-ph].
- [25] B. Cheng, X.-H. Deng, X. Gu, Y. He, G. Hu, P. Huang, J. Li, B.-C. Lin, D. Lu, Y. Lu, C. Qiu, H. Wang, T. Xin, S. Yu, M.-H. Yung, J. Zeng, S. Zhang, Y. Zhong, X. Peng, F. Nori, and D. Yu, *Noisy intermediate-scale quantum computers*, Frontiers of Physics **18**, 21308 (2023), arXiv:2303.04061 [quant-ph].
- [26] J. Eisert, D. Hangleiter, N. Walk, I. Roth, D. Markham, R. Parekh, U. Chabaud, and E. Kashefi, *Quantum certification and benchmarking*, Nature Reviews Physics **2**, 382 (2020), arXiv:1910.06343 [quant-ph].
- [27] M. Kliesch and I. Roth, *Theory of quantum system certification*, PRX Quantum **2**, 010201 (2021), arXiv:2010.05925 [quant-ph].
- [28] S. T. Merkel, J. M. Gambetta, J. A. Smolin, S. Poletto, A. D. Córcoles, B. R. Johnson, C. A. Ryan, and M. Steffen, *Self-consistent quantum process tomography*, Phys. Rev. A **87**, 062119 (2013), arXiv:1211.0322 [quant-ph].
- [29] C. Stark, *Self-consistent tomography of the state-measurement Gram matrix*, Phys. Rev. A **89**, 052109 (2014), arXiv:1209.5737 [quant-ph].
- [30] R. Blume-Kohout, J. King Gamble, E. Nielsen, J. Mizrahi, J. D. Sterk, and P. Maunz, *Robust, self-consistent, closed-form tomography of quantum logic gates on a trapped ion qubit*, arXiv:1310.4492 [quant-ph].

- [31] E. Knill, D. Leibfried, R. Reichle, J. Britton, R. B. Blakestad, J. D. Jost, C. Langer, R. Ozeri, S. Seidelin, and D. J. Wineland, *Randomized benchmarking of quantum gates*, Phys. Rev. A **77**, 012307 (2008), arXiv:0707.0963 [quant-ph].
- [32] J. Helsen, I. Roth, E. Onorati, A. Werner, and J. Eisert, *General framework for randomized benchmarking*, PRX Quantum **3**, 020357 (2022), arXiv:2010.07974 [quant-ph].
- [33] A. W. Cross, L. S. Bishop, S. Sheldon, P. D. Nation, and J. M. Gambetta, *Validating quantum computers using randomized model circuits*, Phys. Rev. A **100**, 032328 (2019), arXiv:1811.12926 [quant-ph].
- [34] R. Blume-Kohout and K. C. Young, *A volumetric framework for quantum computer benchmarks*, Quantum **4**, 362 (2020), arXiv:1904.05546 [quant-ph].
- [35] T. Proctor, S. Seritan, K. Rudinger, E. Nielsen, R. Blume-Kohout, and K. Young, *Scalable Randomized Benchmarking of Quantum Computers Using Mirror Circuits*, Phys. Rev. Lett. **129**, 150502 (2022), arXiv:2112.09853 [quant-ph].
- [36] H.-Y. Huang, R. Kueng, and J. Preskill, *Predicting many properties of a quantum system from very few measurements*, Nature Physics **16**, 1050–1057 (2020), arXiv:2002.08953 [quant-ph].
- [37] A. Elben, S. T. Flammia, H.-Y. Huang, R. Kueng, J. Preskill, B. Vermersch, and P. Zoller, *The randomized measurement toolbox*, Nat. Rev. Phys. 10.1038/s42254-022-00535-2 (2022), arXiv:2203.11374.
- [38] A. Elben, R. Kueng, H.-Y. R. Huang, R. van Bijnen, C. Kokail, M. Dalmonte, P. Calabrese, B. Kraus, J. Preskill, P. Zoller, and B. Vermersch, *Mixed-state entanglement from local randomized measurements*, Phys. Rev. Lett. **125**, 200501 (2020), arXiv:2007.06305 [quant-ph].
- [39] T. Zhang, J. Sun, X.-X. Fang, X.-M. Zhang, X. Yuan, and H. Lu, *Experimental quantum state measurement with classical shadows*, Phys. Rev. Lett. **127**, 200501 (2021), arXiv:2008.05234 [quant-ph].
- [40] G. Struchalin, Y. A. Zagorovskii, E. Kovlakov, S. Straupe, and S. Kulik, *Experimental estimation of quantum state properties from classical shadows*, PRX Quantum **2**, 010307 (2021), arXiv:2008.05234 [quant-ph].
- [41] H.-Y. Huang, R. Kueng, G. Torlai, V. V. Albert, and J. Preskill, *Provably efficient machine learning for quantum many-body problems*, Science **377**, eabk3333 (2022), arXiv:2106.12627 [quant-ph].
- [42] W. J. Huggins, B. A. O’Gorman, N. C. Rubin, D. R. Reichman, R. Babbush, and J. Lee, *Unbiasing fermionic quantum Monte Carlo with a quantum computer*, Nature **603**, 416 (2022), 2106.16235 [quant-ph].
- [43] J. Helsen, M. Ioannou, J. Kitzinger, E. Onorati, A. H. Werner, J. Eisert, and I. Roth, *Shadow estimation of gate-set properties from random sequences*, Nature Communications **14**, 5039 (2023), arXiv:2110.13178 [quant-ph].
- [44] S. Chen, W. Yu, P. Zeng, and S. T. Flammia, *Robust shadow estimation*, PRX Quantum **2**, 030348 (2021), arXiv:2011.09636 [quant-ph].
- [45] R. Brieger, I. Roth, and M. Kliesch, *Compressive gate set tomography*, PRX Quantum **4**, 010325 (2023), arXiv:2112.05176 [quant-ph].

- [46] R. Brieger, I. Roth, and M. Kliesch, *Python implementation of mGST, a compressive gate set tomography algorithm*, <https://github.com/rabrie/mGST> (2021).
- [47] R. Brieger, M. Heinrich, I. Roth, and M. Kliesch, *Stability of classical shadows under gate-dependent noise*, arXiv:2310.19947 [quant-ph] (2023).
- [48] I. A. Luchnikov, M. E. Krechetov, and S. N. Filippov, *Riemannian geometry and automatic differentiation for optimization problems of quantum physics and quantum technologies*, New Journal of Physics **23**, 073006 (2021), arXiv:2007.01287 [quant-ph].
- [49] W. Fulton and J. Harris, *Representation theory*, Vol. 129 (Springer Science & Business Media, 2013).
- [50] B. Simon, *Representations of finite and compact groups*, 10 (Am. Math. Soc., 1996).
- [51] R. Goodman and N. R. Wallach, *Representations and invariants of the classical groups* (Cambridge University Press, 2000).
- [52] J. Emerson, R. Alicki, and K. Życzkowski, *Scalable noise estimation with random unitary operators*, J. Opt. B **7**, S347 (2005), arXiv:quant-ph/0503243.
- [53] E. Magesan, J. M. Gambetta, and J. Emerson, *Characterizing quantum gates via randomized benchmarking*, Phys. Rev. A **85**, 042311 (2012), arXiv:1109.6887.
- [54] M. Heinrich, M. Kliesch, and I. Roth, *General guarantees for randomized benchmarking with random quantum circuits*, arXiv:2212.06181 [quant-ph] (2022).
- [55] D. A. Roberts and B. Yoshida, *Chaos and complexity by design*, Journal of High Energy Physics 10.1007/JHEP04(2017)121, arXiv:1610.04903 [quant-ph].
- [56] Z. Webb, *The Clifford group forms a unitary 3-design*, arXiv:1510.02769 [quant-ph].
- [57] H. Zhu, R. Kueng, M. Grassl, and D. Gross, *The Clifford group fails gracefully to be a unitary 4-design*, arXiv:1609.08172 [quant-ph].
- [58] A. Edelman, T. A. Arias, and S. T. Smith, *The geometry of algorithms with orthogonality constraints*, SIAM Journal on Matrix Analysis and Applications **20**, 303 (1998), arXiv:physics/9806030 [physics.comp-ph].
- [59] P. Absil, R. Mahony, and R. Sepulchre, *Optimization algorithms on matrix manifolds* (Princeton University Press, 2009).
- [60] J. H. Manton, *Optimization algorithms exploiting unitary constraints*, IEEE Trans. Signal Process. **50**, 635 (2002).
- [61] G. Bécigneul and O.-E. Ganea, *Riemannian Adaptive Optimization Methods*, arXiv e-prints (2018), arXiv:1810.00760 [cs.LG].
- [62] J. Li, L. Fuxin, and S. Todorovic, *Efficient Riemannian Optimization on the Stiefel Manifold via the Cayley Transform*, arXiv e-prints (2020), arXiv:2002.01113 [cs.LG].
- [63] N. Boumal, P. A. Absil, and C. Cartis, *Global rates of convergence for nonconvex optimization on manifolds*, arXiv e-prints (2016), arXiv:1605.08101 [math.OC].
- [64] N. Boumal, *An introduction to optimization on smooth manifolds* (Cambridge University Press, 2023).
- [65] S. Wisdom, T. Powers, J. R. Hershey, J. Le Roux, and L. Atlas, *Full-capacity unitary recurrent neural networks*, in *Adv. Neural Inf. Process. Syst.*, Vol. 29 (2016) arXiv:1611.00035 [stat.ML] .

- [66] F. Lotte, L. Bougrain, A. Cichocki, M. Clerc, M. Congedo, A. Rakotomamonjy, and F. Yger, *A review of classification algorithms for EEG-based brain-computer interfaces: a 10 year update*, Journal of Neural Engineering **15**, 031005 (2018).
- [67] E. Chiumiento and M. Melgaard, *Stiefel and grassmann manifolds in quantum chemistry*, Journal of Geometry and Physics **62**, 1866 (2012).
- [68] X. Yu, J.-C. Shen, J. Zhang, and K. B. Letaief, *Alternating minimization algorithms for hybrid precoding in millimeter wave mimo systems*, IEEE Journal of Selected Topics in Signal Processing **10**, 485 (2016).
- [69] F. Liu, C. Masouros, A. Li, H. Sun, and L. Hanzo, *Mu-mimo communications with mimo radar: From co-existence to joint transmission*, IEEE Transactions on Wireless Communications **17**, 2755 (2018).
- [70] E. Knill, *Quantum computing with realistically noisy devices*, Nature **434**, 39 (2005), arXiv:quant-ph/0410199.
- [71] R. Kueng, D. M. Long, A. C. Doherty, and S. T. Flammia, *Comparing experiments to the fault-tolerance threshold*, Phys. Rev. Lett. **117**, 170502 (2016), arXiv:1510.05653 [quant-ph].
- [72] K. Temme, S. Bravyi, and J. M. Gambetta, *Error Mitigation for Short-Depth Quantum Circuits*, Phys. Rev. Lett. **119**, 180509 (2017), arXiv:1612.02058 [quant-ph].
- [73] M. A. Nielsen and I. L. Chuang, *Quantum computation and quantum information* (Cambridge University Press, 2010).
- [74] J. Watrous, *The Theory of Quantum Information* (Cambridge University Press, 2018).
- [75] C. J. Wood, J. D. Biamonte, and D. G. Cory, *Tensor networks and graphical calculus for open quantum systems*, Quant. Inf. Comp. **15**, 0579 (2015), arXiv:1111.6950 [quant-ph].
- [76] K. Życzkowski and I. Bengtsson, *On Duality between Quantum Maps and Quantum States*, Open Systems & Information Dynamics **11**, 3 (2004).
- [77] S. T. Flammia and J. J. Wallman, *Efficient estimation of Pauli channels*, ACM Transactions on Quantum Computing **1**, 1 (2020), arXiv:1907.12976 [quant-ph].
- [78] S. J. Beale, J. J. Wallman, M. Gutiérrez, K. R. Brown, and R. Laflamme, *Quantum error correction decoheres noise*, Phys. Rev. Lett. **121**, 190501 (2018).
- [79] J. J. Wallman and J. Emerson, *Noise tailoring for scalable quantum computation via randomized compiling*, Phys. Rev. A **94**, 052325 (2016), arXiv:1512.01098 [quant-ph].
- [80] A. Hashim, R. K. Naik, A. Morvan, J.-L. Ville, B. Mitchell, J. M. Kreikebaum, M. Davis, E. Smith, C. Iancu, K. P. O'Brien, I. Hincks, J. J. Wallman, J. Emerson, and I. Siddiqi, *Randomized Compiling for Scalable Quantum Computing on a Noisy Superconducting Quantum Processor*, Physical Review X **11**, 041039 (2021), arXiv:2010.00215 [quant-ph].
- [81] M. Ware, G. Ribeill, D. Ristè, C. A. Ryan, B. Johnson, and M. P. da Silva, *Experimental Pauli-frame randomization on a superconducting qubit*, Phys. Rev. A **103**, 042604 (2021), arXiv:1803.01818 [quant-ph].
- [82] G. García-Pérez, M. A. C. Rossi, and S. Maniscalco, *Ibm q experience as a versatile experimental testbed for simulating open quantum systems*, npj Quantum Information **6**, 1 (2020), arXiv:1906.07099 [quant-ph].

- [83] M. M. Wolf and J. I. Cirac, *Dividing quantum channels*, Commun. Math. Phys. **279**, 147 (2008), math-ph/0611057.
- [84] M. Guță, J. Kahn, R. Kueng, and J. A. Tropp, *Fast state tomography with optimal error bounds*, Journal of Physics A Mathematical General **53**, 10.1088/1751-8121/ab8111 (2020), arXiv:1809.11162 [quant-ph].
- [85] M. Cramer, M. B. Plenio, S. T. Flammia, R. Somma, D. Gross, S. D. Bartlett, O. Landon-Cardinal, D. Poulin, and Y.-K. Liu, *Efficient quantum state tomography*, Nat. Commun. **1**, 149 (2010), arXiv:1101.4366 [quant-ph].
- [86] A. Kyrillidis, A. Kalev, D. Park, S. Bhojanapalli, C. Caramanis, and S. Sanghavi, *Provable compressed sensing quantum state tomography via non-convex methods*, npj Quant. Inf. **4**, 36 (2018), arXiv:1711.02524 [quant-ph].
- [87] F. G. S. L. Brandão, R. Kueng, and D. Stilck França, *Fast and robust quantum state tomography from few basis measurements*, arXiv e-prints 10.48550/arXiv.2009.08216 (2020), arXiv:2009.08216 [quant-ph].
- [88] J. A. Smolin, J. M. Gambetta, and G. Smith, *Efficient Method for Computing the Maximum-Likelihood Quantum State from Measurements with Additive Gaussian Noise*, Phys. Rev. Lett. **108**, 10.1103/PhysRevLett.108.070502 (2012), arXiv:1106.5458 [quant-ph].
- [89] M. A. Nielsen, *A simple formula for the average gate fidelity of a quantum dynamical operation*, Phys. Lett. A **303**, 249 (2002), arXiv:quant-ph/0205035.
- [90] T. Proctor, K. Rudinger, K. Young, M. Sarovar, and R. Blume-Kohout, *What randomized benchmarking actually measures*, Phys. Rev. Lett. **119**, 130502 (2017), arXiv:1702.01853 [quant-ph].
- [91] A. W. Cross, D. P. DiVincenzo, and B. M. Terhal, *A comparative code study for quantum fault-tolerance*, Quant. Inf. Comp. **9**, 0541 (2009), arXiv:0711.1556 [quant-ph].
- [92] J. Watrous, *Simpler semidefinite programs for completely bounded norms*, Chicago J. Theo. Comp. Sci. **2013**, 1 (2013), arXiv:1207.5726.
- [93] J. J. Wallman, *Bounding experimental quantum error rates relative to fault-tolerant thresholds*, arXiv:1511.00727 [quant-ph] (2015).
- [94] J. Wallman, C. Granade, R. Harper, and S. T. Flammia, *Estimating the coherence of noise*, New J. Phys. **17**, 10.1088/1367-2630/17/11/113020 (2015), arXiv:1503.07865 [quant-ph].
- [95] E. Magesan, *Characterizing noise in quantum systems* (2012), PhD thesis.
- [96] F. B. Maciejewski, Z. Zimborás, and M. Oszmaniec, *Mitigation of readout noise in near-term quantum devices by classical post-processing based on detector tomography*, Quantum **4**, 257 (2020), arXiv:1907.08518 [quant-ph].
- [97] Z. Puchała, L. Paweł, A. Krawiec, and R. Kukulski, *Strategies for optimal single-shot discrimination of quantum measurements*, Phys. Rev. A **98**, 042103 (2018), arXiv:1804.05856 [quant-ph].
- [98] I. L. Chuang and M. A. Nielsen, *Prescription for experimental determination of the dynamics of a quantum black box*, Journal of Modern Optics **44**, 2455 (1997), arXiv:quant-ph/9610001 [quant-ph].

- [99] J. F. Poyatos, J. I. Cirac, and P. Zoller, *Complete Characterization of a Quantum Process: The Two-Bit Quantum Gate*, Phys. Rev. Lett. **78**, 390 (1997), arXiv:quant-ph/9611013 [quant-ph].
- [100] Y. S. Weinstein, T. F. Havel, J. Emerson, N. Boulant, M. Saraceno, S. Lloyd, and D. G. Cory, *Quantum process tomography of the quantum Fourier transform*, J. Phys. Chem. **121**, 6117 (2004), arXiv:quant-ph/0406239 [quant-ph].
- [101] J. L. O'Brien, G. J. Pryde, A. Gilchrist, D. F. James, N. K. Langford, T. C. Ralph, and A. G. White, *Quantum Process Tomography of a Controlled-NOT Gate*, Phys. Rev. Lett. **93**, 080502 (2004), arXiv:quant-ph/0402166 [quant-ph].
- [102] M. Howard, J. Twamley, C. Wittmann, T. Gaebel, F. Jelezko, and J. Wrachtrup, *Quantum process tomography and Linblad estimation of a solid-state qubit*, New Journal of Physics **8**, 33 (2006), arXiv:quant-ph/0601167 [quant-ph].
- [103] M. Riebe, K. Kim, P. Schindler, T. Monz, P. O. Schmidt, T. K. Körber, W. Hänsel, H. Häffner, C. F. Roos, and R. Blatt, *Process Tomography of Ion Trap Quantum Gates*, Phys. Rev. Lett. **97**, 220407 (2006), arXiv:quant-ph/0609228 [quant-ph].
- [104] R. C. Bialczak, M. Ansmann, M. Hofheinz, E. Lucero, M. Neeley, A. D. O'Connell, D. Sank, H. Wang, J. Wenner, M. Steffen, A. N. Cleland, and J. M. Martinis, *Quantum process tomography of a universal entangling gate implemented with Josephson phase qubits*, Nature Physics **6**, 409 (2010), arXiv:0910.1118 [quant-ph].
- [105] A. Klappenecker and M. Roetteler, *Mutually Unbiased Bases are Complex Projective 2-Designs*, arXiv e-prints, quant-ph/0502031 (2005), arXiv:quant-ph/0502031 [quant-ph].
- [106] J. M. Renes, R. Blume-Kohout, A. J. Scott, and C. M. Caves, *Symmetric informationally complete quantum measurements*, Journal of Mathematical Physics **45**, 2171 (2004), arXiv:quant-ph/0310075 [quant-ph].
- [107] A. Bendersky, F. Pastawski, and J. P. Paz, *Selective Efficient Quantum Process Tomography*, arXiv:0801.0758 [quant-ph] (2008).
- [108] C. T. Schmiegelow, A. Bendersky, M. A. Larotonda, and J. P. Paz, *Selective and efficient quantum process tomography without ancilla*, Phys. Rev. Lett. **107**, 100502 (2011), arXiv:1105.4815 [quant-ph].
- [109] M. Mohseni and D. A. Lidar, *Direct Characterization of Quantum Dynamics*, Phys. Rev. Lett. **97**, 10.1103/PhysRevLett.97.170501 (2006), arXiv:quant-ph/0601033 [quant-ph].
- [110] M. Mohseni and D. A. Lidar, *Direct characterization of quantum dynamics: General theory*, Phys. Rev. A **75**, 10.1103/PhysRevA.75.062331 (2007), arXiv:quant-ph/0601034 [quant-ph].
- [111] G. M. D'Ariano and P. Lo Presti, *Quantum Tomography for Measuring Experimentally the Matrix Elements of an Arbitrary Quantum Operation*, Phys. Rev. Lett. **86**, 4195 (2001), arXiv:quant-ph/0012071 [quant-ph].
- [112] D. W. Leung, *Choi's proof as a recipe for quantum process tomography*, Journal of Mathematical Physics **44**, 528 (2003), arXiv:quant-ph/0201119 [quant-ph].
- [113] J. B. Altepeter, D. Branning, E. Jeffrey, T. C. Wei, P. G. Kwiat, R. T. Thew, J. L. O'Brien, M. A. Nielsen, and A. G. White, *Ancilla-Assisted Quantum Process Tomography*, Phys. Rev. Lett. **90**, 193601 (2003), arXiv:quant-ph/0303038 [quant-ph].

- [114] C. Granade, J. Combes, and D. G. Cory, *Practical Bayesian tomography*, New Journal of Physics **18**, 033024 (2016), arXiv:1509.03770 [quant-ph].
- [115] C. Granade, C. Ferrie, and S. T. Flammia, *Practical adaptive quantum tomography*, New Journal of Physics **19**, 113017 (2017), arXiv:1605.05039 [quant-ph].
- [116] T. Surawy-Stepney, J. Kahn, R. Kueng, and M. Guta, *Projected Least-Squares Quantum Process Tomography*, Quantum **6**, 844 (2022), arXiv:2107.01060 [quant-ph].
- [117] G. C. Knee, E. Bolduc, J. Leach, and E. M. Gauger, *Quantum process tomography via completely positive and trace-preserving projection*, Phys. Rev. A **98**, 062336 (2018), arXiv:1803.10062 [quant-ph].
- [118] A. Oufkir, *Sample-Optimal Quantum Process Tomography with Non-Adaptive Incoherent Measurements*, arXiv e-prints, arXiv:2301.12925 (2023), arXiv:2301.12925 [quant-ph].
- [119] S. Foucart and H. Rauhut, *A mathematical introduction to compressive sensing* (Springer, 2013).
- [120] R. L. Kosut, *Quantum Process Tomography via L1-norm Minimization*, arXiv:0812.4323 [quant-ph] (2008).
- [121] A. Shabani, R. L. Kosut, M. Mohseni, H. Rabitz, M. A. Broome, M. P. Almeida, A. Fedrizzi, and A. G. White, *Efficient measurement of quantum dynamics via compressive sensing*, Phys. Rev. Lett. **106**, 100401 (2011), arXiv:0910.5498 [quant-ph].
- [122] D. Gross, *Recovering low-rank matrices from few coefficients in any basis*, IEEE Trans. Inf. Th. **57**, 1548 (2011), arXiv:0910.1879 [cs.IT].
- [123] D. Gross, Y.-K. Liu, S. T. Flammia, S. Becker, and J. Eisert, *Quantum state tomography via compressed sensing*, Phys. Rev. Lett. **105**, 150401 (2010), arXiv:0909.3304 [quant-ph].
- [124] S. T. Flammia, D. Gross, Y.-K. Liu, and J. Eisert, *Quantum tomography via compressed sensing: error bounds, sample complexity and efficient estimators*, New J. Phys. **14**, 095022 (2012), arXiv:1205.2300 [quant-ph].
- [125] M. Kliesch, R. Kueng, J. Eisert, and D. Gross, *Guaranteed recovery of quantum processes from few measurements*, Quantum **3**, 171 (2019), arXiv:1701.03135 [quant-ph].
- [126] M. Kliesch, R. Kueng, J. Eisert, and D. Gross, *Improving compressed sensing with the diamond norm*, IEEE Trans. Inf. Th. **62**, 7445 (2016), arXiv:1511.01513 [cs.IT].
- [127] C. H. Baldwin, A. Kalev, and I. H. Deutsch, *Quantum process tomography of unitary and near-unitary maps*, Phys. Rev. A **90**, 012110 (2014), arXiv:1404.2877 [quant-ph].
- [128] S. Kimmel and Y. K. Liu, *Phase retrieval using unitary 2-designs*, in *2017 International Conference on Sampling Theory and Applications (SampTA)* (2017) pp. 345–349, arXiv:1510.08887 [quant-ph].
- [129] I. Roth, R. Kueng, S. Kimmel, Y. K. Liu, D. Gross, J. Eisert, and M. Kliesch, *Recovering quantum gates from few average gate fidelities*, Phys. Rev. Lett. **121**, 170502 (2018), arXiv:1803.00572 [quant-ph].
- [130] S. Kimmel, M. P. da Silva, C. A. Ryan, B. R. Johnson, and T. Ohki, *Robust extraction of tomographic information via randomized benchmarking*, Phys. Rev. X **4**, 011050 (2014), arXiv:1306.2348 [quant-ph].
- [131] C. B. Mendl and M. M. Wolf, *Unital quantum channels - convex structure and revivals of Birkhoff's theorem*, Comm. Math. Phys. **289**, 1057 (2009), arXiv:0806.2820 [quant-ph].

- [132] A. J. Scott, *Optimizing quantum process tomography with unitary 2-designs*, Journal of Physics A Mathematical General **41**, 055308 (2008), arXiv:0711.1017 [quant-ph].
- [133] I. Roth, Ph.D. thesis, Freie Universität Berlin.
- [134] J. P. Gaebler, A. M. Meier, T. R. Tan, R. Bowler, Y. Lin, D. Hanneke, J. D. Jost, J. P. Home, E. Knill, D. Leibfried, and D. J. Wineland, *Randomized benchmarking of multiqubit gates*, Phys. Rev. Lett. **108**, 260503 (2012), arXiv:1203.3733 [quant-ph].
- [135] E. Magesan, J. M. Gambetta, B. R. Johnson, C. A. Ryan, J. M. Chow, S. T. Merkel, M. P. da Silva, G. A. Keefe, M. B. Rothwell, T. A. Ohki, M. B. Ketchen, and M. Steffen, *Efficient measurement of quantum gate error by interleaved randomized benchmarking*, Phys. Rev. Lett. **109**, 080505 (2012), arXiv:1203.4550 [quant-ph].
- [136] A. Carignan-Dugas, J. J. Wallman, and J. Emerson, *Bounding the average gate fidelity of composite channels using the unitarity*, New J. Phys. **21**, 053016 (2019), arXiv:1610.05296 [quant-ph].
- [137] J. Helsen, X. Xue, L. M. K. Vandersypen, and S. Wehner, *A new class of efficient randomized benchmarking protocols*, npj Quant. Inf. **5**, 71 (2019), arXiv:1806.02048 [quant-ph].
- [138] R. Harper, S. T. Flammia, and J. J. Wallman, *Efficient learning of quantum noise*, Nat. Phys. 10.1038/s41567-020-0992-8 (2020), arXiv:1907.13022 [quant-ph].
- [139] C. Rouzé and D. Stilck França, *Efficient learning of the structure and parameters of local Pauli noise channels*, arXiv:2307.02959 [quant-ph] (2023).
- [140] J. Helsen, F. Battistel, and B. M. Terhal, *Spectral quantum tomography*, npj Quantum Information **5**, 74 (2019), arXiv:1904.00177 [quant-ph].
- [141] T. Sarkar and O. Pereira, *Using the matrix pencil method to estimate the parameters of a sum of complex exponentials*, IEEE Antennas and Propagation Magazine **37**, 48 (1995).
- [142] H.-Y. Huang, S. Chen, and J. Preskill, *Learning to Predict Arbitrary Quantum Processes*, PRX Quantum **4**, 10.1103/PRXQuantum.4.040337 (2023), arXiv:2210.14894 [quant-ph].
- [143] J. Kunjummen, M. C. Tran, D. Carney, and J. M. Taylor, *Shadow process tomography of quantum channels*, Phys. Rev. A **107**, 10.1103/PhysRevA.107.042403 (2023), arXiv:2110.03629 [quant-ph].
- [144] R. Levy, D. Luo, and B. K. Clark, *Classical shadows for quantum process tomography on near-term quantum computers*, Physical Review Research **6**, 10.1103/PhysRevResearch.6.013029 (2024), arXiv:2110.02965 [quant-ph].
- [145] E. Onorati, T. Kohler, and T. S. Cubitt, *Fitting time-dependent Markovian dynamics to noisy quantum channels*, arXiv:2303.08936 [quant-ph] (2023).
- [146] D. Hangleiter, I. Roth, J. Fuksa, J. Eisert, and P. Roushan, *Robustly learning the Hamiltonian dynamics of a superconducting quantum processor*, arXiv e-prints, arXiv:2108.08319 (2021), arXiv:2108.08319 [quant-ph].
- [147] W. Yu, J. Sun, Z. Han, and X. Yuan, *Robust and Efficient Hamiltonian Learning*, Quantum **7**, 1045 (2023), arXiv:2201.00190 [quant-ph].
- [148] D. Stilck França, L. A. Markovich, V. V. Dobrovitski, A. H. Werner, and J. Borregaard, *Efficient and robust estimation of many-qubit Hamiltonians*, Nature Communications **15**, 311 (2024), arXiv:2205.09567 [quant-ph].

- [149] A. Gu, L. Cincio, and P. J. Coles, *Practical Hamiltonian learning with unitary dynamics and Gibbs states*, arXiv:2206.15464 [quant-ph] (2024).
- [150] F. Wilde, A. Kshetrimayum, I. Roth, D. Hangleiter, R. Sweke, and J. Eisert, *Scalably learning quantum many-body Hamiltonians from dynamical data*, arXiv e-prints (2022), arXiv:2209.14328 [quant-ph].
- [151] H.-Y. Huang, Y. Tong, D. Fang, and Y. Su, *Learning Many-Body Hamiltonians with Heisenberg-Limited Scaling*, Phys. Rev. Lett. **130**, 200403 (2023), arXiv:2210.03030 [quant-ph].
- [152] H. Li, Y. Tong, H. Ni, T. Gefen, and L. Ying, *Heisenberg-limited Hamiltonian learning for interacting bosons*, arXiv e-prints , arXiv:2307.04690 (2023), arXiv:2307.04690 [quant-ph].
- [153] S. Kimmel, G. H. Low, and T. J. Yoder, *Robust calibration of a universal single-qubit gate set via robust phase estimation*, Phys. Rev. A **92**, 062315 (2015), arXiv:1502.02677 [quant-ph].
- [154] F. Verstraete, V. Murg, and J. Cirac, *Matrix product states, projected entangled pair states, and variational renormalization group methods for quantum spin systems*, Adv. Phys. **57**, 143 (2008), arXiv:0907.2796 [quant-ph].
- [155] J. I. Cirac, D. Pérez-García, N. Schuch, and F. Verstraete, *Matrix product states and projected entangled pair states: Concepts, symmetries, theorems*, Rev. Mod. Phys. **93**, 045003 (2021), arXiv:2011.12127 [quant-ph].
- [156] T. Baumgratz, D. Gross, M. Cramer, and M. B. Plenio, *Scalable reconstruction of density matrices*, Phys. Rev. Lett. **111**, 020401 (2013).
- [157] B. P. Lanyon, C. Maier, M. Holzäpfel, T. Baumgratz, C. Hempel, P. Jurcevic, I. Dhand, A. S. Buyskikh, A. J. Daley, M. Cramer, M. B. Plenio, R. Blatt, and C. F. Roos, *Efficient tomography of a quantum many-body system*, Nature Physics **13**, 1158 (2017).
- [158] M. Holzäpfel, T. Baumgratz, M. Cramer, and M. B. Plenio, *Scalable reconstruction of unitary processes and hamiltonians*, Phys. Rev. A **91**, 042129 (2015), arXiv:1411.6379 [quant-ph].
- [159] G. Torlai, C. J. Wood, A. Acharya, G. Carleo, J. Carrasquilla, and L. Aolita, *Quantum process tomography with unsupervised learning and tensor networks*, Nature Communications **14**, 2858 (2023), arXiv:2006.02424 [quant-ph].
- [160] A. H. Werner, D. Jaschke, P. Silvi, M. Kliesch, T. Calarco, J. Eisert, and S. Montangero, *Positive tensor network approach for simulating open quantum many-body systems*, Phys. Rev. Lett. **116**, 237201 (2016), arXiv:1412.5746 [quant-ph].
- [161] F. A. Pollock, C. Rodríguez-Rosario, T. Frauenheim, M. Paternostro, and K. Modi, *Non-markovian quantum processes: Complete framework and efficient characterization*, Phys. Rev. A **97**, 012127 (2018), arXiv:1512.00589 [quant-ph].
- [162] G. A. L. White, F. A. Pollock, L. C. L. Hollenberg, C. D. Hill, and K. Modi, *From many-body to many-time physics*, arXiv e-prints 10.48550/arXiv.2107.13934 (2021), arXiv:2107.13934 [quant-ph].
- [163] G. A. L. White, F. A. Pollock, L. C. L. Hollenberg, K. Modi, and C. D. Hill, *Non-Markovian Quantum Process Tomography*, PRX Quantum **3**, 020344 (2022), arXiv:2106.11722 [quant-ph].

- [164] Z.-T. Li, C.-C. Zheng, F.-X. Meng, H. Zeng, T. Luan, Z.-C. Zhang, and X.-T. Yu, *Non-Markovian Quantum Gate Set Tomography*, arXiv:2307.14696 [quant-ph] (2023).
- [165] C. Guo, *Reconstructing non-Markovian open quantum evolution from multiple time measurements*, Phys. Rev. A **106**, 10.1103/PhysRevA.106.022411 (2022), arXiv:2205.06521 [quant-ph].
- [166] C. Giarmatzi, T. Jones, A. Gilchrist, P. Pakkiam, A. Fedorov, and F. Costa, *Multi-time quantum process tomography of a superconducting qubit*, arXiv:2308.00750 [quant-ph] (2023).
- [167] G. A. L. White, P. Jurcevic, C. D. Hill, and K. Modi, *Unifying non-Markovian characterisation with an efficient and self-consistent framework*, arXiv e-prints 10.48550/arXiv.2312.08454 (2023), arXiv:2312.08454 [quant-ph].
- [168] E. Nielsen, J. K. Gamble, K. Rudinger, T. Scholten, K. Young, and R. Blume-Kohout, *Gate set tomography*, Quantum **5**, 557 (2021), arXiv:2009.07301 [quant-ph].
- [169] C. Stark, *Simultaneous estimation of dimension, states and measurements: Computation of representative density matrices and POVMs*, arXiv:1210.1105 [quant-ph].
- [170] M. Takahashi, S. D. Bartlett, and A. C. Doherty, *Tomography of a spin qubit in a double quantum dot*, Phys. Rev. A **88**, 022120 (2013), arXiv:1306.1013 [quant-ph].
- [171] R. Blume-Kohout, J. K. Gamble, E. Nielsen, K. Rudinger, J. Mizrahi, K. Fortier, and P. Maunz, *Demonstration of qubit operations below a rigorous fault tolerance threshold with gate set tomography*, Nat. Commun. **8**, 14485 (2017), arXiv:1605.07674 [quant-ph].
- [172] E. Nielsen, R. Blume-Kohout, L. Saldyt, J. Gross, T. L. Scholten, K. Rudinger, T. Proctor, J. K. Gamble, and A. Russo, *pygstio/pygsti: Version 0.9.9.3* (2020).
- [173] E. Nielsen, K. Rudinger, T. Proctor, A. Russo, K. Young, and R. Blume-Kohout, *Probing quantum processor performance with pyGSTi*, Quantum Sci. Technol. **5**, 044002 (2020), arXiv:2002.12476 [quant-ph].
- [174] D. Greenbaum, *Introduction to quantum gate set tomography*, arXiv:1509.02921 [quant-ph] (2015).
- [175] G. De las Cuevas, J. I. Cirac, N. Schuch, and D. Perez-Garcia, *Irreducible forms of matrix product states: Theory and applications*, Journal of Mathematical Physics **58**, 121901 (2017), arXiv:1708.00029 [quant-ph].
- [176] L. Rudnicki, Z. Puchala, and K. Zyczkowski, *Gauge invariant information concerning quantum channels*, Quantum **2**, 60 (2018), arXiv:1707.06926 [quant-ph].
- [177] T. Sugiyama, S. Imori, and F. Tanaka, *Self-consistent quantum tomography with regularization*, Phys. Rev. A **103**, 062615 (2021).
- [178] J. Lin, B. Buonacorsi, R. Laflamme, and J. J. Wallman, *On the freedom in representing quantum operations*, New Journal of Physics **21**, 023006 (2019), arXiv:1810.05631 [quant-ph].
- [179] C. Neill, P. Roushan, K. Kechedzhi, S. Boixo, S. V. Isakov, V. Smelyanskiy, A. Megrant, B. Chiaro, A. Dunsworth, K. Arya, R. Barends, B. Burkett, Y. Chen, Z. Chen, A. Fowler, B. Foxen, M. Giustina, R. Graff, E. Jeffrey, T. Huang, J. Kelly, P. Klimov, E. Lucero, J. Mutus, M. Neeley, C. Quintana, D. Sank, A. Vainsencher, J. Wenner, T. C. White, H. Neven, and J. M. Martinis, *A blueprint for demonstrating quantum supremacy with superconducting qubits*, Science **360**, 195 (2018), arXiv:1709.06678 [quant-ph].

- [180] S. Boixo, S. V. Isakov, V. N. Smelyanskiy, R. Babbush, N. Ding, Z. Jiang, M. J. Bremner, J. M. Martinis, and H. Neven, *Characterizing quantum supremacy in near-term devices*, Nature Physics **14**, 595 (2018), arXiv:1608.00263 [quant-ph].
- [181] Y. Gu, R. Mishra, B.-G. Englert, and H. K. Ng, *Randomized linear gate-set tomography*, PRX Quantum **2**, 030328 (2021), arXiv:2010.12235 [quant-ph].
- [182] T. J. Evans, W. Huang, J. Yoneda, R. Harper, T. Tanttu, K. W. Chan, F. E. Hudson, K. M. Itoh, A. Saraiva, C. H. Yang, A. S. Dzurak, and S. D. Bartlett, *Fast Bayesian tomography of a two-qubit gate set in silicon*, Phys. Rev. Applied **17**, 024068 (2022), arXiv:2107.14473 [quant-ph].
- [183] H.-Y. Huang, S. T. Flammia, and J. Preskill, *Foundations for learning from noisy quantum experiments*, arXiv:2204.13691 [quant-ph] (2022).
- [184] S. Aaronson, *Shadow tomography of quantum states*, in *Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing* (2018) pp. 325–338, arXiv:1711.01053 [quant-ph] .
- [185] A. Gresch and M. Kliesch, *Guaranteed efficient energy estimation of quantum many-body Hamiltonians using ShadowGrouping*, arXiv e-prints , arXiv:2301.03385 (2023), arXiv:2301.03385 [quant-ph].
- [186] A. Dutt, W. Kirby, R. Raymond, C. Hadfield, S. Sheldon, I. L. Chuang, and A. Mezza-capo, *Practical Benchmarking of Randomized Measurement Methods for Quantum Chemistry Hamiltonians*, arXiv:2312.07497 [quant-ph] (2023).
- [187] G. Lugosi and S. Mendelson, *Mean estimation and regression under heavy-tailed distributions—a survey*, Found Comput. Math. **19**, 1145–1190 (2019), arXiv:1906.04280 [math.ST].
- [188] J. Helsen and M. Walter, *Thrifty Shadow Estimation: Reusing Quantum Circuits and Bounding Tails*, Phys. Rev. Lett. **131**, 240602 (2023), arXiv:2212.06240 [quant-ph].
- [189] A. Zhao, N. C. Rubin, and A. Miyake, *Fermionic Partial Tomography via Classical Shadows*, Phys. Rev. Lett. **127**, 110504 (2021), arXiv:2010.16094 [quant-ph].
- [190] M. Arienzo, M. Heinrich, I. Roth, and M. Kliesch, *Closed-form analytic expressions for shadow estimation with brickwork circuits*, Quantum Inf. Comp. **23**, 961 (2023), arXiv:2211.09835 [quant-ph].
- [191] C. Bertoni, J. Haferkamp, M. Hinsche, M. Ioannou, J. Eisert, and H. Pashayan, *Shallow shadows: Expectation estimation using low-depth random Clifford circuits*, arXiv:2209.12924 [quant-ph].
- [192] A. A. Akhtar, H.-Y. Hu, and Y.-Z. You, *Scalable and flexible classical shadow tomography with tensor networks*, Quantum **7**, 1026 (2023), arXiv:2209.02093 [quant-ph].
- [193] M. Ippoliti, Y. Li, T. Rakovszky, and V. Khemani, *Operator Relaxation and the Optimal Depth of Classical Shadows*, Phys. Rev. Lett. **130**, 230403 (2023), arXiv:2212.11963 [quant-ph].
- [194] O. Gühne and G. Tóth, *Entanglement detection*, Phys. Rep. **474**, 1 (2009), arXiv:0811.2803 [quant-ph].
- [195] B. Huang, C. Mu, D. Goldfarb, and J. Wright, *Provable low-rank tensor recovery*, Optimization-Online **4252**, 455 (2014).

- [196] H. Rauhut, R. Schneider, and Z. Stojanac, *Low rank tensor recovery via iterative hard thresholding*, Linear Algebra and its Applications **523**, 220 (2017), arXiv:1602.05217 [cs.IT].
- [197] Y. N. Dauphin, R. Pascanu, C. Gulcehre, K. Cho, S. Ganguli, and Y. Bengio, *Identifying and attacking the saddle point problem in high-dimensional non-convex optimization*, in *Advances in Neural Information Processing Systems*, Vol. 27 (2014) arXiv:1406.2572 [cs.LG].
- [198] C. Piltz, T. Sriarunothai, S. S. Ivanov, S. Wo lk, and C. Wunderlich, *Versatile microwave-driven trapped ion spin system for quantum information processing*, Science Advances **2**, e1600093 (2016), arXiv:1509.01478 [quant-ph].
- [199] D. E. Koh and S. Grewal, *Classical shadows with noise*, Quantum **6**, 776 (2022), arXiv:2011.11580 [quant-ph].
- [200] V. Vitale, A. Rath, P. Jurcevic, A. Elben, C. Branciard, and B. Vermersch, *Estimation of the quantum Fisher information on a quantum processor*, arXiv:2307.16882 [quant-ph] (2023).
- [201] H. Jnane, J. Steinberg, Z. Cai, H. Chau Nguyen, and B. Koczor, *Quantum Error Mitigated Classical Shadows*, (2023), arXiv:2305.04956 [quant-ph].
- [202] J. R. Seddon, B. Regula, H. Pashayan, Y. Ouyang, and E. T. Campbell, *Quantifying Quantum Speedups: Improved Classical Simulation From Tighter Magic Monotones*, PRX Quantum **2**, 010345 (2021), arXiv:2002.06181 [quant-ph].
- [203] A. Zhao and A. Miyake, *Group-theoretic error mitigation enabled by classical shadows and symmetries*, (2023), arXiv:2310.03071 [quant-ph].

Chapter 6

Appendix

A Paper - Compressive gate set tomography

Title: Compressive gate set tomography

Authors: Raphael Brieger, Ingo Roth, Martin Kliesch

Journal: PRX Quantum

Publication status: Published

Contribution by RB: First author (input approx 85%)

A summary of this publication is presented in Chapter 3.

The initial idea to revisit the problem of gate set tomography in a tensor network picture was brought forward by IR and jointly discussed by all authors. I derived all analytical results and developed all code for the optimization algorithm in python, with suggestions from my co-authors in periodic discussions. I wrote an initial draft of the manuscript, whereafter parts of the introduction and main text were added or rewritten by IR and MK. IR further contributed about half of Appendix E. Finally, the manuscript was proofread by all authors and several paragraphs in the main text were jointly edited.

Compressive gate set tomography

Raphael Brieger,^{1,*} Ingo Roth,^{2,3} and Martin Kliesch^{1,4,†}

¹*Institute for Theoretical Physics, Heinrich Heine University Düsseldorf, Germany*

²*Quantum Research Centre, Technology Innovation Institute, Abu Dhabi, UAE*

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

⁴*Institute for Quantum-Inspired and Quantum Optimization, Hamburg University of Technology, Germany*

Flexible characterization techniques that provide a detailed picture of the experimental imperfections under realistic assumptions are crucial to gain actionable advice in the development of quantum computers. Gate set tomography self-consistently extracts a complete tomographic description of the implementation of an entire set of quantum gates, as well as the initial state and measurement, from experimental data. It has become a standard tool for this task but comes with high requirements on the number of sequences and their design, making it experimentally challenging already for only two qubits.

In this work, we show that low-rank approximations of gate sets can be obtained from significantly fewer gate sequences and that it is sufficient to draw them at random. This coherent noise characterization however still contains the crucial information for improving the implementation. To this end, we formulate the data processing problem of gate set tomography as a rank-constrained tensor completion problem. We provide an algorithm to solve this problem while respecting the usual positivity and normalization constraints of quantum mechanics. For this purpose, we combine methods from Riemannian optimization and machine learning and develop a saddle-free second-order geometrical optimization method on the complex Stiefel manifold. Besides the reduction in sequences, we demonstrate numerically that the algorithm does not rely on structured gate sets or an elaborate circuit design to robustly perform gate set tomography. Therefore, it is more flexible than traditional approaches. We also demonstrate how coherent errors in shadow estimation protocols can be mitigated using estimates from gate set tomography.

CONTENTS

I. Introduction	1
II. The data processing problem of gate set tomography	3
A. Compressive gate set description	3
B. Gauge freedom and gate set metrics	5
III. GST data processing via Riemannian optimization	6
A. The complex Stiefel manifold	6
B. The mGST estimation algorithm	7
IV. Numerical analysis	8
A. Gate set and measurement structure	9
B. Number of samples per sequence	10
C. Number of sequences	11
D. Characterizing unitary errors using prior knowledge	12
E. Implementation details and calibration	13
F. Runtime and scaling	14
V. Noise-mitigation for shadow estimation with compressive GST	14
VI. Conclusion and outlook	16

VII. Appendix	18
A. Geodesics on the Stiefel manifold	18
B. Complex Newton equation	21
C. Complex Euclidean gradient and Hessian	25
D. Mean variation error dependence on the choice of objective function	26
E. Noise-mitigation of shadow estimation with GST characterization	27

I. INTRODUCTION

The precise characterization of digital quantum devices is crucial for several reasons: (i) to obtain ‘actionable advice’ on how imperfections on their implementation can be reduced, e.g. by experimental control, (ii) to tailor applications to unavoidable device errors so that their effect can be mitigated, and (iii) to benchmark the devices for the comparison of different physical platforms and implementations. There is already a wide variety of protocols to characterize components of a digital quantum computing device with a trade-off between the information gained about the system and the associated resource requirements and assumptions of the scheme [1, 2].

One particular important requirement for practical characterization protocols for quantum gates is their robustness against errors in the state preparation and measurement (SPAM). There are two general approaches that SPAM-robustly characterize the implementation of entire *gate sets* of a quantum computer. On the low complexity side there is randomized benchmarking (RB) [3–5]

* brieger@hhu.de

† martin.kliesch@tuhh.de

and variants thereof [6], that typically aim at determining a single measure of quality for an experiment, though with the exception of RB tomography protocols [7–10]. Yet for the targeted improvement of individual quantum operations, protocols which provide more detailed information beyond mere benchmarking are crucial.

This is the motivation of self-consistent gate set tomography (GST) [11–18]. GST estimates virtually all parameters describing a noisy implementation of a quantum computing device simultaneously from the measurements of many gate sequences. This comprises tomographic estimates for all channels implementing the gate set elements, the initial state(s) and the measurement(s). The full tomographic information can then be used to compute arbitrary error measures for verification and to provide advice on error mitigation and device calibration [15, 19–23]. Concomitant with the massive amount of inferred information and minimal assumptions, these protocols come with enormous resource requirements in terms of the necessary number of measurement rounds and the time and storage consumption of the classical post-processing. Standard GST, as described, e.g. by Nielsen et al. [18], uses many carefully designed gate sequences in the experiment and a sophisticated and challenging data processing pipeline in post. To arrive at physically interpretable estimates, i.e. completely positive and trace preserving (CPT) maps, additional post-processing is required. The massive amount of specific data consumed by standard GST limits its practical applicability already for two-qubit gate sets. The focus of Nielsen et al. [18] and their implementation `pyGSTi` lies on so-called *long sequence GST*, a method to improve an initial GST estimate by using gate sequences in which a building block is repeated many times. The resulting error amplification of the building block is then used to significantly improve the accuracy of the GST estimate at the cost of larger measurement effort. In our work we focus on *short sequences* and the problem of finding an initial estimate without assuming any prior knowledge on the gate set and minimal experimental requirements.

The most important diagnostic information for a quantum computing device is often already contained in a low-rank approximation of the processes, states and measurements. Coherent errors are typically the ones that can be corrected by experimental control and are of interest for refining calibration models. The strength of incoherent noise on the other hand is arguably well-captured by average error measures as provided by RB outputs. Moreover, current fault tolerance thresholds often rely on worst-case error measures for which no good direct estimation technique exists [24–27] and coherent errors in particular hinder their indirect inference from average error measures [28–30]. For standard state and process tomography, it was realized that low-rank assumptions can crucially reduce the sample complexity, the required number of measurements and the post-processing complexity [9, 31–40] as well as improve the stability against imperfections in the measurements [41] using compressed

sensing techniques [42–44].

In this work, we take a fresh look at the data processing problem of GST from a compressed sensing perspective and regard it as a highly-structured tensor completion problem. We develop a reconstruction method, called `mGST`, that exploits the geometric structure of CPT maps with low Kraus ranks. In numerical simulations we demonstrate that our structure-exploiting `mGST` approach (i) allows for maximal flexibility in the design of gate sequences, so that standard GST gate sequences and random sequences work equally well, and (ii) obtains low-rank approximations of the implemented gate set from a significantly reduced number of sequences and samples. This allows us to successfully perform GST with gate sets and sequences that are not amenable to the standard GST implementation `pyGSTi` [17, 45]. As one example, while the sequence design of `pyGSTi` uses at least 907 specific sequences to reconstruct a two-qubit gate set, we numerically demonstrate low-rank reconstruction from 200 random sequences of maximal length $\ell = 7$ with runtimes of less than an hour on a standard desktop computer. Thus, compressive GST significantly lowers the experimental resource requirements for maybe the most prominent use-cases of GST making it a tool that can be more easily and routinely applied. At the same time, for the default gate sets and sequences from the standard GST implementation, the novel algorithm matches state-of-the-art results. The runtime and storage requirements of `mGST` still scale exponentially in the number of qubits as does the amount parameters of the gate set it identifies. This limits the feasibility of the classical post-processing of compressive GST to gate sets acting on only a few qubits without further assumptions. Nonetheless, we demonstrate that coherent errors and depolarizing noise of a 3-qubit gate set can be completely characterized, from as little as 128 sequences of length $\ell = 7$ on desktop hardware in a few hours.

To give a novel example of how information about coherent errors can be used, we simulate a 10 qubit system and perform GST on neighboring 2-qubit pairs. The resulting gate set estimates then allow us to calibrate the post-processing step of the shadow estimation protocol [46], which is widely used for the sample efficient estimation of observables. More concretely, we find in simulations with moderate coherent errors that shadow estimates of the ground state energy of a 10 qubit Heisenberg Hamiltonian are heavily biased when knowledge of the noisy gate implementation is limited. Information from GST on 2 qubit pairs allows us to reduce this bias by about an order of magnitude.

Our `mGST` reconstruction method relies on manifold optimization over complex Stiefel manifolds [47–53] in order to include the low-rank CPT constraints. Such constraints emerge in several optimization problems [54–57] with applications in machine learning [52, 58], quantum chemistry [59], signal processing in wireless communication [60, 61] and more recently in the quantum information literature, see e.g. [62–64]. In order to deal with

the non-convex optimization landscape we adopt a second order saddle-free Newton method [65] to this setting. This involves the derivation of an analytic expression for geodesics, as well as an expression for the Riemannian Hessian in the respective product manifolds. Another important motivation for phrasing GST as a randomly subsampled tensor completion problem is to bring it closer to potential analytical recovery guarantees common for related tensor completion problems [66–73], opening up a new research direction.

Finally, being able to perform GST from random sequences enables one to use the same type of data for different increasingly refined characterization tasks from filtered RB [6], cross-entropy benchmarking (XEB) [74] and RB tomography [7–9] to GST. Unifying these approaches, random gate sequences can be regarded as the ‘classical shadow’ of a gate set from which many properties can be estimated efficiently [75]. Compressive gate set tomography provides more detailed diagnostic information and only requires to further increase the amount of data without changing the experimental instructions.

With randomized linear GST [76] and fast Bayesian tomography [77] related alternatives to tackle the GST data processing problem have been proposed. Here, the gates are assumed to be well-approximated by an a priori known unitary followed by a noise channel that is either linearized around the identity [76] or around a prior noise estimate [77]. This allows for a treatment of the outcome probabilities as approximately linear functions. The resulting scheme already works for random sequence data but comes at the expense of much stronger assumptions compared to the compressed sensing approach taken with mGST.

The rest of the paper is structured into three parts. In the following section, we formalize the data processing problem of GST as a constrained reconstruction problem. In Section III, we formulate the data processing problem as a geometric optimization task and derive the mGST algorithm. In Section IV we demonstrate the performance of the novel algorithm in numerical simulations and compare our results with the standard GST processing pipeline of pyGSTi.

II. THE DATA PROCESSING PROBLEM OF GATE SET TOMOGRAPHY

In GST a quantum computing device is modeled as follows. The device is initialized with a state $\rho \in \mathcal{S} := \{\sigma \in \mathcal{H} : \sigma \succeq 0, \text{Tr}[\sigma] = 1\}$ on a finite dimensional Hilbert space $\mathcal{H} = \mathbb{C}^d$. Subsequently, a sequence of noisy operations from a fixed gate set $(\mathcal{G}_i)_{i \in [n]}$ can be applied, where we use the notation $[n] := \{1, 2, 3, \dots, n\}$. The noisy operations $\mathcal{G}_i : \mathcal{L}(\mathcal{H}) \rightarrow \mathcal{L}(\mathcal{H})$ are CPT maps on $\mathcal{L}(\mathcal{H})$, the set of linear operators on $\mathcal{H}$. We define $[n]_\ell^* := \bigcup_{k=0}^\ell [n]^k$ such that $\mathbf{i} \in [n]_\ell^*$ defines a gate sequences with length of at most ℓ and associated CPT map $\mathcal{G}_{\mathbf{i}} := \mathcal{G}_{i_\ell} \circ \dots \circ \mathcal{G}_{i_1}$, the concatenation of the gates in the sequence $\mathbf{i}$.

In the end, a measurement is performed described by a positive operator valued measure (POVM) with elements $(E_j)_{j \in [n_E]}$, satisfying $\sum_j E_j = \mathbb{1}$ and $0 \preceq E_j \preceq \mathbb{1}$ for all $j \in [n_E]$. The full description of the noisy quantum computing device is, thus, given by the triple

$$\mathcal{X} = ((E_j)_{j \in [n_E]}, (\mathcal{G}_i)_{i \in [n]}, \rho) \quad (1)$$

of a quantum state, a physical gate set and a POVM. Let $I \subseteq [n]_\ell^*$ be the set of accessible gate sequences with $n_{\text{seq}} := |I|$ denoting the number of sequences. The probability of measuring outcome j upon applying a gate sequence $\mathbf{i} \in I$ is

$$p_{j|\mathbf{i}}(\mathcal{X}) = \text{Tr}[E_j \mathcal{G}_{\mathbf{i}}(\rho)]. \quad (2)$$

By $(\mathbf{p}_{\mathbf{i}}(\mathcal{X}))_j := p_{j|\mathbf{i}}(\mathcal{X})$ we denote the corresponding vector and, moreover, often omit the argument $\mathcal{X}$. While being a fairly general description, this gate set model relies on a couple of assumptions:

- (i) the physical system needs to be well-characterized by a Hilbert space of fixed dimension,
- (ii) the system parameters need to be time independent over different experiments, and
- (iii) a gate’s action is independent of the gates applied before and after (Markovianity).

There exist multiple descriptions of quantum computing devices within the gate set model that yield the same measurement probabilities on all sequences. Below, we provide a more detailed description of this freedom in terms of *gauge transformations*. These are linear transformations that, when simultaneously applied to all gates, input state and POVM elements, leave the measurement statistics (2) invariant.

The task of GST is to infer the device’s full description (1) from measured data. To this end, one estimates the output probabilities for a set of different sequences $I \subset [n]_\ell^*$ by repeatedly performing the measurements of the corresponding sequences. Thus, we can state the *data-processing problem of GST* as follows.

Problem (GST data-processing). *Let $\mathcal{X}$ be a gate set and $I \subset [n]_\ell^*$ a set of sequences. Given empirical estimates $\{y_{j|\mathbf{i}}\}_{\mathbf{i} \in I, j \in [n_E]}$ of $\{p_{j|\mathbf{i}}(\mathcal{X})\}_{\mathbf{i} \in I, j \in [n_E]}$, find the device description $\mathcal{X} = ((E_j)_{j \in [n_E]}, (\mathcal{G}_i)_{i \in [n]}, \rho)$ up to the gauge freedom.*

Note that GST aims at solving an *identification problem*. That is, for sufficiently much data, find the unique device description of the device compatible with the data. In particular, the input data $\{\hat{p}_{j|\mathbf{i}}\}_{\mathbf{i} \in I, j}$ is required to uniquely single out the device description. This is related but distinct from the corresponding *learning task* to find a description that generalizes on unseen data.

A. Compressive gate set description

At the heart of our approach is to capture this data-processing problem as a highly structured *tensor completion problem*. The structure allows us to reduce the

and ranks r_ρ , r_K , and r_E , find the compressive device description $\mathcal{X}_c = (A, \mathcal{K}, B)$ of dimension r_ρ , r_K and r_E respectively, so that the normalization constraints (7), (10), and (11) are satisfied.

As before, the set of sequences needs to be large enough so that this identification problem is well-defined. Again a desired compressive device description $\mathcal{X}_c$ can only be determined up to gauge freedom. Note that for the identification problem to be well-defined, it is not required that the true gate set $\mathcal{X}$ that generated the data is of low-ranks itself. As one usually aims to implement unit rank states, unitary gates, and basis measurements, i.e., for $r_\rho = r_K = r_E = 1$, it can be expected however that a compressive device description $\mathcal{X}_c$ is often also a good approximation to the true gate set. Moreover, coherent errors are arguably the most relevant, since they give actionable advice on error mitigation and are complementary to the incoherent error measures provided by randomized benchmarking experiments. By choosing the ranks in $\mathcal{X}_c$, a problem specific decision can be made that balances the information gained with the computational and sample complexity of model reconstruction. Since the $p_{j|i}(\mathcal{X})$ are high degree polynomials in the gate set parameters, the compressive GST data processing problem is different from compressed sensing for standard state and process tomography, where the map from the model parameters to the outcome probabilities is linear.

Next, we discuss another unique problem of GST, the gauge freedom, more explicitly and introduce relevant error measures for gates sets $\mathcal{X}$.

B. Gauge freedom and gate set metrics

So far we have not made explicit what ‘finding a device description’ actually means. What is well studied in the GST and RB literature [6, 14, 16, 84–86], is that without additional prior assumptions, there is a freedom in representing a device in the gate set model. In particular, this freedom needs to be considered when defining a metric for gate sets [84] w.r.t. which we want to recover the device description.

Gauge freedom refers to the following observation. The observable measurement probabilities $p_{j|i}$ of the form (4) are invariant under the transformation

$$\rho \mapsto \mathcal{T}^{-1}(\rho) \quad (13)$$

$$\mathcal{G}_i \mapsto \mathcal{T}^{-1} \circ \mathcal{G}_i \circ \mathcal{T} \quad \forall i \quad (14)$$

$$E_j \mapsto \mathcal{T}^\dagger(E_j) \quad \forall j \quad (15)$$

for any invertible super operator $\mathcal{T} : \mathcal{L}(\mathcal{H}) \rightarrow \mathcal{L}(\mathcal{H})$, where $\mathcal{T}^\dagger$ denotes the adjoint of $\mathcal{T}$ w.r.t. the Hilbert-Schmidt inner product. This invariance is also the well-known *gauge freedom of MPS* [87].

If the gate set is universal and the initial state is pure then the gauge transformation $\mathcal{T}$ has to be either a uni-

tary or anti-unitary channel. This statement can be seen as follows.

In our case the $\mathcal{G}_i$ are constrained to be CPT. Hence, $\mathcal{G} \mapsto \mathcal{T}^{-1}(\mathcal{G})\mathcal{T}$ has to map CPT maps to CPT maps. Similarly, $\mathcal{T}^{-1}\rho$ has to be a density operator and $\{\mathcal{T}^\dagger(E_j)\}_j$ a valid POVM.

A more explicit condition on $\mathcal{T}$ can be obtained by considering gauge action on entire sequences. For all sequences $\mathbf{i}$, we have

$$p_{j|\mathbf{i}} = \text{Tr}[E_j \mathcal{T} \circ \mathcal{T}^{-1} \circ \mathcal{G}_{\mathbf{i}}(\rho)]$$

where $\mathcal{G}_{\mathbf{i}}(\rho)$ is a positive operator if the gates $\mathcal{G}_i$ are CPT. Now $\mathcal{T}^{-1}\mathcal{G}_{\mathbf{i}}(\rho)$ has to be positive as well for all sequences $\mathbf{i}$. Thus if the gate set is universal, the map $\mathcal{T}^{-1}$ has to be positive and trace preserving for all states. An analogous statement can be made for $\mathcal{T}^\dagger$, by considering that $\mathcal{T}^\dagger \mathcal{G}_{\mathbf{i}}^\dagger(E_j)$ has to be positive-definite for all POVM elements E_j . This implies that $\mathcal{T}$ has to be a positive map as well.

It has been shown that any positive invertible map $\mathcal{T}$ with a positive inverse can be written either as $\mathcal{T}(\rho) = PU\rho U^\dagger P^\dagger$ or $\mathcal{T}(\rho) = PU\rho^T U^\dagger P^\dagger$ for $U \in \text{GL}(d, \mathbb{C})$ and $P \succeq 0$ [88, Theorem 2]. The condition that $\mathcal{T}$ needs to be trace preserving then yields $P = \mathbb{1}$, as can be seen from the Kraus decomposition. Hence, $\mathcal{T}$ is indeed either a unitary or anti-unitary channel.

We note that the map $\rho \mapsto U\rho^T U^\dagger$ is positive but not completely positive. However, it has the property that $\mathcal{T}^{-1}\mathcal{G}_i\mathcal{T}$ is CPT whenever $\mathcal{G}_i$ is CPT. This can be seen by observing that the Choi matrix of $\mathcal{T}^{-1}\mathcal{G}_i\mathcal{T}$ is given by $(U^* \otimes U)\text{Choi}(\mathcal{G}_i)^T(U^T \otimes U^\dagger)$, which is positive definite for $\mathcal{G}_i$ being CPT.

However, actual gate set implementations are noisy and hence not universal in the sense that they cannot prepare any pure state. Therefore, in practice, the gauge freedom can be larger [18]. For instance, if all gate implementations are given by unital channels then an additional freedom exists: depolarizing noise can be commuted through the circuit. Therefore, it can be distributed arbitrarily among initial state, gates, and measurement. Meaningful distance measures for gate sets should have the same gauge freedom as the GST data. The problem of finding gauge invariant distance has been studied by Lin et al. [84]. For individual gate sequences, any measure that compares only the ideal and observed outcome probabilities is naturally gauge invariant. The authors thus propose to use the total variation error, a natural error measure to compare probability distributions, for individual gate sequences. Let

$$p_{j|\mathbf{i}}(E_j, \mathcal{G}_{\mathbf{i}}, \rho) := \text{Tr}[E_j \mathcal{G}_{\mathbf{i}}(\rho)] \quad (16)$$

denote the probabilities of measuring the j th output of the POVM with elements E_j after applying the sequence $\mathbf{i}$ of gates in $\mathcal{G}_{\mathbf{i}}$ to the state ρ . The *total variation error* for sequence $\mathbf{i}$ between two gate sets $\hat{\mathcal{X}} = \{(\hat{E}_j), \hat{\mathcal{G}}, \hat{\rho}\}$

and $\mathcal{X} = ((E_j), \mathcal{G}, \rho)$ is defined as

$$\delta d_{\mathbf{i}}(\hat{\mathcal{X}}, \mathcal{X}) := \frac{1}{2} \sum_j \left| \text{Tr} \left[\hat{E}_j^\dagger \hat{\mathcal{G}}_{\mathbf{i}}(\hat{\rho}) \right] - \text{Tr} \left[E_j^\dagger \mathcal{G}_{\mathbf{i}}(\rho) \right] \right|. \quad (17)$$

The *mean variation error (MVE)* is defined as [84]

$$\text{MVE}_I(\hat{\mathcal{X}}, \mathcal{X}) := \mathbb{E}_{\mathbf{i} \sim I} [\delta d_{\mathbf{i}}(\hat{\mathcal{X}}, \mathcal{X})] \quad (18)$$

w.r.t. a set of sequences I , where $\mathbf{i} \sim I$ means that $\mathbf{i}$ is drawn uniformly from I . Often, we omit the subscript I in the following. The MVE corresponds to taking the natural worst case error measure over the measurement outcomes (the total variation distance) and averaging it over the available gate sequences. Often I is chosen as the set of all gate sequences up to some length ℓ . Then the expectation value (18) contains a sum over exponentially many terms. However, since they are all non-negative, they can be estimated sampling efficiently via Monte Carlo sampling [84].

A closely related error measure is the *mean squared error (MSE)*

$$\mathcal{L}_I(\hat{\mathcal{X}}, \mathcal{X}) := \mathbb{E}_{\mathbf{i} \sim I} \sum_{j \in [n_E]} \left(\text{Tr} \left[\hat{E}_j^\dagger \hat{\mathcal{G}}_{\mathbf{i}}(\hat{\rho}) \right] - \text{Tr} \left[E_j^\dagger \mathcal{G}_{\mathbf{i}}(\rho) \right] \right)^2, \quad (19)$$

which averages the squared deviation over all sequences and POVM elements.

III. GST DATA PROCESSING VIA RIEMANNIAN OPTIMIZATION

In the previous section, we defined the compressive GST data processing problem and introduced metrics for the quality of reconstruction. We now turn to devising a concrete algorithm for the data processing problem. To this end, we formulate the reconstruction problem as a constraint optimization problem of a loss-function for the data fitting. A natural candidate for the loss-function is the MVE restricted to the set of measured sequences. As a proxy we instead minimize the MSE which depends smoothly on the gate set and is, therefore, more suitable for local optimization. In terms of the compressive device description, the MSE (19) can be written as

$$\mathcal{L}_I(A, \mathcal{K}, B|\mathbf{y}) := \frac{1}{|I|} \sum_{\mathbf{i} \in I} \sum_j (p_{j|\mathbf{i}}(A, \mathcal{K}, B) - y_{j|\mathbf{i}})^2 \quad (20)$$

where $y_{j|\mathbf{i}}$ is the empirical estimate of $\text{Tr}[E_j \mathcal{G}_{\mathbf{i}}(\rho)]$. Correspondingly, the compressive GST data processing prob-

lem can be cast as the constraint optimization problem:

$$\begin{aligned} & \underset{A, \mathcal{K}, B}{\text{minimize}} && \mathcal{L}_I(A, \mathcal{K}, B|\mathbf{y}) \\ & \text{subject to} && \sum_{l=1}^{r_K} \mathcal{K}_{il}^\dagger \mathcal{K}_{il} = \mathbb{1} \quad \forall i \in [n], \\ & && \sum_{j=1}^{r_E} A_j^\dagger A_j = \mathbb{1}, \\ & && \|B\|_F = 1. \end{aligned} \quad (21)$$

The constraints restrict the objective variables to embedded matrix manifolds. Therefore, algorithms for the optimization problem can be derived by generalizing standard optimization algorithms for functions on the Euclidean space to the geometric structure of these manifolds.

A. The complex Stiefel manifold

In order to formulate our main reconstruction algorithm we need to understand the matrix manifold that encompasses the physicality constraints mentioned in Section II A. We start by summarizing the elementary properties of these manifolds, to then derive a parametrization of geodesics and the Riemannian Hessian, thereby extending what was previously done for their real counterparts in Ref. [48]. For a comprehensive introduction to optimization on matrix manifolds we refer to the book by Absil, Mahony and Sepulchre [89].

Let $(K_l)_{l \in [r]}$ be the Kraus operators of a fixed gate. By stacking them along their row dimension to a new matrix $K \in \mathbb{C}^{dr \times d}$, we can write the CPT constraint as $K^\dagger K = \mathbb{1}$. In the following we set $D = dr$. The set

$$\text{St}(D, d) := \{K \in \mathbb{C}^{D \times d} : K^\dagger K = \mathbb{1}_d\} \quad (22)$$

is called the $D \times d$ *complex Stiefel manifold*. This manifold is the set of isometries of the Euclidean space and contains the special cases of the sphere $\text{St}(D, 1)$ and the unitary matrices $U(D) = \text{St}(D, D)$. We regard it here as a submanifold of $\mathbb{C}^{dr \times d}$.

The tangent space of $\text{St}(D, d)$ at K is given by

$$T_K \text{St}(D, d) = \{\Delta \in \mathbb{C}^{D \times d} : K^\dagger \Delta = -\Delta^\dagger K\}. \quad (23)$$

The canonical inner product of $\Delta_1, \Delta_2 \in T_K \text{St}(D, d)$ can be defined as

$$\langle \Delta_1, \Delta_2 \rangle_K = \text{Re} \left\{ \text{Tr}(\Delta_1^\dagger \Gamma \Delta_2) \right\} \quad (24)$$

with $\Gamma = \mathbb{1} - \frac{1}{2} K K^\dagger$. Another choice is the standard Hilbert-Schmidt inner product of the embedding matrix space. However, the advantage of the canonical inner product is that it weights all degrees of freedom on the tangent space equally. The Stiefel $\text{St}(D, d)$ together with

the metric given by (24) is a Riemannian manifold. The *normal space* is defined by

$$N_K \text{St}(D, d) = \{ \Delta_\perp \in \mathbb{C}^{D \times d} : \langle \Delta, \Delta_\perp \rangle_K = 0 \\ \forall \Delta \in T_K \text{St}(D, d) \}.$$

The projector onto the normal space at position K is given by

$$P_N(X) = K(K^\dagger X + X^\dagger K)/2 \quad (25)$$

for $X \in \mathbb{C}^{dr \times d}$ and we can write the projector onto the tangent space at K as

$$P_T(X) = X - P_N(X). \quad (26)$$

We wish to optimize the MSE over $\text{St}(D, d)$. In analogy to the optimization over $U(n)$ in [47], we will move along geodesics, which are the locally length minimizing curves. In Appendix A we show that within $\text{St}(D, d)$, a geodesic starting at $K_{t=0} \equiv K$ and going in the direction $\Delta \in T_K \text{St}(D, d)$ can be written as

$$K_t(K, \Delta) = (K \ Q) \exp \left[t \begin{pmatrix} K^\dagger \Delta & -R^\dagger \\ R & 0 \end{pmatrix} \right] \begin{pmatrix} \mathbf{1} \\ 0 \end{pmatrix}, \quad (27)$$

with Q, R given by the QR decomposition of $(\mathbf{1} - K K^\dagger) \Delta$. Note that $\dot{K}_t|_{t=0} = \Delta$. Often simpler curves that just satisfy $K_0 = K$ and $\dot{K}_t|_{t=0} = \Delta$ are used instead of the geodesic in order to save computation time [89]. However, computing the exponential of the $2d$ -dimensional matrix in Eq. (27) provides no bottleneck in our scenario as the inversion of the $2nd^2 r_K$ -dimensional Hessian is more costly (see Section IV F).

In order to identify the Riemannian gradient and Hessian, we generalize results from the real case [48] to the complex case. Then we use the second order Taylor approximation of the objective function, which will be given below in terms of the MSE (19) along geodesics (see Appendix B). The same treatment can be applied to the POVM given by the matrices A_j from the decomposition (9), where we define A as the matrix obtained from stacking the A_j along their row dimension. The physicality constraint on A is then equivalent to $A \in \text{St}(dn_E, r_E)$ with n_E being the number of POVM elements and r_E their maximal rank. Finally, the constraint $\|B\|_F^2 = \text{vec}(B)^\dagger \text{vec}(B) = 1$ on the initial state (9) can also be captured by the Stiefel manifold via the requirement $\text{vec}(B) \in \text{St}(dr_\rho, 1)$, where $\text{vec}(B) \in \mathbb{C}^{dr_\rho}$ is the vectorization of $B \in \mathbb{C}^{d \times r_\rho}$.

B. The mGST estimation algorithm

With a better understanding of the underlying manifold structure we can now formulate a concrete optimization approach to tackle the estimation problem (21). The least squares cost function (19) is a polynomial of order at most the sequence length squared in the parameters of

$\mathcal{G}$, with a highly degenerate global minimum due to the gauge freedom. In analogy to the alternating minimization techniques which are successful for matrix product state completion [66, 90, 91] we alternate between updates on $A, \mathcal{K}$ and B . Each update would naively be done via a local optimization approach such as gradient descent. However, we observe that following the gradient direction on the respective manifolds is problematic around saddle points, which are frequently encountered in our optimization problem. In principle the gradient direction points away from saddle points, yet the norm of the gradient can be arbitrarily small. There are different approaches in the literature to deal with this problem. For instance information about the curvature can be included [65] or, if a saddle point is encountered, random update directions can be chosen to escape the area of vanishing gradient [50, 92]. We find that the so-called saddle free Newton (SFN) method [65] yields considerably better results than first order methods. There the update direction is given by $-|H|^{-1}g$ with H being the Hessian and g the gradient and the absolute value $|H|$ define by spectral calculus. An instructive way to see why this leads to a speedup is to write the Hessian H as $H = \sum_i \lambda_i |v_i\rangle\langle v_i|$, where v_i is the eigenvector to eigenvalue λ_i . The update direction of the SFN method then reads $-|H|^{-1}g = -\sum_i |\lambda_i|^{-1} |v_i\rangle\langle v_i|g$. Since the vectors $|v_i\rangle$ form a basis, this can be interpreted as a rescaling of $|g\rangle$ by $|\lambda_i|^{-1}$ in the directions $|v_i\rangle$. As with the standard Newton method, this leads to a large rescaling if the curvature in a particular direction is small, resulting in large steps even close to the saddle point. Taking the absolute value of the eigenvalues then ensures that saddle points are repulsive. For numerical stability it is beneficial to introduce a damping term that offsets the eigenvalues of H that are very close to zero before the inversion.

Algorithm 1 describes a single step of the damped saddle-free Newton method with damping parameter λ and is a generalization of the original SFN method [65] to manifolds.

Algorithm 1: SFN update

input: Curve parametrization $Y_t(Y_0, \Delta)$, objective function $\mathcal{L}_I(Y_t)$, damping parameter λ

- 1 Compute the gradient $\mathbf{G}$ and Hessian H of $\mathcal{L}_I(Y_t)$ at Y_0 .
- 2 Determine the update direction

$$\begin{pmatrix} \Delta \\ \Delta^* \end{pmatrix} = (|H| + \lambda \mathbf{1})^{-1} \begin{pmatrix} \mathbf{G}^* \\ \mathbf{G} \end{pmatrix}.$$

- 3 Determine the step size $\tau = \underset{t}{\text{argmin}} \mathcal{L}_I(Y_t(Y_0, \Delta))$.

return $Y_\tau(Y_0, \Delta)$

Algorithm 1 is formulated in a way that is compatible with an update in Euclidean space as well as an update on the Stiefel manifold. In Euclidean space we update along the curve $Y_t(\cdot, \cdot) : \mathbb{C}^{D \times d} \times \mathbb{C}^{D \times d} \rightarrow \mathbb{C}^{D \times d}$ with

$Y_t(Y_0, \Delta) = Y_0 + t\Delta$ for an update direction Δ . On the Stiefel manifold we have $Y_t(\cdot, \cdot) : \text{St}(D, d) \times \mathfrak{T} \rightarrow \text{St}(D, d)$ with the curve given by the geodesic (27) and $\mathfrak{T}$ being the tangent bundle on $\text{St}(D, d)$. The step size is determined by locally optimizing over the parameter t using standard gradient free optimizers. We derive an expression for the Hessian on the Stiefel manifold in Appendix B. In Appendix C, we also provide a detailed discussion and expressions for the optimization in complex Euclidean space.

Algorithm 2: mGST

input: Data $\{\mathbf{y}_{j|i}\}_{i \in I, j \in [n_E]}$, batch size κ , Kraus rank r_K , initialization $(A^0, \mathcal{K}^0, B^0)$, stopping criterion

```

1  $i \leftarrow 0$ 
2 repeat
3   Select batch  $J \subset I$  of size  $|J| = \kappa$  at random
4    $A^{i+1} \leftarrow$  update  $A^i$  with objective  $\mathcal{L}_J(\cdot, K^i, B^i; \mathbf{y})$ 
    along geodesic on  $\text{St}(dn_E, d)$ 
5    $\mathcal{K}^{i+1} \leftarrow$  update  $\mathcal{K}^i$  with objective
     $\mathcal{L}_J(A^{i+1}, \cdot, B^i; \mathbf{y})$  along geodesic on  $\text{St}(r_K d, d)^{\times n}$ 
6    $B^{i+1} \leftarrow$  update  $B^i$  with objective
     $\mathcal{L}_J(A^{i+1}, \mathcal{K}^{i+1}, \cdot; \mathbf{y})$  along geodesic on  $\text{St}(d^2, 1)$ 
7    $i \leftarrow i + 1$ 
8 until stopping criterion is met at  $i = i_*$ ;
9 return  $(A^{i_*}, \mathcal{K}^{i_*}, B^{i_*})$ 
```

Algorithm 2 describes the main mGST routine. It can be run with different choices of smooth objective functions, and we use the MSE (19) by default. In our numerics we often find that optimizing the log-likelihood function after the MSE can improve estimates (see Appendix D for a discussion).

The algorithm alternates updates on A , $\mathcal{K}$ and B . Updates are performed using Algorithm 1 on the tangent spaces of the respective Stiefel manifolds. In order to achieve good convergence, we run the optimization with mGST in two consecutive steps: we start from a random initialization and perform a coarse grained optimization with a small batch size κ , i.e. only using κ many random gate sequences from I for each update step. The batching of data results in lower computation time for the derivatives and adds a factor of randomness to the optimization, which avoids getting stuck at suboptimal points to a certain degree. We terminate the first optimization loop when the objective function $\mathcal{L}_I(A^i, \mathcal{K}^i, B^i | \mathbf{y})$ is smaller than an *early stopping value* δ , which is obtained from the data as follows.

For a number of m samples per sequence the outcome probabilities of each sequence for the true gate set are given by

$$y_{j|i} = k_{j|i}/m, \quad (28)$$

where $k_{j|i}$ is the number of times outcome j is measured upon applying the gate sequence i . Due to Born's rule, $k_{j|i}$ is distributed according to the multinomial distribution $M(m, (p_{1|i}, \dots, p_{n_E|i}))$ with probabilities $\{p_{j|i}\}_j$ and

m trials. We estimate the expectation value of the objective function from the values $y_{j|i}$. This provides us with a rough estimate for how low the objective function value can become, given the sample counts $k_{j|i}$. Then we set the early stopping value to be twice that estimate,

$$\delta := 2 \mathbb{E}_{\tilde{k}_{j|i} \sim M(m, (y_{j|i}))} \frac{1}{|I|} \sum_{i \in I} \sum_j \left(y_{j|i} - \tilde{k}_{j|i}/m \right)^2. \quad (29)$$

Hence, we require the objective function on the full data set to be close to its expectation value for the measured probabilities $y_{j|i}$ obtained from m samples.

While computationally inexpensive the mini-batch stochastic optimization does not converge to an optimal point on the full data set I . In a second optimization loop, we initialize the mGST algorithm with the result from the first run and use all the data for the updates. Formally, we choose the batch size $\kappa = |I|$ and, thereby, make the random batch selection obsolete. We perform these more costly update steps until the change in objective function reaches a desired relative precision ϵ ,

$$\mathcal{L}_I(A^i, \mathcal{K}^i, B^i | \mathbf{y}) - \mathcal{L}_I(A^{i-1}, \mathcal{K}^{i-1}, B^{i-1} | \mathbf{y}) \leq \delta \epsilon \quad (30)$$

or a maximal number of iterations is exceeded.

The first optimization run is initialized with a random gate set parameterized by $A^0, \mathcal{K}^0$ and B^0 (see Section II A). For the random initialization we make use of the *Gaussian unitary ensemble* (GUE). A matrix H belongs to the GUE if $H = (M + M^\dagger)/2$, where M is a *complex Gaussian matrix*, i.e. real and imaginary part of each M_{ij} are independently drawn from $\mathcal{N}(0, 1)$, the normal distribution with zero mean and unit variance. In this case we write $H \sim \text{GUE}$. For A^0 and each gate in $\mathcal{K}^0$ we take the first d columns of e^{iH} with $H \sim \text{GUE}$ to obtain a random isometry $\mathcal{K}^0$. For B_0 we take a complex Gaussian matrix and normalize it such that $\text{Tr}[B^{0\dagger} B^0] = 1$.

Importantly, due to the nature of non-convex optimization, several initializations can be needed to converge to a satisfactory minimum.

IV. NUMERICAL ANALYSIS

In this section, we evaluate the performance of mGST in different scenarios in numerical simulations. In particular, we compare its performance to the state-of-the-art implementation for gate set tomography, pyGSTi [18], in the regimes where both methods can be applied.

For pyGSTi to be applicable one has to use structured gate sequences inspired by standard quantum process tomography. In Section IV A we evaluate the performance of mGST and pyGSTi on minimal measurement sequences and different models to find that mGST benefits from flexibility in the sequence design and a fully general model parametrization. Section IV B numerically validates the expected inverse square-root scaling of the reconstruction error with the number of measurement samples per sequence for different noise regimes.

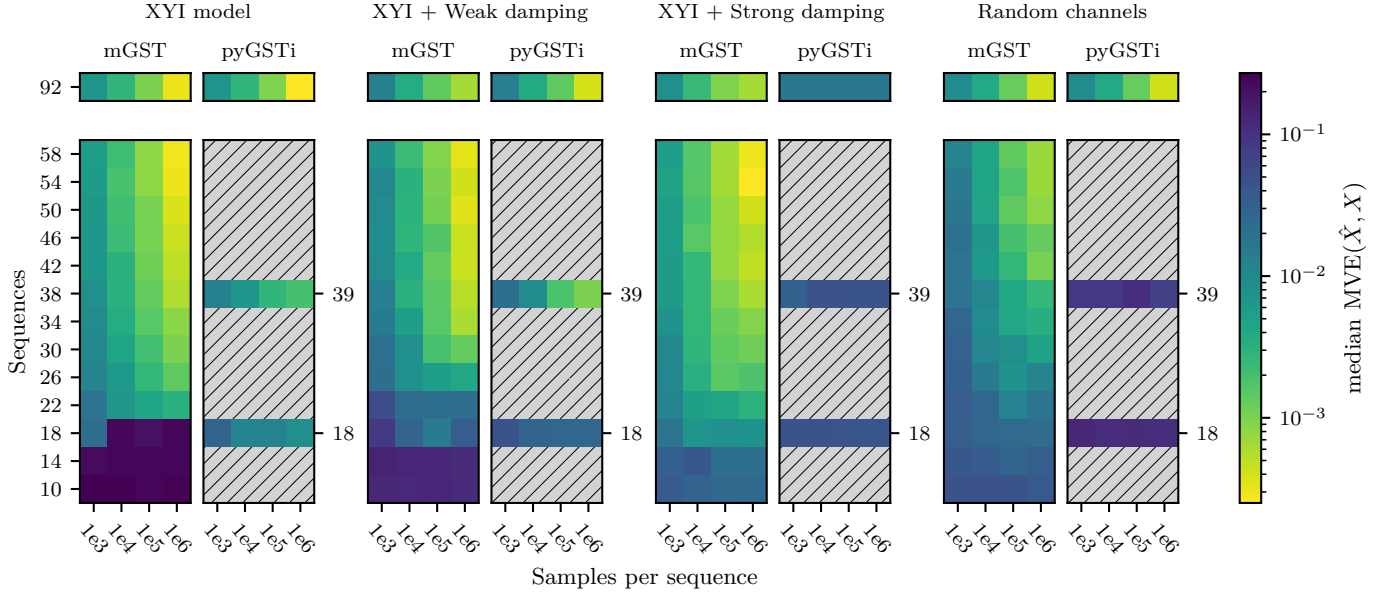


Figure 1. mean variation error (MVE) comparison between mGST (with log-likelihood cost function and $r_K = 4$) and pyGSTi showing the dependence on the number of sequences, the number of samples per sequence and the gate set on a single qubit. The number of sequences used by mGST in the range 10-58 are drawn uniformly at random, while the sequences for pyGSTi need to follow the pyGSTi fiducial design and are limited to the fixed sequence counts (18 and 39). We choose independently drawn random fiducials for each instance. The 92 sequences used by both mGST and pyGSTi are taken from the standard pyGSTi sequence design for the XYI-model. All sequences are of length $\ell = 7$. The MVE depicted in each square is the median result for 10 instances, each with random statistical measurement noise and a random sequence drawn from the uniform distribution. In the random channel scenario, a new random channel is used for each instance. The XYI-model is a simple unitary model used in the GST literature and the weak damping model consists of amplitude damping noise on each gate with $\Gamma = 0.94$, while the strong damping model uses $\Gamma = -0.6$. A complete description of the models used can be found in the main text. Each model has additional depolarizing noise of strength $p = 0.01$ on the initial state.

Section IV C numerically determines the required number of random sequences to accurately reconstruct simple and random gate set models with mGST for different Kraus ranks. In Section IV D we follow up with a numerical demonstration of unitary noise characterization for a three-qubit gate set using a priori knowledge in the initialization. Finally, in Sections IV E and IV F we discuss the choice of initialization and hyperparameters, as well as the runtime of mGST.

For a model of n gates reconstructed from m measurements of sequence length ℓ , we validate the performance of mGST by computing the MVE (18) over all possible n^ℓ sequences, or 10^4 random sequences of length ℓ if $n^\ell > 10^4$. Usually $m \ll \min(n^\ell, 10^4)$ and the MVE can be thought of as a generalization error on the predicted output probabilities of the gate set estimate. The gate sets studied in this section all use the same target initial state $|0\rangle\langle 0|$ and computational basis measurement, although with different levels of noise applied to them. For instance, we often use global depolarizing noise, which acts on a quantum state ρ as $\rho \mapsto (1 - p)\rho + p\mathbb{I}/d$. For the numerics presented here, we use a maximum of 100 reinitializations (if not stated otherwise). A discussion of the required number of initializations is given in Section IV E. A Python implementation of mGST and a short

tutorial can be found on GitHub [93].

A. Gate set and measurement structure

We compare mGST and pyGSTi for the minimal number of sequences doable with each method and for gate sets of different conditionings, without using the compression capabilities of mGST yet. We find that mGST is more flexible in the sequence design and model parametrization, while generating estimators with lower mean variation errors in several regimes.

The traditional strategy for GST, akin to standard quantum process tomography, is to generate a frame for $L(\mathcal{H})$, measure each gate in that frame and generate an estimate for each gate by applying the pseudo-inverse of the measurement operator.

This is particularly important for the first reconstruction step in pyGSTi where the sequences that generate the frame are called *fiducials*. The strategy of pyGSTi is to obtain an initial estimate via the pseudo-inverse, followed up by local optimization of a particular cost function [18]. In contrast, we perform mGST using random initializations and, thereby, not rely on designated fiducial sequences.

Figure 1 compares mGST to pyGSTi focusing on the regime of very few gate sequences, showing what is needed in terms measurement effort to obtain low mean variation errors for different gate sets. In order to test pyGSTi in the regime of low sequence counts, we replace the 5 standard fiducial sequences with 2 or 3 fiducial sequences drawn uniformly at random, thereby reducing the total sequence number from 92 to 18 or 39 sequences. Since mGST is compatible with any sequences design, we use between 10 and 58 random sequences for mGST to explore the low sequence count region.

The first gate set we study is the so-called *XYI model*, the standard single qubit example in the pyGSTi package [45]. The XYI model consists of the identity gate, a $\pi/2$ X-rotation and a $\pi/2$ Y-rotation on the Bloch sphere, with initial state $|0\rangle\langle 0|$ and measurement in the computational basis. Results for the ideal XYI-model can be seen on the left in Figure 1, with mGST and pyGSTi performing identically for 92 sequences. Comparing the results for 18 sequences we find that mGST does not converge on more than 50% of trials, which reflects in the median MVE being above 10^{-1} , while pygsti achieves lower median errors. Comparing the 38 and 39 sequence medians however, we find that mGST yields lower error models than pyGSTi.

The subsequent models analyzed successively deviate from the simple unitary XYI-model and highlight the versatility of our manifold approach. Since the full CPT parametrization used in our optimization (21) is agnostic to any special gate set properties we expect it to perform well for all possible CPT maps as gate implementations. For instance, for random and specific non-Markovian channels. pyGSTi on the other hand uses a parametrization of Lindblad type and is therefore based on a more limited model space.

To illustrate this comparison, we perturb the XYI-model by adding amplitude damping noise to each gate. The amplitude damping channel can be written in terms of the Kraus operators $K_1 = \begin{pmatrix} 1 & 0 \\ 0 & \Gamma \end{pmatrix}$ and $K_2 = \begin{pmatrix} 0 & \sqrt{1 - |\Gamma|^2} \\ 0 & 0 \end{pmatrix}$, which arise e.g. from the Jaynes-Cummings model of a qubit system interacting with a quantized bosonic field [94]. How well mGST and pyGSTi perform on a model with $\Gamma = 0.94$ can be seen in the center left block of Figure 1 (XYI + Weak damping). We find generally similar performance, with mGST being a bit more accurate on 38 sequences, and a bit less accurate on 92 sequences with 10^6 samples.

Increasing the interaction time between qubit and environment leads to memory effects and strong non-Markovianity of the amplitude damping channel at $\Gamma = -0.6$. This scenario is shown in the center right plot of Figure 1, and we see that while the accuracy of mGST is the same as before, the model parametrization of pyGSTi cannot fit the model with MVEs below 10^{-2} , independent of the sequence or sample count.

For the last comparison (rightmost block in Figure 1)

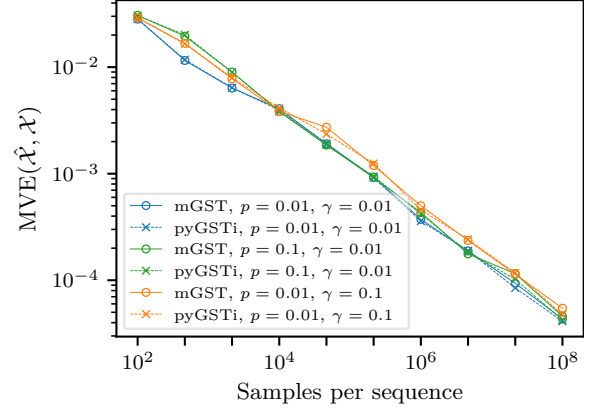


Figure 2. Reconstruction of the XYI gate set for different levels of depolarizing noise with strength p and unitary noise with strength γ on each gate. The unitary noise is given by $e^{i\gamma H}$ with $H \sim \text{GUE}$. Additional depolarizing noise with $p = 0.01$ is applied before measurement. The mGST-algorithm is run on the log-likelihood cost function with $r_K = 4$ (max.), which is the same as for pyGSTi. As gate sequences we used again the standard pyGSTi fiducial sequences, with the number of measurements per sequence between 10^2 and 10^8 . The lines connect data points of which each is the median over 10 runs. For each run a new random overrotation is drawn and new measurements are simulated. The measurement sequences are the 92 sequences provided by the pyGSTi software, with a maximum sequence length of $\ell \leq 7$.

we look at the performance for random full Kraus rank channels. Each channel is constructed by drawing a Haar random $d^3 \times d^3$ unitary and then taking its first d columns. The resulting $d^3 \times d$ matrix is an isometry and therefore constitutes a valid set of Kraus operators. Note that this construction is different from the previous construction of random channels via the Gaussian unitary ensemble. The results show that mGST can reconstruct these models from low sequence counts, while pyGSTi does not yield good estimators. Using the standard sequence design of 92 sequences, mGST and pyGSTi have identical accuracy again, suggesting that random channels are typically well within the model space of pyGSTi after all. These demonstrations show that mGST is indeed flexible in the sequence design with state-of-the-art performance for arbitrary gate set implementations.

B. Number of samples per sequence

The probability associated to every sequence is estimated from a finite number of samples. Here, we study the resulting effect on the reconstruction accuracy as measured by the MVE more closely.

For a high number m of samples per sequence, each probability y_i^j in the objective function is estimated with an error of order $1/\sqrt{m}$. Therefore, we expect the MVE to also decrease as $1/\sqrt{m}$ if the algorithm converges to

the global minimum. This scaling was observed to hold true for pyGSTi [18]. In order to be able to compare the scaling of mGST directly to the one of pyGSTi, we use a standard pyGSTi setting: The gate set is the XYI-model (with $\pi/2$ -rotations) and the gate sequences are the standard pyGSTi sequences for this model with a maximum sequence length of $\ell = 7$.

We add noise to the gate set by varying the amount of depolarizing noise with strength p on each gate and also overrotating each gate by a random unitary. The random unitaries are given by $e^{i\gamma H}$ with $H \sim \text{GUE}$. In particular, this means that H can be bounded on average as follows. We can write $H = (M_1 + M_1^T + i(M_2 - M_2^T))/2$ with M_i being independent Gaussian matrices. Next, we use Gordon's theorem for Gaussian matrices (see e.g. [95, Theorem 5.32]), which tells us that $\mathbb{E} \|M_1\|_\infty \leq 2\sqrt{d}$. The relevant magnitude of the random generator H is then in expectation upper bounded as

$$\mathbb{E} \|H\|_\infty \leq 2 \mathbb{E} \|M_1\|_\infty \leq 4\sqrt{d}. \quad (31)$$

State preparation and measurement are assumed to be noise-free in this setup, however for a fixed sequence length the depolarizing noise per gate is equivalent to a global depolarizing channel applied before measurement, since it commutes with the unitary gates. Figure 2 depicts the resulting MVE-scaling of the reconstruction where data was generated using different numbers of samples per sequence m .

We observe that mGST follows the expected scaling in m , matching the scaling of pyGSTi for different levels of unitary and depolarizing noise.

C. Number of sequences

The arguably most challenging experimental requirement of GST is the number of measurement settings (sequences) that are required for a successful gate reconstruction. One of the main motivations of compressive GST is to employ structure constraints, i.e. to reduce the number of degrees of freedom of the reconstruction problem, in order to reduce the required number of measurements. Instead of reconstructing arbitrary quantum channels we aim at reconstructing low-rank approximations of the gate set elements. In addition, we expect that by using the mGST algorithm, compressive recovery is possible from already a 'few' randomly selected sequences. We here numerically demonstrate that this is indeed the case.

The top row of Figure 3 shows the median performance in MVE against the number of randomly chosen sequences for different Kraus ranks. On the left are the results for the single qubit XYI model as defined in Section IV A. On the right are the results for the XYICNOT gate set that is based on the identity, CNOT and Pauli-X and -Y rotations on each qubit individually, with rotation angle $\pi/2$.

We observe a phase transition in the MVE that indicates a minimal number of sequences that are required for the successful reconstructions of the gate sets. As expected, constraining the reconstruction to a lower Kraus rank indeed reduces the amount of required sequences in the reconstruction in most cases.

An intriguing exception is the $r_K = 1$ reconstruction of the two qubit gate set that exhibits the worst reconstruction performance compared to higher rank constraints. We suspect that this is due to the optimization problem being more dependent on the initialization for $r_K = 1$. In more general settings, it has been observed that the optimization over matrix-product states with fixed Kraus rank can be unstable and using rank-adaptive optimization techniques yield much better performance [96, 97]. This motivates to use a slightly higher rank in the optimization than the expected rank of an effective approximation of the gate set. In accordance with this intuition, we find that it is beneficial to constrain the optimization to $r_K = 2$ in order to achieve an accurate unit-rank approximation. The same effect is also observed in the single qubit example when taking a detailed look at the number of required initializations (see section IV E), yet less pronounced. In the bottom row of Figure 3 we show the recovery rates for random unitary models, with the reconstruction now using a fixed Kraus rank of $r_K = 2$. Note that there are three sources of randomness present in the data, first the Haar-random unitary gates, then the random drawing of gate sequences and finally the random initialization of the algorithm. Each shade of green corresponds to one random gate set, and the recovery rate tells us how many of the 10 random sequence sets lead to a successful reconstruction, given a budget of 33 initializations. We find that the random single qubit gate sets all have similar recovery rates, with a successful reconstruction possible from $n_{\text{seq}} = 20$ to $n_{\text{seq}} = 30$ sequences, and a high rate of recovery at $n_{\text{seq}} = 100$ sequences. In the two qubit case a different picture emerges, where two of the random gate sets show a high recovery rate at $n_{\text{seq}} = 200$ sequences (akin to the XYICNOT-model), while the least favorable random gate set was only recoverable at $n_{\text{seq}} = 500$ sequences. This shows that random gate sets for two qubits can have very different conditioning.

Sequence number comparison to pyGSTi

Comparing the number of random sequences needed for mGST and the number of sequences for pyGSTi is not straightforward. The standard pyGSTi data-processing pipelines crucially relies on specific, fixed sequence construction. For this reason pyGSTi cannot be applied to the type of data that we use here. We can however compare the number of random sequences with the number of deterministic sequences that the standard implementation of pyGSTi uses. For the single qubit XYI-model, the minimal number of sequences given in the pyGSTi

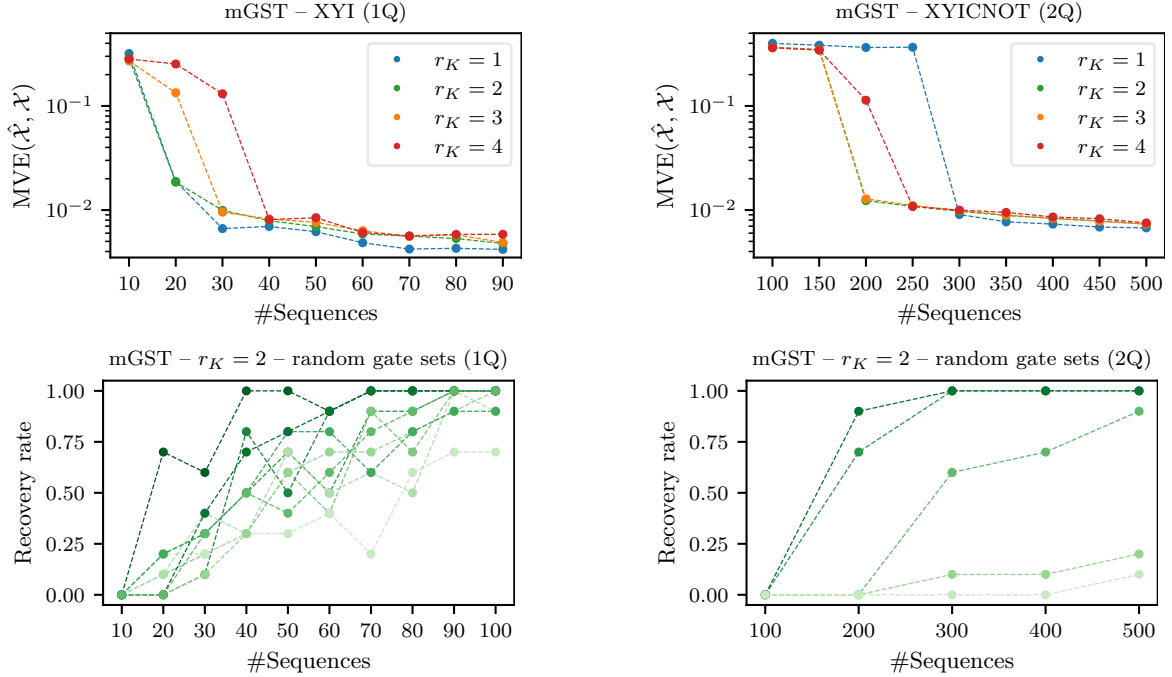


Figure 3. Median error of mGST run on the least squares cost function, plotted over the number of sequences for a single qubit model (**top left**) and a two-qubit model (**top right**). Each data point is the median over the results from 10 different random sequences. The measurement data for the XYI - and the XYICNOT gate set is taken from a noisy version with depolarizing noise of strength $p = 0.001$ on each gate, depolarizing noise with strength $p = 0.01$ before measurement, as well as independent random unitary rotations $e^{i\gamma H}$ with $\gamma = 0.001$ and $H \sim \text{GUE}$ on each gate.

On the **bottom left** the recovery rates for the reconstruction of different models of 3 Haar random unitaries are shown. For each gate set the average over 10 draws of random sequences is shown. A gate set is classified as recovered if the MVE falls below 0.03. The **bottom right** depicts the recovery rate for random two qubit gate sets of the form $\mathcal{G} = \{\mathbb{1}_4, \mathbb{1}_2 \otimes U_1, \mathbb{1}_2 \otimes U_2, U_1 \otimes \mathbb{1}_2, U_2 \otimes \mathbb{1}_2, U_{12}\}$ where U_1 and U_2 are Haar random single qubit unitaries and U_{12} is a Haar random two qubit unitary. Additionally, each single qubit and two qubit gate contains depolarizing noise of strength $p = 0.001$ and depolarizing noise of strength 0.01 is applied before measurement. The recovery rate is averaged over 10 random sequence draws. For all gate sets the sequences are drawn uniformly at random with sequence length $\ell = 7$ and $m = 1000$ samples per sequences. The maximum number of initializations are 80, 33, 17 and 10 for Kraus ranks $1, \dots, 4$ respectively. They are chosen such that the maximal computation time is equal among different ranks.

implementation is $n_{\text{seq}} = 92$. This is significantly larger than the number of random sequences at which the phase transition of mGST in Figure 3 appears. However, the $n_{\text{seq}} = 92$ sequences are overcomplete by design, and we find that pyGSTi can also reconstruct the XYI model with $n_{\text{seq}} = 48$ sequences. Yet we find the same sequence design not to be successful for three Haar-random single qubit gates, indicating that the choice of sequences is well-tailored to the XYI-model. The reduction in sequences becomes more pronounced for the two-qubit gate set studied in the top right of Figure 3. For this gate set, the minimal number of gate sequences that pyGSTi uses is $n_{\text{seq}} = 907$, which is significantly larger than what mGST needs.

D. Characterizing unitary errors using prior knowledge

In the previous section we demonstrated compressive gate set tomography for one- and two-qubit gate sets using agnostic random initializations. A major obstacle in going beyond reconstructing entire two-qubits gate sets even compressively on desktop hardware is that besides run-time and storage also the number of required random initializations until proper convergence grows in principle with the number of qubits. This is due to longer sequences being required for tomographic completeness, leading to a higher order polynomial in the cost function. This situation can be remedied by using prior knowledge, such as the target gate set, for the initialization. In this case, a gate set in the vicinity of the initial point will be found which is in better agreement with the data. In a conceivable experimental scenario the gates are more or less known due to the physical setup and previous benchmarking rounds, but further calibration requires infor-

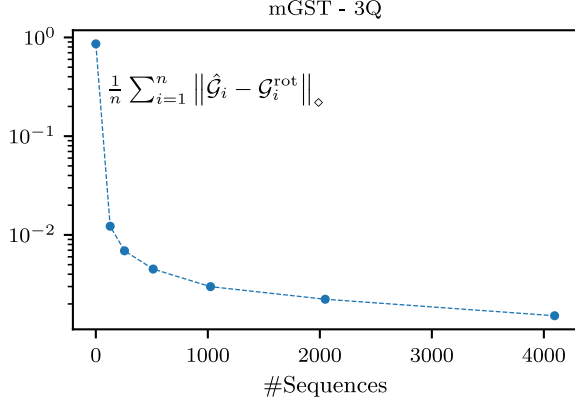


Figure 4. Average diamond distance between 3-qubit rotated target gates $\mathcal{G}_i^{\text{rot}}$ and their unitary ($r_K = 1$) mGST estimators $\hat{\mathcal{G}}_i$ as a function of the number of sequences. The 0-sequence data marks the average diamond distance between initialization $\mathcal{G}_i$ and $\mathcal{G}_i^{\text{rot}}$. The rotated gates $\mathcal{G}_i^{\text{rot}}$ are related to their counterparts $\mathcal{G}_i$ by independent random (global) overrotations on each gate, given by $\exp(i\gamma H)$ with $H \sim \text{GUE}$ and $\gamma = 0.05$, leading to $\mathbb{E} \|\gamma H\|_\infty \leq 2\sqrt{2}/5$. The data is simulated from the gates $\mathcal{G}_i^{\text{rot}}$ with depolarizing noise of strength $p = 0.01$ on each gate and the measurements are taken from random sequences of length $\ell = 7$, with $m = 10^5$ samples per sequence.

mation about present coherent errors. The use of prior knowledge can be seen as a situational tool to further reduce runtime when applicable, but for general purpose verification and characterization no initial point is to be assumed for compressive GST with the mGST algorithm. To showcase the characterization of unitary errors using prior knowledge, we take a three-qubit gate set that is the direct generalization of the previous two-qubit XY-ICNOT model, by adding the local X- and Y-rotations as individual gates to the 3rd qubit and adding a CNOT between qubits 2 and 3. We then apply a global random rotation to each gate individually, as well as depolarizing noise on each gate. From random sequences of fixed sequence length we can then, in theory, fit the noisy model perfectly via an $r_K = 1$ approximation, as the depolarizing channels commute with the unitary gates and can be pulled into initial state or measurement. In Figure 4 we see that mGST is indeed able to precisely reconstruct the rotated gates, as shown by the average diamond norm error. We chose a comparatively high number of 10^5 samples per sequence to showcase that high accuracy can indeed be realized using this method: for instance, only 256 sequences are enough to achieve an average diamond norm distance of around 0.007 between the reconstructed unitary gates and the true unitary gates, which include overrotations. The fact that these overrotations were modelled as being global on all 3 qubits suggests that we can efficiently characterize unitary crosstalk as well, by capturing the effect of single and two qubit gate on their neighbours within a three qubit region.

E. Implementation details and calibration

We now provide more details on the simulations, the criteria for successful recovery and the required number of initializations. To simulate measurements on a gate sequence $\mathbf{i}$, we first compute the outcome probabilities $p_{j|\mathbf{i}}$ from Eq. (16) of the POVM elements according to the model gate set in question. Afterwards we draw m samples from the multinomial distribution $M(m, (p_{1|\mathbf{i}}, \dots, p_{n_E|\mathbf{i}}))$, where $\sum_j p_{j|\mathbf{i}} = 1$. Let $k_{j|\mathbf{i}}$ be the number of times outcome j occurred for sequence $\mathbf{i}$. Then Algorithm 2 optimizes the objective function (20) on the estimated probabilities $y_{j|\mathbf{i}} = k_{j|\mathbf{i}}/m$.

For the single qubit examples the batch size $\kappa = 50$ was chosen, while for the two qubit example we use $\kappa = 120$. The choice of batch size determines the number of values summed over in Eq. (20). Therefore, the computation time of the objective function and its derivatives scales linearly in κ , making a small batch size favorable. However, it cannot be set too small, otherwise the update directions become highly erratic, and no convergence is reached. A general rule of thumb is to set the batch size close to the number of free parameters in the model. Another hyperparameter is the damping value λ for the saddle-free Newton method described in Algorithm 1. We find that a fixed value of $\lambda = 10^{-3}$ leads to the best results across the models tested.

Judging whether mGST recovers a gate set by looking at the attained objective function value can only be done if the set of measured sequences is informationally complete. Then there is a unique (up to gauge) global minimum in the least squares minimization problem and the minimum corresponds to the true gate set in the limit of infinitely many samples per sequence.

In Figure 5 we take a look at the correlation between the final least square objective function value $\mathcal{L}(\hat{\mathcal{X}}, \mathbf{y})$ and the mean variation error $\text{MVE}(\hat{\mathcal{X}}, \mathcal{X})$. We see that for a low number of sequences (10-20), a low objective function value does not imply a low MVE, yet for higher numbers of sequences, an objective function value below 10^{-3} implies an MVE around 10^{-2} . For sufficiently many sequences, the gray line indicating our success criterion clearly separates two clusters of points, meaning that no intermediate quality fits are found in our model space. In this sequence regime either the algorithm converges to a fit as good as the sample count allows, or it does not converge at all. Therefore, restarting the algorithm when an initialization turns out to be bad yields practically optimal results. A thorough analysis of the probability of obtaining an informationally complete set of random sequences is left for future work.

To give an intuition on how many initializations are required for mGST to converge, we can take a look at data from a modified XYI-model with gates $\{\mathbb{1}, e^{i\frac{\pi}{2}\sigma_y}, e^{i\frac{\pi}{2}\sigma_x}\}$, simulating more difficult gate set conditioning. Figure 6 shows histograms for the number of reinitializations needed for convergence. The data combines the results for α between $\pi/18$ and $\pi/2$, with depolarizing noise of

strength $p = 0.001$ on each gate and $p = 0.01$ on the initial state, as well as a maximum of 100 reinitializations.

We observe that in 48.4% of all cases convergence is reached on the first attempt, and in 90.4% of cases 4 or fewer reinitializations are required. The histogram indicates that the chance of needing multiple reinitializations rapidly decays and that rank-1 optimization is more sensitive to bad initializations compared to rank-4 optimization.

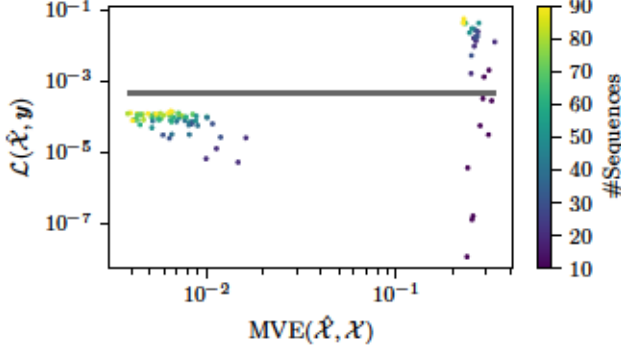


Figure 5. Least square objective function vs. MVE at the end of the mGST-optimization, with the gray bar indicating the range of stopping values below which a run is considered successful. The plot illustrates that results with a large MVE also have a large objective function if enough sequences are measured. The experiment is the same as in the top left of Figure 3, for $r_K = 4$. The color of each point indicates the number of measured sequences on which mGST was run.

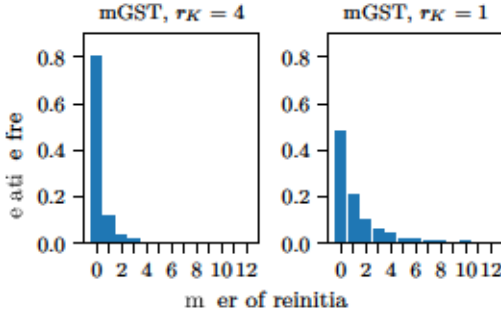


Figure 6. Relative frequencies of reinitialization counts for mGST run on modified XYI models (see main text for details).

F. Runtime and scaling

In order to assess the runtime scaling of Algorithm 2 in the problem parameters we identify the two most time-consuming steps as the computation of the second derivative terms in H and the diagonalization of H for the SFN update on the gates.

Recall that r_K denotes the Kraus rank of the gate estimators, n the number of gates in the gate set, ℓ the number of gates per sequence and κ the number of sequences

per batch. The computation of the second derivative terms scales as $O(\kappa \ell^3 d^6)$, while the eigendecomposition of H scales as $O((2nr_K d^2)^3)$. A computationally less expensive variant is to not optimize over all variables of the full gate tensor at once, but over the individual gates one after another. The complexity of the eigendecompositions then reduces to $O(n(2r_K d^2)^3)$, which is beneficial for large gate sets. However, this approach also leads to slower convergence, and we choose to optimize over the full tensor by default.

Table I contains runtimes for the system sizes studied in our numerical experiments. We find that reconstruction of single qubit gate sets can be achieved within seconds and low rank 2-qubit gate sets can be reconstructed within minutes on a single modern 32-core CPU. The runtimes for the gradient descent method on the 2-qubit example were very fast up to high ranks, however we generally find that the default Newton method is more reliable across different gate set and sequence scenarios. If a good initialization is known, the runtime is reduced drastically. For instance the 3-qubit reconstruction done with prior knowledge about the gates in Section IV D took 5 hours and 30 minutes to complete.

mGST: Default (Newton)				
	$r_K = 1$	$r_K = 4$	$r_K = 8$	$r_K = 16$
1Q	4"	14"	/	/
2Q	8'8"	39'50"	2h26'10"	9h1'43"
3Q	3d2h	/	/	/
mGST: Gradient descent				
	$r_K = 1$	$r_K = 4$	$r_K = 8$	$r_K = 16$
1Q	26"	17"	/	/
2Q	5'40"	6'0"	5'51"	4'47"
3Q	No conv.	/	/	/
pyGSTi				
1Q	5"			
2Q	1h6'48"			
3Q	No result after 4d14h			

Table I. Average runtimes for 1, 2 and 3 qubits with selected Kraus ranks on a modern 32 core CPU. The average runtime was calculated as the average time until the first successful reconstruction with a random initialization was achieved. The one qubit gate set is the same as in Figure 2, the two qubit gate set the same as in Figure 3 and the three qubit gate set is the same as in Figure 4 but without using any prior knowledge. For the gradient descent method in the three qubit scenario, no convergence was achieved within the maximum iteration limit. The average runtimes for pyGSTi were determined on the same models and with the same sequences.

V. NOISE-MITIGATION FOR SHADOW ESTIMATION WITH COMPRESSIVE GST

Gate set estimates provide a detailed picture of the imperfections that can inform prioritization and further

experimental efforts [15, 19, 21–23]. Coherent error estimates can often be directly corrected for by adjusting or optimizing the control. We expect that the more economically accessible compressive estimates from mGST can be used in place of traditional GST estimates in the above applications, in particular when complemented with RB estimates of incoherent noise effects. However, it is beyond the scope of this work to demonstrate mGST in a full engineering cycle of a quantum computing device.

To still showcase the value of GST estimates, we sketch another novel application that can be demonstrated without simulating a whole engineering cycle. Generally speaking, noise characterizations can be used to mitigate noise-induced biases in other quantum characterization protocols by adapting the classical post-processing [98, 99]. Given a device that can repeatedly prepare a quantum state, a fundamental characterization task is to estimate the expectation values of observables from measurements. In particular, an informationally complete measurement allows one to estimate arbitrary observables from the same data in the post-processing. Such an informationally complete measurement can be implemented on a quantum computing device by measuring in sufficiently many random bases, the prototypical example being measurements in random Pauli bases. Ref. [46] showed how to derive optimal guarantees with exponential confidence for estimating multiple observables simultaneously using a median-of-means estimator and explicit bounds on the variance for random basis measurement that constitute unitary 3-designs. They introduced the term ‘classical shadow’ to refer to the elements of a dual frame of the informationally complete POVM corresponding to observed samples, see Appendix E for a brief summary. Importantly, one can often arrive at high precision estimates of observables long before one has measured all the informationally complete bases in multi qubit systems.

In practice however, implementing a random bases measurement, say, by applying a unitary rotation followed by a computational bases measurement will suffer from noise from the gates and read-out. This has motivated the development of robust variants of shadow estimation that either make use of simple depolarizing noise-models of known strength [100] or perform a separate RB-style experiment that estimates the depolarizing noise-strength induced by a gate-independent channel acting between the rotation and the measurement [101]. Using GST estimates provides a complimentary, flexible approach to mitigate even highly gate-dependent noise with finite correlations in shadow estimation.

We demonstrate how GST estimates on 2-qubit pairs can be used to calculate noise-robust classical shadow estimators in post-processing. Our robust estimation scheme consists of two distinct stages each consisting of multiple steps: (I) *calibration stage*: (i) the local channels implementing each combination of two local gates are reconstructed with mGST; (ii) the gauge of these gate estimates is matched to the gauge in which the ideal gates

and the observables are given. For this step we use the gauge optimization provided by the pyGSTi package [45]. (iii) From the gauge-optimized channel estimates, we numerically calculate the *effective measurement map* when implementing random Pauli measurements with the characterized noisy gates-set.

(II) After calibration, the second stage is a *shadow estimation* protocol consists of two separate phases: (i) the *data acquisition* by repeatedly measuring the unknown state of the quantum device in a randomly selected Pauli basis; (ii) the *classical post-processing* where estimators of the observables are calculated using the data. We use the inverse of the effective measurement map from the calibration stage to calculate the empirical estimators. We give a more detailed description of the individual steps of the procedure in Appendix E.

Using an empirically estimated effective measurement map instead of the ideal theoretical result is the essential modification compared to standard shadow estimation. In this way, we also ‘invert’ the effect of the noise on our estimator. The right column of Figure 7 shows an effective measurement map implemented with imperfect gates.

As a proof of concept, we chose the following simple but practically relevant setup: Random local Pauli basis measurements are implemented by native measurements in the computational basis after rotating with a Hadamard gate H (if the Pauli-X basis is to be measured) or a phase gate S followed by a Hadamard gate (for measurement in the Pauli-Y basis). Since throughout the protocol the S -gate only turns up before application of the Hadamard gate, we treat the sequence HS as a single gate.

We assume that the dominant noise associated with the single qubit rotations of each local gate in the experiment stays confined to two neighboring qubits. This assumption makes both the gate set estimation and the post-processing of the shadow estimation highly scalable.

Figure 7 shows the results of our scheme in simulations of the energy estimation of a Heisenberg Hamiltonian on a 10-qubit system. We observe that using the estimated effective measurement map instead of the ideal theoretical one significantly reduces the relative error $|(\hat{E} - E)/E|$ between the estimated energy $\hat{E}$ and the true energy E of a given state. There are two contributions to the relative error in a shadow estimation protocol: First, the statistical fluctuation from the randomness of both the Pauli basis selection and the single shot measurements. Second, the systematic bias introduced in the post-processing due to imperfect implementations of the measurements. The histograms in Figure 7 show the infinite measurement limit of the relative error and thus directly reflect the bias. We observe that for a fixed noise model, the magnitude of the bias depends heavily on the selected initial state, with relative errors being distributed over two orders of magnitude. When comparing the most likely errors between standard shadow estimation and GST-mitigated shadow estimation we find that using GST data leads to a reduction in relative error by

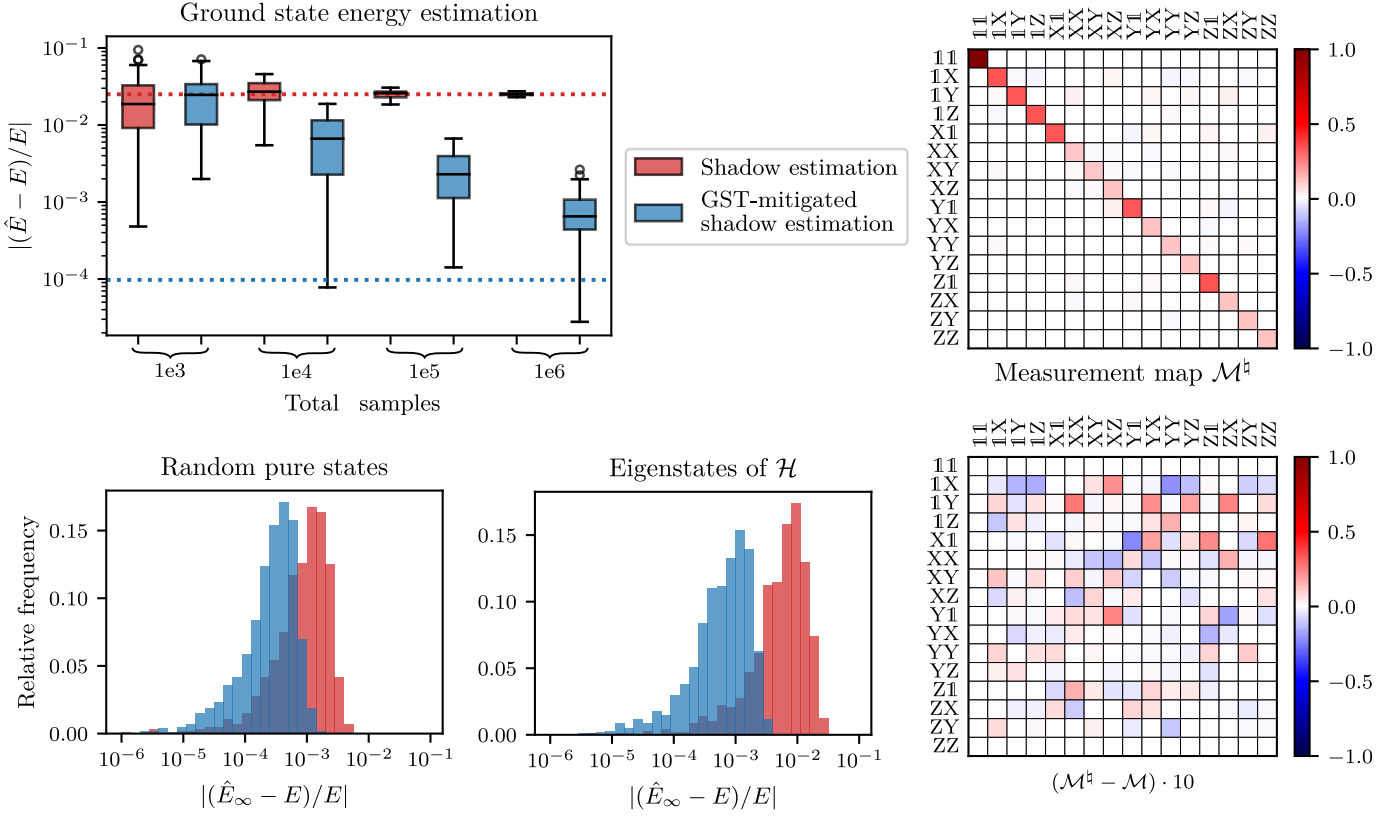


Figure 7. Energy estimation for the 10 qubit Hamiltonian $\mathcal{H} = \frac{1}{2} \sum_{j=1}^{10} (\sigma_x^j \sigma_x^{j+1} + \sigma_y^j \sigma_y^{j+1} + \sigma_z^j \sigma_z^{j+1} - \sigma_z^j)$ with periodic boundary conditions. The **top left** shows the sample dependence of the relative accuracy $|(\hat{E} - E)/E|$ for estimating the ground state energy. The colored blocks extend from the 1st quartile to the 3rd quartile around the median (black line) of the data (50 repetitions per sample value). The whiskers extend from the 5th to the 95th percentile and the dashed lines indicate the infinite sample expectation values. On the **bottom left** two histograms are shown that compare the theoretical infinite sample energy estimates $\hat{E}_\infty$ (biases) for 1000 random pure states and for all 1024 eigenstates of the Hamiltonian, respectively. All simulations were done with noisy Clifford gates, whose average gate fidelity to their ideal counterparts is at 0.99 ± 10^{-3} . On the **top right** the Pauli transfer matrix of a two-qubit effective measurement map $\mathcal{M}^\sharp$ under this noise model is shown. The **bottom right** plot displays the difference between $\mathcal{M}^\sharp$ and its noise-free counterpart $\mathcal{M}$. Gate estimates for GST-mitigated shadow estimation were produced with mGST ($r_K = 2$), using 400 random sequences equally distributed among sequences lengths $\{6, 7, 8, 9\}$ with 10^4 samples per sequence. The noise in the simulations is given by two-qubit random unitary noise $e^{i\gamma K}$ with $K \sim \text{GUE}$. The error parameter is $\gamma = 0.14$ on H and on HS , leading to the aforementioned average gate fidelities of ~ 0.99 . For the bottom left histogram, random pure state were generated as $U|0\rangle$, with U drawn according to the Haar measure.

half an order of magnitude for random pure states and an order of magnitude for eigenstates of the Hamiltonian. The simulation of the protocol in the top left of Figure 7 includes statistical fluctuations and showcases how estimates spread for different sample counts when the ground state energy is estimated. We find that from 10^4 samples on, the GST-mitigated protocol yields significantly more accurate estimate.

VI. CONCLUSION AND OUTLOOK

We have revisited the data processing task of GST from a compressed sensing perspective regarding it as a highly

structured and constrained tensor completion problem. In this formulation, we can naturally require the reconstructed gate set to be physical and, moreover, of low rank. Compressive gate set tomography, thus, aims at extracting considerably fewer parameters of the gate set. At the same time we have argued that the low-rank approximation to the implementation of a gate set contains the most valuable information about experimental imperfections for the practitioner.

The set of Kraus-operators of a low-rank gate can be regarded as isometries that make up the complex Stiefel manifold. This observation has motivated the solution of the compressive GST data processing problem via geometrical optimization on the respective product mani-

folds. We have devised the optimization algorithm **mGST** that performs an adapted saddle-free Newton method on the manifold. To this end, we have derived the Riemannian Newton equation, Hessian equation and geodesic curves.

In numerical experiments we have studied the performance of the **mGST** algorithm. We have compared it to **pyGSTi**, the state-of-the-art approach to the GST data processing problem, in settings where both algorithms can be applied and using full rank **mGST** estimates. We have found that in these settings **mGST** matches the performance of **pyGSTi**, while offering a larger model space, more flexibility in the sequence design and allowing for low rank assumptions. Moreover, we have demonstrated numerically that making use of the low-rank constraints significantly reduces the required number of measured sequences and the run-time of the reconstruction algorithm for a standard single and two qubit model. Importantly, we have found that we can successfully reconstruct generic unitary channels and depolarizing noise of one- and two-qubit gate sets from *random* gate sequences. This reduces the demands of GST both for experiments and classical post-processing: the data that compressive GST requires is virtually identical with the experimental data produced by randomized benchmarking experiments. The classical post-processing of **mGST** for a low-rank reconstruction of two qubit gate sets takes only minutes even on desktop hardware, compared to over an hour with **pyGSTi**. We expect that this speedup and the low number of sequences required can lift 2-qubit GST from being a protocol that is unpractical in many situations to one that is routinely applied, thus enabling it to be used in the engineering cycle for the design and calibration of gate sets.

Furthermore, compressive GST makes it feasible to perform self-consistent tomography on 3-qubit systems. Making use of often available prior knowledge about an initialization can further reduce the computing time. We demonstrated this by performing tomography of unitary errors for 3-qubit gate sets, using only a small number of random gate sequences, and with the post-processing still running on desktop hardware in a few hours. We expect that even going slightly beyond three qubits is feasible by simply using more computing power. We leave it to future work to further tweak the numerical implementation in order to improve the scalability of the classical post-processing. We also expect that progressively longer sequences can be added at the end of our optimization method, much in the same fashion as in **pyGSTi**, in order to further improve reconstruction accuracy of gate set estimates.

To demonstrate the use of compressive GST for error mitigation, we have introduced one novel application where low-rank **mGST** reconstructions are used to alleviate the effect of coherent errors in classical shadow estimation. The protocol uses a set of gates to implement basis changes before the measurement. We have demonstrated that with tomographic information

on these gates through two-qubit compressive GST, more accurate ground state energy estimates are obtained in practically relevant regimes. This constitutes just one example where the full information of a low rank GST estimate is used to correct errors, and we expect that the reduced runtime requirements of **mGST** enable frequent use of GST for error diagnosis and mitigation.

Finally, besides making GST more applicable and flexible in practice, our reformulation is motivated by bringing it closer to theoretical recovery guarantees quantifying a required and sufficient number of random sequences for accurate reconstruction. Regarding the data processing of GST as a translation-invariant matrix-product-state/tensor-train completion problem makes it more amenable to prove techniques from compressed sensing. For example, establishing local convergence guarantees for **mGST** would allow one to quantify the assumptions on the experimental implementation that justify certain initialization of the algorithm. We hope that our work can serve as a foundation and inspiration in the quest of establishing mathematically rigorous guarantees for GST.

ACKNOWLEDGMENTS

We thank Lennart Bittel for helpful discussions and hints regarding optimization methods and their implementation and Markus Heinrich for a discussion on inverses of positive maps on operator spaces. We are particularly thankful to Kenneth Rudinger and Robin Blume-Kohout for valuable feedback on a previous version of the manuscript and help with finding an error in our code. The work of RB and MK has been funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) within the Emmy Noether program (grant number 441423094) and by the German Federal Ministry of Education and Research (BMBF) within the funding program “quantum technologies – from basic research to market” via the joint project MIQRO (grant number 13N15522). IR acknowledges funding from the BMBF (DAQC) and the Einstein Foundation (Einstein Research Unit).

VII. APPENDIX

In this appendix, we provide the mathematical details required for the saddle-free Newton method within the Riemannian optimization framework, see Appendices A, B, and C. Moreover, we compare the dependence of the mGST MVE on the choice of objective function (mean squared error vs. maximum likelihood) in Appendix D.

A. Geodesics on the Stiefel manifold

Edelman, Arias, and Smith [48] derived the geodesic on the real Stiefel manifold by solving the respective geodesic equation. We now show that the simple generalization given in Eq. (27) is indeed the correct geodesic in the complex case. For a curve $K_t \equiv K(t)$ the general geodesic equation is [89, Chapter 5.4, Proposition 5.3.2]

$$P_{T(K_t)} \left(\ddot{K}_t + C_{K_t}(\dot{K}_t, \dot{K}_t) \right) = 0, \quad (32)$$

where the Christoffel symbol C_{K_t} depends on the chosen metric. Here, we use the canonical metric

$$\langle \Delta_1, \Delta_2 \rangle_K = \text{Re} \left\{ \text{Tr}(\Delta_1^\dagger \Gamma \Delta_2) \right\} =: g(\Delta_1, \Delta_2) \quad (33)$$

with $\Gamma = \mathbb{1} - \frac{1}{2} K_t K_t^\dagger$. Using the Einstein summation convention, the Christoffel symbol at K can be computed as

$$(C_{K_t}^k)_{ij} = \frac{1}{2} g_{kl}^{-1} \left(\frac{\partial g_{lj}}{\partial K_{ti}} + \frac{\partial g_{li}}{\partial K_{tj}} - \frac{\partial g_{ij}}{\partial K_{tl}} + \text{c.c.} \right), \quad (34)$$

where $C_{K_t}^k$ is the k -th component of the Christoffel symbol at K_t with respect to a basis $\{E_k, E_k^*\}_{k \in [Dd]}$ on the ambient space $\mathbb{C}^{D \times d}$.

Lemma 1. *The geodesic equation on the complex Stiefel manifold $\text{St}(D, d)$ equipped with the canonical metric for the curve $K_t : \mathbb{R} \rightarrow \text{St}(D, d)$ is given by*

$$P_{T(K_t)} \left(\ddot{K}_t + \dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t \right) = 0. \quad (35)$$

Proof. By noting that $\Gamma^{-1} = \mathbb{1} + K_t K_t^\dagger$ we can determine the function $g^{-1}(\Delta_1, \Delta_2)$ via the condition $g(g^{-1}(\Delta_1, \cdot), \Delta_2) = \text{Tr}[\Delta_1^\dagger \Delta_2 + \Delta_2 \Delta_1^\dagger]$, meaning the inverse g^{-1} would recover the standard symmetric inner product on $T_K \text{St}(D, d)$. One can quickly verify that $g^{-1}(\Delta_1, \cdot) = \Gamma^{-1} \Delta_1$ satisfies this condition.

We determine the derivatives of g needed for the Christoffel symbol by explicitly writing out g as

$$g(\Delta_1, \Delta_2) = \text{Tr} \left[\Delta_1^\dagger \left(\mathbb{1} - \frac{1}{2} K_t K_t^\dagger \right) \Delta_2 + \Delta_2^\dagger \left(\mathbb{1} - \frac{1}{2} K_t K_t^\dagger \right) \Delta_1 \right], \quad (36)$$

from where we can find the derivatives by K and K^* as

$$\frac{\partial g_{ij}}{\partial K_{tl}}(E_i, E_j) = -\frac{1}{2} \text{Tr} \left[E_i^\dagger E_l K^\dagger E_j + E_j^\dagger E_l K^\dagger E_i \right], \quad (37)$$

$$\frac{\partial g}{\partial K_{tl}^*}(E_i, E_j) = \left(\frac{\partial g}{\partial K_{tl}}(E_i, E_j) \right)^* \quad (38)$$

$$= -\frac{1}{2} \text{Tr} \left[E_i^\dagger K E_l^\dagger E_j + E_j^\dagger K E_l^\dagger E_i \right]. \quad (39)$$

With these derivatives we can calculate

$$C_{K_t}^k(\dot{K}_t, \dot{K}_t) = (C_{K_t}^k)_{ij} (\dot{K}_t)_i (\dot{K}_t)_j \quad (40)$$

$$= \frac{1}{2} g_{kl}^{-1} \left(\frac{\partial g_{lj}}{\partial K_{ti}} + \frac{\partial g_{li}}{\partial K_{tj}} - \frac{\partial g_{ij}}{\partial K_{tl}} + \text{c.c.} \right) \dot{K}_{ti} \dot{K}_{tj} \quad (41)$$

$$= \frac{1}{2} \left(\frac{\partial g}{\partial K_{ti}} \left(g^{-1}(E_k, \cdot), \dot{K}_t \right) \dot{K}_{ti} + \frac{\partial g}{\partial K_{tj}} \left(g^{-1}(E_k, \cdot), \dot{K}_t \right) \dot{K}_{tj} - \frac{\partial g}{\partial K_{tl}} \left(\dot{K}_t, \dot{K}_t \right) (g^{-1}(E_k, \cdot))_l + \text{c.c.} \right) \quad (42)$$

$$= -\frac{1}{2} \text{Re} \left\{ \text{Tr} \left[2E_k^\dagger \Gamma^{-1} \dot{K}_t K_t^\dagger \dot{K}_t + 2\dot{K}_t^\dagger \dot{K}_t K_t^\dagger \Gamma^{-1} E_k - 2\dot{K}_t^\dagger \Gamma^{-1} E_k K_t^\dagger \dot{K}_t \right] \right\} \quad (43)$$

$$= -\text{Re} \left\{ \text{Tr} \left[\dot{K}_t K_t \dot{K}_t^\dagger \Gamma^{-1} E_k + \dot{K}_t^\dagger \dot{K}_t K_t^\dagger \Gamma^{-1} E_k - K_t^\dagger \dot{K}_t \dot{K}_t^\dagger \Gamma^{-1} E_k \right] \right\} \quad (44)$$

$$= \left\langle \left(K_t^\dagger \dot{K}_t K_t^\dagger - \dot{K}_t^\dagger \dot{K}_t K_t^\dagger - \dot{K}_t^\dagger K_t \dot{K}_t^\dagger \right)^\dagger, E_k \right\rangle \quad (45)$$

$$= \left\langle \dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t, E_k \right\rangle, \quad (46)$$

where we have used that $(\Gamma^{-1})^\dagger = \Gamma^{-1}$ and $\text{Re Tr}[X] = \text{Re Tr}[X^\dagger]$. We now first write out the geodesic equation (32) on the ambient space,

$$\langle \ddot{K}_t, E_k \rangle + C_{K_t}^k(\dot{K}_t, \dot{K}_t) = \langle \ddot{K}_t, E_k \rangle + \left\langle \dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t, E_k \right\rangle = 0 \quad \forall E_k, \quad (47)$$

$$(48)$$

which is equivalent to

$$\ddot{K}_t + \dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t = 0. \quad (49)$$

□

To arrive at the geodesic equation (27), it remains to project the above equation onto the tangent space. Indeed, with the explicit form of the geodesic equation from Lemma 1 we can show that the immediate generalization from the geodesic in the real case [48] gives a valid geodesic for the complex case.

Lemma 2. *The curve given by*

$$K_t = \begin{pmatrix} K & Q \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} \mathbb{1} \\ 0 \end{pmatrix} \quad (50)$$

is a geodesic on $\text{St}(D, d)$, determined through the initial conditions $K_{t=0} = K$ and $\dot{K}_{t=0} = \Delta$, with Q, R given by the QR decomposition of $(\mathbb{1} - KK^\dagger)\Delta$ and $A = K^\dagger \Delta$.

Proof. We recall that the ambient space splits into the tangent space and its orthogonal complement, the normal space. Therefore, the condition that the projection of the left-hand side onto the tangent space in Eq. (35) vanishes is equivalent to demanding that it lies solely in the normal space. If it is in the normal space, $K_t^\dagger$ applied from the left will yield a Hermitian matrix. We will now show that this is indeed the case. For that we first need to determine the first and second derivatives of K_t :

$$\dot{K}_t = \begin{pmatrix} K & Q \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \begin{pmatrix} \mathbb{1} \\ 0 \end{pmatrix} \quad (51)$$

$$= \underbrace{K_t A}_{\dot{K}_1} + \underbrace{\begin{pmatrix} K & Q \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} R}_{\dot{K}_2}, \quad (52)$$

$$\ddot{K}_t = \begin{pmatrix} K & Q \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix}^2 \begin{pmatrix} \mathbb{1} \\ 0 \end{pmatrix} \quad (53)$$

$$= \underbrace{K(t)(A^2 - R^\dagger R)}_{\ddot{K}_1} + \underbrace{\begin{pmatrix} K & Q \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} R A}_{\ddot{K}_2}. \quad (54)$$

We immediately see that $K_t^\dagger \ddot{K}_1 = A^2 - R^\dagger R$, which is Hermitian, as A is skew Hermitian. We will now show that $K_t^\dagger \ddot{K}_2 = 0$ starting with

$$K_t^\dagger \ddot{K}_2 = (\mathbb{1} \ 0) \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} K^\dagger \\ Q^\dagger \end{pmatrix} (K \ Q) \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} RA \quad (55)$$

$$= (\mathbb{1} \ 0) \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} \mathbb{1} & K^\dagger Q \\ Q^\dagger K & \mathbb{1} \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} RA \quad (56)$$

$$= (\mathbb{1} \ 0) \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \left(\begin{pmatrix} \mathbb{1} & 0 \\ 0 & \mathbb{1} \end{pmatrix} + \begin{pmatrix} 0 & K^\dagger Q \\ Q^\dagger K & 0 \end{pmatrix} \right) \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} RA \quad (57)$$

$$= (\mathbb{1} \ 0) \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 & K^\dagger Q \\ Q^\dagger K & 0 \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} RA. \quad (58)$$

To simplify the last expression, we set

$$\begin{pmatrix} U_{00} & U_{01} \\ U_{10} & U_{11} \end{pmatrix} = \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \quad (59)$$

and obtain $K_t^\dagger \ddot{K}_2 = (U_{00}^\dagger K^\dagger Q U_{11} + U_{10}^\dagger Q^\dagger K U_{01}) RA$. From the series representation of the matrix exponential we gather that $U_{11} = \mathbb{1} + R \cdot X$ for some matrix X . Moreover $U_{10} = R \tilde{X}$ and $U_{10}^\dagger = \tilde{X}^\dagger R^\dagger$ for some $\tilde{X}$, leading to

$$K_t^\dagger \ddot{K}_2 = (U_{00}^\dagger K^\dagger Q + U_{00}^\dagger K^\dagger Q R X + \tilde{X}^\dagger R^\dagger Q^\dagger K U_{01}) RA = 0, \quad (60)$$

since $K^\dagger Q R = K^\dagger (\mathbb{1} - K K^\dagger) \Delta = 0$.

This shows that the $\ddot{K}_t$ lies in the normal space, leaving us with the terms in the geodesic equation (35) that depend only on $\dot{K}_t$:

$$\begin{aligned} K_t^\dagger (\dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t) &= K_t^\dagger \dot{K}_t \dot{K}_t^\dagger K_t - \dot{K}_t^\dagger \dot{K}_t - (K_t^\dagger \dot{K}_t)^2 \\ &= (A + K_t^\dagger \dot{K}_2)(A + K_t^\dagger \dot{K}_2)^\dagger - (A^\dagger A + A^\dagger K_t^\dagger \dot{K}_2 + \dot{K}_2^\dagger K_t A + \dot{K}_2^\dagger \dot{K}_2) \\ &\quad - (A + K_t^\dagger \dot{K}_2)^2 \\ &= AA^\dagger - A^\dagger A - \dot{K}_2^\dagger \dot{K}_2 - A^2. \end{aligned}$$

The last line follows from $\dot{K}_2 A = \dot{K}_2$ and our previous observation that $K_t^\dagger \ddot{K}_2 = 0$, which implies that $K_t^\dagger \dot{K}_2 = 0$ as well. The remaining term $\dot{K}_2^\dagger \dot{K}_2$ can be computed similarly to $K_t^\dagger \ddot{K}_2$ and we obtain

$$\dot{K}_2^\dagger \dot{K}_2 = R^\dagger \begin{pmatrix} 0 & \mathbb{1} \end{pmatrix} \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} \mathbb{1} & K^\dagger Q \\ Q^\dagger K & \mathbb{1} \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} R \quad (61)$$

$$= R^\dagger \begin{pmatrix} 0 & \mathbb{1} \end{pmatrix} \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \left(\begin{pmatrix} \mathbb{1} & 0 \\ 0 & \mathbb{1} \end{pmatrix} + \begin{pmatrix} 0 & K^\dagger Q \\ Q^\dagger K & 0 \end{pmatrix} \right) \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} R \quad (62)$$

$$= R^\dagger R + R^\dagger \begin{pmatrix} 0 & \mathbb{1} \end{pmatrix} \exp \left(-t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 & K^\dagger Q \\ Q^\dagger K & 0 \end{pmatrix} \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} 0 \\ \mathbb{1} \end{pmatrix} R \quad (63)$$

$$= R^\dagger R + R^\dagger (U_{11}^\dagger Q^\dagger K U_{01} + U_{01}^\dagger K^\dagger Q U_{11}) R \quad (64)$$

$$= R^\dagger R + R^\dagger \left((\mathbb{1} + X^\dagger R^\dagger) Q^\dagger K U_{01} + U_{01}^\dagger K^\dagger Q (\mathbb{1} + R X) \right) R \quad (65)$$

$$= R^\dagger R, \quad (66)$$

where we used again that $K^\dagger Q R = R^\dagger Q^\dagger K = 0$ in the last line.

We can now put all the terms obtained by multiplying Eq. (49) with $K_t^\dagger$ from the left together and find

$$K_t^\dagger \left(\ddot{K}_t + \dot{K}_t \dot{K}_t^\dagger K_t - K_t \dot{K}_t^\dagger \dot{K}_t - \dot{K}_t K_t^\dagger \dot{K}_t \right) = A^2 - R^\dagger R - AA^\dagger - A^\dagger A - R^\dagger R - A^2 \quad (67)$$

$$= -2R^\dagger R - A^\dagger A - AA^\dagger. \quad (68)$$

We see that these remaining terms are Hermitian and therefore the left-hand side of Eq. (49) is in the normal space. $\square$

B. Complex Newton equation

In this section, we derive the Riemannian Hessian operator and solve the Hessian equation to obtain an update direction on the tangent space, which we can follow along the geodesic defined in Eq. (27). This can be done for each gate individually, or simultaneously over all gates, in which case we operate on the Cartesian product $\text{St}(D, d)^{\times n}$ of single Stiefel manifolds. We consider the latter case, whereby we obtain the single Stiefel Newton equation (Eq. (77)) as a byproduct. The method is based on the real case [48]. See also [49] for a recent treatment of second order optimization on the complex Stiefel manifold, where instead of following geodesics, each optimization step is done in Euclidean space followed by a projection onto the manifold.

First let us make a general observation that will be useful at several points.

Lemma 3 ([49], Theorem 14). *Let $f : T_K \text{St}(D, d) \rightarrow \mathbb{C}$ be a $\mathbb{C}$ -linear function and let $\langle \cdot, \cdot \rangle_K$ be the canonical metric on $\text{St}(D, d)$ as defined in Eq. (24). Then the solution to*

$$\text{Re} \{f(\Delta)\} = \langle X, \Delta \rangle_K \quad \forall \Delta \in T_K \text{St}(D, d) \quad (69)$$

is given by $X = F^* - KF^T K \in T_K \text{St}(D, d)$, where F is chosen such that $f(\Delta) = \text{Tr}(F^T \Delta)$.

Proof. It is straightforward to see that $X \in T_K \text{St}(D, d)$ by applying the projector onto the tangent space: $P_{T_K}(X) = F^* - KF^T K - \frac{1}{2}K(K^\dagger F^* + F^T K) + \frac{1}{2}K(F^T K + K^\dagger F^*) = F^* - KF^T K$.

To show that X solves Eq. (3), we will use that $K^\dagger \Delta$ is skew Hermitian, as well as the fact that $\text{Re Tr}[HS] = 0$ for any skew Hermitian matrix S and Hermitian matrix H . Plugging $X = F^* - KF^T K$ into Eq. (3) we obtain

$$\begin{aligned} \langle X, \Delta \rangle_K &= \text{Re Tr}[X^\dagger (\mathbb{1} - \frac{1}{2}KK^\dagger)\Delta] \\ &= \text{Re Tr}[(F^T - K^\dagger F^* K^\dagger)(\mathbb{1} - \frac{1}{2}KK^\dagger)\Delta] \\ &= \text{Re Tr}\left[\left(F^T - \frac{1}{2}(F^T K + K^\dagger F^*)K^\dagger\right)\Delta\right] \\ &= \text{Re Tr}[F^T \Delta] - \text{Re Tr}[\text{herm}(F^T K)K^\dagger \Delta] \\ &= \text{Re Tr}[F^T \Delta]. \end{aligned}$$

□

Our goal is to simultaneously update all gates along the geodesic $\bigoplus_{i=1}^n \mathcal{K}_i(t) \in \text{St}(D, d)^{\times n}$, with the single Stiefel geodesics $\mathcal{K}_i(t)$ being given by Eq. (27). We define the initial directions as $\Delta_i = \dot{\mathcal{K}}_i(0)$. The first step to identify the Riemannian gradient and Hessian is to compute the second order Taylor series expansion of $\mathcal{L}$ in t at $t = 0$. Using $(\mathcal{K}_i)_{lm} \equiv \mathcal{K}_{ilm}$ and Einstein notation we find

$$\begin{aligned} \mathcal{L}(\mathcal{K}_1 \oplus \dots \oplus \mathcal{K}_n; \mathcal{K}_1^* \oplus \dots \oplus \mathcal{K}_n^*) &= \mathcal{L}|_{t=0} + 2 \text{Re} \left\{ \frac{\partial \mathcal{L}}{\partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{ilm}}{\partial t} \right\} \Big|_{t=0} \cdot t \\ &+ 2 \text{Re} \left\{ \frac{\partial^2 \mathcal{L}}{\partial \mathcal{K}_{jop} \partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{jop}}{\partial t} \frac{\partial \mathcal{K}_{ilm}}{\partial t} + \frac{\partial^2 \mathcal{L}}{\partial \mathcal{K}_{jop}^* \partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{jop}^*}{\partial t} \frac{\partial \mathcal{K}_{ilm}}{\partial t} + \frac{\partial \mathcal{L}}{\partial \mathcal{K}_{ilm}} \frac{\partial^2 \mathcal{K}_{ilm}}{\partial t^2} \right\} \Big|_{t=0} \cdot t^2/2 + \mathcal{O}(t^3). \end{aligned} \quad (70)$$

We have $\frac{\partial \mathcal{K}_{ilm}}{\partial t} \Big|_{t=0} = (\Delta_i)_{lm}$ and define $(\mathcal{L}_{\mathcal{K}_i})_{lm} := \frac{\partial \mathcal{L}}{\partial \mathcal{K}_{ilm}}$, so that we can write $\frac{\partial \mathcal{L}}{\partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{ilm}}{\partial t} = \text{Tr}(\mathcal{L}_{\mathcal{K}_i}^T \Delta_i)$ and $\frac{\partial \mathcal{L}}{\partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{ilm}}{\partial t} =: \mathcal{L}_{\mathcal{K}_i}[\Delta_i]$. In a similar fashion we define $\mathcal{L}_{\mathcal{K}_j \mathcal{K}_i}[\Delta_j, \Delta_i] := \frac{\partial^2 \mathcal{L}}{\partial \mathcal{K}_{jop} \partial \mathcal{K}_{ilm}} \frac{\partial \mathcal{K}_{jop}}{\partial t} \frac{\partial \mathcal{K}_{ilm}}{\partial t}$, where $\mathcal{L}_{\mathcal{K}_j \mathcal{K}_i}[\cdot, \cdot]$ is a bilinear function, which is symmetric per definition via the second derivative. For more details on how to compute these derivatives for the objective function used in the main text, see Appendix C.

Before determining the relevant terms for the update on $\text{St}(D, d)^{\times n}$ we first consider the gradient and Hessian, as well as the Newton equation for a single variable $K \in \text{St}(D, d)$, leaving all others constant.

The Riemannian gradient $G \in T_K \text{St}(D, d)$ can be identified from the first order term in the Taylor expansion via its definition [49]

$$2 * \text{Re} \{ \mathcal{L}_K[\Delta] \} = \langle G, \Delta \rangle_K \quad \forall \Delta \in T_K \text{St}(D, d). \quad (71)$$

The solution for G in the canonical metric (24) is given by

$$G = 2 (\mathcal{L}_K^* - K \mathcal{L}_K^T K), \quad (72)$$

as per Lemma 3.

Lemma 4. *The Riemannian Hessian $\text{Hess} : T_K \text{St}(D, d) \times T_K \text{St}(D, d) \rightarrow \mathbb{R}$ of a function $\mathcal{L} : \text{St}(D, d) \rightarrow \mathbb{R}$ with respect to the canonical metric on $\text{St}(D, d)$ is given by*

$$\begin{aligned} \text{Hess}(\Delta, \Omega) = & 2 \text{Re} \{ \mathcal{L}_{KK}[\Delta, \Omega] + \mathcal{L}_{K^*K}[\Delta^*, \Omega] \} \\ & + \text{Re} \{ \text{Tr} [\mathcal{L}_K^T(\Delta K^\dagger \Omega + \Omega K^\dagger \Delta)] - \text{Tr} [\mathcal{L}_K^T K(\Delta^\dagger \Pi \Omega + \Omega^\dagger \Pi \Delta)] \} . \end{aligned} \quad (73)$$

Proof. According to [89, Proposition 5.5.5], we can compute $\text{Hess}(\Delta, \Omega)$ via

$$\text{Hess}(\Delta, \Omega) = \frac{1}{2} \frac{d^2}{dt^2} [\mathcal{L}(K(t(\Delta + \Omega))) - \mathcal{L}(K(t\Delta)) - \mathcal{L}(K(t\Omega))] \Big|_{t=0} \quad (74)$$

where $K(t\Delta)$ satisfies $\dot{K}(t\Delta)|_{t=0} = \Delta$ (see [48] for a discussion of the real case). The individual terms in Eq. (74) can be determined from our general Taylor approximation in Eq. (70), i.e. with $i = j = 1$, if we take $K = K_1$. The term $\frac{\partial \mathcal{L}}{\partial K} \left[\frac{\partial^2 K(t\Delta)}{\partial t^2} \right] \Big|_{t=0}$ contains second derivatives of the geodesic given in Lemma 2, which we write out next. $\ddot{K}(t)$ is given by

$$\begin{aligned} \ddot{K}(t) &= (K \ Q) \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix}^2 \begin{pmatrix} \mathbf{1} \\ 0 \end{pmatrix} \\ &= (K \ Q) \exp \left(t \begin{pmatrix} A & -R^\dagger \\ R & 0 \end{pmatrix} \right) \begin{pmatrix} A^2 - R^\dagger R \\ RA \end{pmatrix} . \end{aligned}$$

It follows using $QR = (\mathbf{1} - KK^\dagger)\Delta =: \Pi\Delta$ and $A = K^\dagger\Delta$ from the definition of the geodesic, that

$$\begin{aligned} \ddot{K}(0) &= K(A^2 - R^\dagger R) + QRA \\ &= K(A^2 - R^\dagger Q^\dagger QR) + \Pi\Delta K^\dagger \Delta \\ &= K(K^\dagger \Delta K^\dagger \Delta - \Delta^\dagger \Pi^\dagger \Pi \Delta) + \Pi\Delta K^\dagger \Delta \\ &= K(K^\dagger \Delta K^\dagger \Delta - \Delta^\dagger \Pi \Delta) + \Delta K^\dagger \Delta - KK^\dagger \Delta K^\dagger \Delta \\ &= \Delta K^\dagger \Delta - K\Delta^\dagger \Pi \Delta . \end{aligned}$$

Putting the terms together, we arrive at

$$\frac{d^2}{dt^2} \mathcal{L}(K(t\Delta)) \Big|_{t=0} = 2 \text{Re} \{ \mathcal{L}_{KK}[\Delta, \Delta] + \mathcal{L}_{K^*K}[\Delta^*, \Delta] \} + \text{Re} \{ \text{Tr} [\mathcal{L}_K^T(\Delta K^\dagger \Delta - K\Delta^\dagger \Pi \Delta)] \} . \quad (75)$$

The terms involving $\mathcal{L}_{KK}$ and $\mathcal{L}_{K^*K}$ satisfy $\mathcal{L}_{KK}[\Delta, \Omega] = \mathcal{L}_{KK}[\Omega, \Delta]$ and $\mathcal{L}_{K^*K}[\Delta^*, \Omega] = \mathcal{L}_{K^*K}[\Omega^*, \Delta]$, by the symmetry of second derivatives. Using this symmetry property we obtain the full Hessian (74), which turns out to be

$$\begin{aligned} \text{Hess}(\Delta, \Omega) = & 2 \text{Re} \{ \mathcal{L}_{KK}[\Delta, \Omega] + \mathcal{L}_{K^*K}[\Delta^*, \Omega] \} \\ & + \text{Re} \{ \text{Tr} [\mathcal{L}_K^T(\Delta K^\dagger \Omega + \Omega K^\dagger \Delta)] - \text{Tr} [\mathcal{L}_K^T K(\Delta^\dagger \Pi \Omega + \Omega^\dagger \Pi \Delta)] \} , \end{aligned} \quad (76)$$

where it is helpful to note that Eq. (76) is related to Eq. (75) via a symmetrization of the $\mathcal{L}_K$ term. $\square$

Theorem 5. *Let $\text{vec}(\Delta)$ be the row major vectorization of $\Delta \in T_K \text{St}(D, d)$. Furthermore, let T and $\tilde{\mathcal{L}}_{KK}$ be defined by $T \text{vec}(X) = \text{vec}(X^T)$ and $\mathcal{L}_{KK}(\Delta, \cdot) = \tilde{\mathcal{L}}_{KK}^T \text{vec}(\Delta)$. Then the solution Δ of the linear equation in $\text{vec}(\Delta)$ and $\text{vec}(\Delta^*)$ given by*

$$\left(\tilde{\mathcal{L}}_{K^*K}^\dagger - (K \otimes K^T) T \tilde{\mathcal{L}}_{KK}^T - \frac{1}{2} \mathbf{1} \otimes (K^T \mathcal{L}_K) - \frac{1}{2} (K \mathcal{L}_K^T) \otimes \mathbf{1} - \frac{1}{2} \Pi \otimes (\mathcal{L}_K^\dagger K^*) \right) \text{vec}(\Delta) \quad (77)$$

$$+ \left(\tilde{\mathcal{L}}_{KK}^\dagger - (K \otimes K^T) T \tilde{\mathcal{L}}_{K^*K}^T + \frac{1}{2} (\mathcal{L}_K^* \otimes K^T) T + \frac{1}{2} (K \otimes \mathcal{L}_K^\dagger) T \right) \text{vec}(\Delta^*) = -\frac{1}{2} \text{vec}(G) \quad (78)$$

is the update direction along the geodesic given in Lemma 2 for the complex Newton method of a real function $\mathcal{L}$ at position $K \in \text{St}(D, d)$.

Proof. The update direction Δ for the standard Newton method [48] is determined through the equation

$$\text{Hess}(\Delta, \Omega) = -\langle G, \Omega \rangle_K \quad \forall \Omega \in T_K \text{St}(D, d) , \quad (79)$$

which can be solved by rewriting the left-hand side as $\text{Hess}(\Delta, \Omega) = \langle f(\Delta), \Omega \rangle_K$ (for some yet to be determined f) and setting $\Omega = P_T(X)$ with arbitrary matrix X . This leads us to

$$\langle f(\Delta), P_T(X) \rangle_K = -\langle G, P_T(X) \rangle_K \quad (80)$$

$$\langle P_T(f(\Delta)), X \rangle_K = -\langle G, X \rangle_K \quad \forall X \in \mathbb{C}^{D \times d}; \quad (81)$$

in the second line the scalar product is extended from the canonical scalar product initially defined on $T_K \text{St}(D, d)$ to $\mathbb{C}^{D \times d}$ and we will use the same notation for both. The second line follows from the fact that any matrix X can be decomposed as $X = P_T(X) + P_N(X)$ and from $\langle A, B \rangle_K = 0$ for $A \in T_K \text{St}(D, d)$ and $B \in N_K \text{St}(D, d)$. To determine $f(\Delta)$ we split it into three terms $f(\Delta) = f_{KK}(\Delta) + f_{K^*K}(\Delta) + f_K(\Delta)$, where $f_{KK}(\Delta)$, $f_{K^*K}(\Delta)$ and $f_K(\Delta)$ depend only on $\mathcal{L}_{KK}$, $\mathcal{L}_{K^*K}$ and $\mathcal{L}_K$ respectively (compare Eq. (76)).

We first look at the term $2 \text{Re} \{ \text{Tr} (\mathcal{L}_{KK}[\Delta, \Omega]) \} = 2 \text{Re} \{ \text{Tr} (\mathcal{L}_{KK}[\Delta, \cdot]^T \Omega) \}$, where $\mathcal{L}_{KK}[\Delta, \cdot]$ is in $\mathbb{C}^{D \times d}$.

To solve $2 \text{Re} \{ \text{Tr} (\mathcal{L}_{KK}[\Delta, \cdot]^T \Omega) \} = \langle f_{KK}(\Delta), \Omega \rangle_K$ for all $\Omega \in T_K \text{St}(D, d)$ and to find f_{KK} we use Lemma 3 and obtain

$$f_{KK}(\Delta) = 2 (\mathcal{L}_{KK}[\Delta, \cdot]^* - K \mathcal{L}_{KK}[\Delta, \cdot]^T K). \quad (82)$$

The same argument can be made for the $\mathcal{L}_{K^*K}$ term, leading to

$$f_{K^*K}(\Delta^*) = 2 (\mathcal{L}_{K^*K}[\Delta^*, \cdot]^* - K \mathcal{L}_{K^*K}[\Delta^*, \cdot]^T K). \quad (83)$$

To identify $f_K(\Delta)$ we rewrite the second line in (76) as follows:

$$\begin{aligned} & \text{Re} \{ \text{Tr} [\mathcal{L}_K^T (\Delta K^\dagger \Omega + \Omega K^\dagger \Delta)] - \text{Tr} [\mathcal{L}_K^T K (\Delta^\dagger \Pi \Omega + \Omega^\dagger \Pi \Delta)] \} \\ &= \text{Re} \{ \text{Tr} [(\mathcal{L}_K^T \Delta K^\dagger + K^\dagger \Delta \mathcal{L}_K^T - \mathcal{L}_K^T K \Delta^\dagger \Pi - (\Pi \Delta \mathcal{L}_K^T K)^\dagger) \Omega] \} \\ &\stackrel{!}{=} \text{Re} \{ \text{Tr} [f_K(\Delta)^\dagger \Gamma \Omega] \}, \end{aligned}$$

where we used $\text{Re} \{ \text{Tr} [AB^\dagger] \} = \text{Re} \{ \text{Tr} [A^\dagger B] \}$. Thus we find

$$\begin{aligned} f_K(\Delta) &= [(\mathcal{L}_K^T \Delta K^\dagger + K^\dagger \Delta \mathcal{L}_K^T - \mathcal{L}_K^T K \Delta^\dagger \Pi - (\Pi \Delta \mathcal{L}_K^T K)^\dagger) \Gamma^{-1}]^\dagger \\ &= 2K \Delta^\dagger \mathcal{L}_K^* + \Gamma^{-1} \mathcal{L}_K^* \Delta^\dagger K - \Pi \Delta K^\dagger \mathcal{L}_K^* - \Pi \Delta \mathcal{L}_K^T K \end{aligned}$$

by using $\Pi = \Pi^\dagger$, $\Pi \Gamma^{-1} = \Gamma^{-1} \Pi = \Pi$, $K^\dagger \Gamma^{-1} = 2K^\dagger$.

Eq. (81) implies $P_T(f(\Delta)) = -G$, and it remains to compute $P_T(f_K(\Delta))$, as $P_T(f_{KK}(\Delta)) = f_{KK}(\Delta)$ and $P_T(f_{K^*K}(\Delta^*)) = f_{K^*K}(\Delta^*)$ (see Lemma 3). After a straightforward computation using $\Gamma^{-1} = \mathbb{1} + K K^\dagger$, $\Pi \Gamma^{-1} = \Pi$, $P_T(\Pi Z) = P_T(Z)$, as well as $K^\dagger \Pi = \Pi K = 0$, we get

$$P_T(f_K(\Delta)) = -\Pi \Delta K^\dagger \mathcal{L}_K^* - 2 \text{skew}(\Delta \mathcal{L}_K^T) K - 2K \text{skew}(\mathcal{L}_K^T \Delta), \quad (84)$$

with $\text{skew}(A) = (A - A^\dagger)/2$.

Finally, by plugging Eqs. (82), (83) and (84) into Eq. (81), we obtain the Newton equation

$$\mathcal{L}_{KK}[\Delta, \cdot]^* - K \mathcal{L}_{KK}[\Delta, \cdot]^T K + \mathcal{L}_{K^*K}[\Delta^*, \cdot]^* - K \mathcal{L}_{K^*K}[\Delta^*, \cdot]^T K \quad (85)$$

$$- \frac{1}{2} \Pi \Delta K^\dagger \mathcal{L}_K^* - \text{skew}(\Delta \mathcal{L}_K^T) K - K \text{skew}(\mathcal{L}_K^T \Delta) = -G/2. \quad (86)$$

This is a linear equation in Δ and Δ^* that can be solved via rewriting it as an equation in $\text{vec}(\Delta)$. Using row-major vectorization with $\text{vec}(AXB) = (A \otimes B^T) \text{vec}(X)$, the matrices T and $\tilde{\mathcal{L}}_{KK}$ defined by $T \text{vec}(X) = \text{vec}(X^T)$ and $\mathcal{L}_{KK}(\Delta, \cdot) = \tilde{\mathcal{L}}_{KK}^T \text{vec}(\Delta)$, we arrive at the final equation for the single gate case

$$\begin{aligned} & \left(\tilde{\mathcal{L}}_{K^*K}^\dagger - (K \otimes K^T) T \tilde{\mathcal{L}}_{KK}^T - \frac{1}{2} \mathbb{1} \otimes (K^T \mathcal{L}_K) - \frac{1}{2} (K \mathcal{L}_K^T) \otimes \mathbb{1} - \frac{1}{2} \Pi \otimes (\mathcal{L}_K^\dagger K^*) \right) \text{vec}(\Delta) \\ & + \left(\tilde{\mathcal{L}}_{KK}^\dagger - (K \otimes K^T) T \tilde{\mathcal{L}}_{K^*K}^T + \frac{1}{2} (\mathcal{L}_K^* \otimes K^T) T + \frac{1}{2} (K \otimes \mathcal{L}_K^\dagger) T \right) \text{vec}(\Delta^*) = -\frac{1}{2} \text{vec}(G). \end{aligned}$$

□

Now for the simultaneous optimization over all gates on $\text{St}(D, d)^{\times n}$, the Hessian as defined in Eq. (74) is determined by including all terms in Eq. (70). The Newton equation reads

$$\text{Hess}(\Delta_1 \oplus \dots \oplus \Delta_n; \Omega_1 \oplus \dots \oplus \Omega_n) = - \sum_{i=1}^n \langle G_i, \Omega_i \rangle_{\mathcal{K}_i} \quad (87)$$

for all $\Omega_i \in T_{\mathcal{K}_i} \text{St}(D, d)$. The terms in Eq. (70) where $i = j$ are obtained from the single variable case. The mixed variable terms $f_{\mathcal{K}_i \mathcal{K}_j}(\Delta_i)$ and $f_{\mathcal{K}_i^* \mathcal{K}_j}(\Delta_i^*)$ still need to be determined. Analogous to Eq. (83) we need to solve

$$\begin{aligned} 2 \text{Re} \{ \text{Tr} (\mathcal{L}_{\mathcal{K}_i^* \mathcal{K}_j} [\Delta_i^*, \cdot]^T \Omega_j) \} &= \langle f_{\mathcal{K}_i^* \mathcal{K}_j}(\Delta_i), \Omega_j \rangle_{\mathcal{K}_j} \quad \forall \Omega_j \in T_{\mathcal{K}_j} \text{St} \quad \text{and} \\ 2 \text{Re} \{ \text{Tr} (\mathcal{L}_{\mathcal{K}_i \mathcal{K}_j} [\Delta_i, \cdot]^T \Omega_j) \} &= \langle f_{\mathcal{K}_i \mathcal{K}_j}(\Delta_i), \Omega_j \rangle_{\mathcal{K}_j} \quad \forall \Omega_j \in T_{\mathcal{K}_j} \text{St} . \end{aligned}$$

We can use Lemma 3 again and obtain

$$\begin{aligned} f_{\mathcal{K}_i^* \mathcal{K}_j}(\Delta_i) &= 2 \left(\mathcal{L}_{\mathcal{K}_i^* \mathcal{K}_j}^* [\Delta_i, \cdot] - \mathcal{K}_j \mathcal{L}_{\mathcal{K}_i^* \mathcal{K}_j} [\Delta_i^*, \cdot]^T \mathcal{K}_j \right) \quad \text{and} \\ f_{\mathcal{K}_i \mathcal{K}_j}(\Delta_i) &= 2 \left(\mathcal{L}_{\mathcal{K}_i \mathcal{K}_j}^* [\Delta_i^*, \cdot] - \mathcal{K}_j \mathcal{L}_{\mathcal{K}_i \mathcal{K}_j} [\Delta_i, \cdot]^T \mathcal{K}_j \right) , \end{aligned}$$

which satisfy $f_{\mathcal{K}_i^* \mathcal{K}_j}(\Delta_i^*) \in T_{\mathcal{K}_j} \text{St}(D, d)$ and $f_{\mathcal{K}_i \mathcal{K}_j}(\Delta_i) \in T_{\mathcal{K}_j} \text{St}(D, d)$. The full Newton equation on $\text{St}(D, d)^{\times n}$ in vectorized form then reads

$$\begin{aligned} \bigoplus_{i=1}^n \left(\tilde{\mathcal{L}}_{\mathcal{K}_i^* \mathcal{K}_i}^\dagger - (\mathcal{K}_i \otimes \mathcal{K}_i^T) T \tilde{\mathcal{L}}_{\mathcal{K}_i \mathcal{K}_i}^T - \frac{1}{2} \mathbb{1} \otimes (\mathcal{K}_i^T \mathcal{L}_{\mathcal{K}_i}) - \frac{1}{2} (\mathcal{K}_i \mathcal{L}_{\mathcal{K}_i}^T) \otimes \mathbb{1} - \frac{1}{2} \Pi \otimes (\mathcal{L}_{\mathcal{K}_i}^\dagger \mathcal{K}_i^*) \right) \text{vec}(\Delta_i) \\ + \bigoplus_{i=1}^n \sum_{j:j \neq i} \left(\tilde{\mathcal{L}}_{\mathcal{K}_j^* \mathcal{K}_i}^\dagger - (\mathcal{K}_i \otimes \mathcal{K}_i^T) T \tilde{\mathcal{L}}_{\mathcal{K}_j \mathcal{K}_i}^T \right) \text{vec}(\Delta_j) \\ + \bigoplus_{i=1}^n \left(\tilde{\mathcal{L}}_{\mathcal{K}_i \mathcal{K}_i}^\dagger - (\mathcal{K}_i \otimes \mathcal{K}_i^T) T \tilde{\mathcal{L}}_{\mathcal{K}_i^* \mathcal{K}_i}^T + \frac{1}{2} (\mathcal{L}_{\mathcal{K}_i}^* \otimes \mathcal{K}_i^T) T + \frac{1}{2} (\mathcal{K}_i \otimes \mathcal{L}_{\mathcal{K}_i}^\dagger) T \right) \text{vec}(\Delta_i^*) \\ + \bigoplus_{i=1}^n \sum_{j:j \neq i} \left(\tilde{\mathcal{L}}_{\mathcal{K}_j \mathcal{K}_i}^\dagger - (\mathcal{K}_i \otimes \mathcal{K}_i^T) T \tilde{\mathcal{L}}_{\mathcal{K}_j^* \mathcal{K}_i}^T \right) \text{vec}(\Delta_j^*) \\ = - \frac{1}{2} \bigoplus_{i=1}^n \text{vec}(G_i) . \end{aligned} \quad (88)$$

We are now faced with an equation of the type $A\mathbf{x} + B\mathbf{x}^* = \mathbf{c}$, a solution to which can be obtained by solving

$$\begin{pmatrix} A & B \\ B^* & A^* \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ \mathbf{x}^* \end{pmatrix} = \begin{pmatrix} \mathbf{c} \\ \mathbf{c}^* \end{pmatrix} . \quad (89)$$

Eq. (88) is an equation on the tangent space and can be solved by finding a basis therein. However, in order to avoid a basis change at every step, we choose to solve it on the ambient space by setting $\Delta_i = P_{T_i}(\Delta_i)$ and $\Delta_i^* = P_{T_i}^*(\Delta_i^*)$. The matrix equation for the update directions Δ_i is now given by

$$\underbrace{\begin{pmatrix} H_{G \leftarrow \Delta} \oplus_i P_{T_i} & H_{G \leftarrow \Delta^*} \oplus_i P_{T_i}^* \\ H_{G \leftarrow \Delta^*}^* \oplus_i P_{T_i}^* & H_{G \leftarrow \Delta} \oplus_i P_{T_i} \end{pmatrix}}_{=:H} \begin{pmatrix} \text{vec}(\Delta_1) \\ \vdots \\ \text{vec}(\Delta_n) \\ \text{vec}(\Delta_1^*) \\ \vdots \\ \text{vec}(\Delta_n^*) \end{pmatrix} = - \frac{1}{2} \begin{pmatrix} \text{vec}(G_1) \\ \vdots \\ \text{vec}(G_n) \\ \text{vec}(G_1^*) \\ \vdots \\ \text{vec}(G_n^*) \end{pmatrix} , \quad (90)$$

where the sub matrices $H_{G \leftarrow \Delta}$ and $H_{G \leftarrow \Delta^*}$ can be identified from Eq. (88).

The update directions for the saddle-free Newton method 1 are calculated by applying

$$(|H| + \lambda \mathbb{1})^{-1} \quad (91)$$

to the right-hand side of Eq. (90). In practice, we use $\frac{1}{2}(H + H^\dagger)$ as the Hessian, since there exist more efficient methods for diagonalizing a Hermitian matrix compared to an arbitrary matrix.

C. Complex Euclidean gradient and Hessian

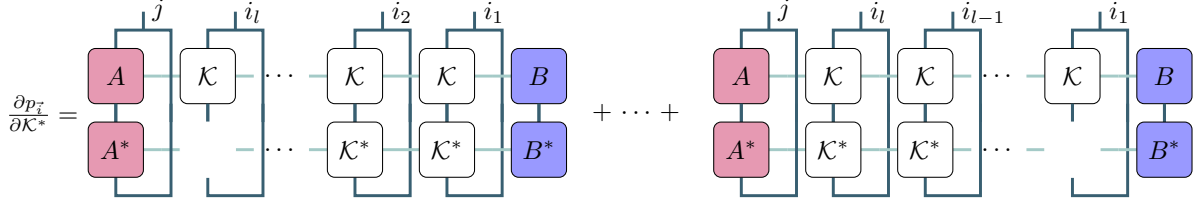


Figure 8. Tensor network representation of first derivative.

To compute the Riemannian Hessians for optimization on the Stiefel manifolds the complex Euclidean gradients and Euclidean Hessians are needed. We shall here go into more detail of their derivations for the least squares objective function. This section also goes into detail as to how the terms $\mathcal{L}_{\mathcal{K}\mathcal{K}}$, $\mathcal{L}_{\mathcal{K}^*\mathcal{K}}$ and their conjugates from Appendix B are calculated.

As the tensor network whose contraction yields $p_{j|i}$ in the cost function (20) is parameterized in terms of matrix variables and their conjugates, we use Wirtinger calculus for the derivatives and treat the conjugate variables as independent. A method for finding all the relevant terms in the Hessian for a scalar function of complex matrix variables is outlined, e.g. in Ref. [102], and we will summarize it in the following. See also Ref. [103] for a short derivation of the complex derivative and Hessian in the context of optimization in complex Euclidean space.

Our objective function is in general not analytic, as is the case with real valued functions of complex variables. This can be seen for the simplest case with one unitary gate U and $\rho = E = |0\rangle\langle 0|$, where $\mathcal{L} = |\langle 0|U|0\rangle|^4 = |U_{00}|^4 = (U_{00} * U_{00}^*)^2$. However the derivatives w.r.t. the real and imaginary parts of the matrix variables exist and one can define formal derivatives for $f : \mathbb{C}^{M \times N} \times \mathbb{C}^{M \times N} \rightarrow \mathbb{R}$ via

$$\begin{aligned} \frac{\partial f(Z, Z^*)}{\partial Z} &:= \frac{\partial f(Z, Z^*)}{\partial \text{Re}[Z]} - i \frac{\partial f(Z, Z^*)}{\partial \text{Im}[Z]}, \\ \frac{\partial f(Z, Z^*)}{\partial Z^*} &:= \frac{\partial f(Z, Z^*)}{\partial \text{Re}[Z]} + i \frac{\partial f(Z, Z^*)}{\partial \text{Im}[Z]}, \end{aligned}$$

where $\frac{\partial f(Z, Z^*)}{\partial Z} \in \mathbb{C}^{M \times N}$ with $\left(\frac{\partial f(Z, Z^*)}{\partial Z}\right)_{ij} = \frac{\partial f(Z, Z^*)}{\partial Z_{ij}}$. These formal derivatives have nice properties, for instance $\frac{\partial f(Z, Z^*)}{\partial Z^*}$ is the direction of maximum increase of f and $\frac{\partial f(Z, Z^*)}{\partial Z^*} = 0$ identifies a stationary point of f , see e.g. Ref. [102, Theorems 3.2 and 3.4]. Furthermore, the product rule and the chain rule apply as they do for real valued matrix variables.

As laid out in Ref. [102, Lemma 5.2], we can write the second order Taylor series of f as

$$f(Z + dZ, Z^* + dZ^*) = f(Z, Z^*) + \left(\frac{\partial}{\partial \text{vec}(Z)} f(Z, Z^*)\right) d \text{vec}(Z) + \left(\frac{\partial}{\partial \text{vec}(Z^*)} f(Z, Z^*)\right) d \text{vec}(Z^*) \quad (92)$$

$$+ \frac{1}{2} [d \text{vec}^T(Z^*) d \text{vec}^T(Z)] \begin{bmatrix} f_{ZZ^*} & f_{Z^*Z^*} \\ f_{ZZ} & f_{Z^*Z} \end{bmatrix} \begin{bmatrix} d \text{vec}(Z) \\ d \text{vec}(Z^*) \end{bmatrix} + r(dZ, dZ^*), \quad (93)$$

where the higher order contribution $r(dZ, dZ^*)$ satisfies

$$\lim_{(dZ, dZ^*) \rightarrow 0} \frac{r(dZ, dZ^*)}{\|(dZ, dZ^*)\|_F^2} = 0. \quad (94)$$

The second order derivatives are defined via

$$f_{ZZ} = \frac{\partial}{\partial \text{vec}(Z)^T} \frac{\partial}{\partial \text{vec}(Z)} f(Z, Z^*, \dots; y),$$

and similarly $f_{Z^*Z^*}$, f_{Z^*Z} and f_{ZZ^*} . The vectorization is to be understood as joining together of indices in a fixed order. For instance $\text{vec} : \mathbb{C}^{n \times d^2 \times d \times d} \rightarrow \mathbb{C}^{nd^4}$ vectorizes $\mathcal{K}$, where the individual d -dimensional legs are the matrix indices of the Kraus operators, and the d^2 index numbers the different Kraus operators. Note that $\text{vec}\left(\frac{\partial}{\partial Z} f(Z, Z^*, y)\right) = \frac{\partial}{\partial \text{vec}(Z)} f(Z, Z^*, y)$.

For the optimizations over A , $\mathcal{K}$, and B , we need the first and second derivatives of $\mathcal{L}$ by the respective variables and their conjugates. Let $Z \in \{A, A^*, \mathcal{K}, \mathcal{K}^*, B, B^*\}$ and $Y \in \{Z, Z^*\}$. Then

$$\begin{aligned} \frac{\partial}{\partial Z} \mathcal{L}(Z, \dots; y) &= \frac{2}{m} \sum_i (p_i(Z, \dots) - y_i) \frac{\partial p_i}{\partial Z}, \\ \frac{\partial}{\partial Y} \frac{\partial}{\partial Z} \mathcal{L}(Z, Y, \dots; y) &= \frac{\partial}{\partial Y} \frac{2}{m} \sum_i (p_i(Z, Y, \dots) - y_i) \frac{\partial p_i}{\partial Z} \\ &= \frac{2}{m} \sum_i \frac{\partial p_i(Z, Y, \dots)}{\partial Y} \frac{\partial p_i(Z, Y, \dots)}{\partial Z} + \frac{2}{m} \sum_i (p_i(Z, Y, \dots) - y_i) \frac{\partial^2 p_i(Z, Y, \dots)}{\partial Y \partial Z}, \end{aligned}$$

meaning that derivatives of the objective function reduce to the derivatives of the tensor p . Taking the derivative of a tensor network w.r.t. one of its constituent tensors can be easily done in the pictorial representation by removing the respective tensor. For instance, $\frac{\partial p_i}{\partial \mathcal{K}^*}$ can be calculated as shown in Figure 8, using the product rule. Care has to be taken for the order of open indices when removing a tensor. In practice, we do not calculate the full tensor $\frac{\partial p}{\partial \mathcal{K}^*}$ of size n^ℓ and only compute $\frac{\partial p_i}{\partial \mathcal{K}^*}$ for $i \in I$, since usually $|I| \ll n^\ell$.

D. Mean variation error dependence on the choice of objective function

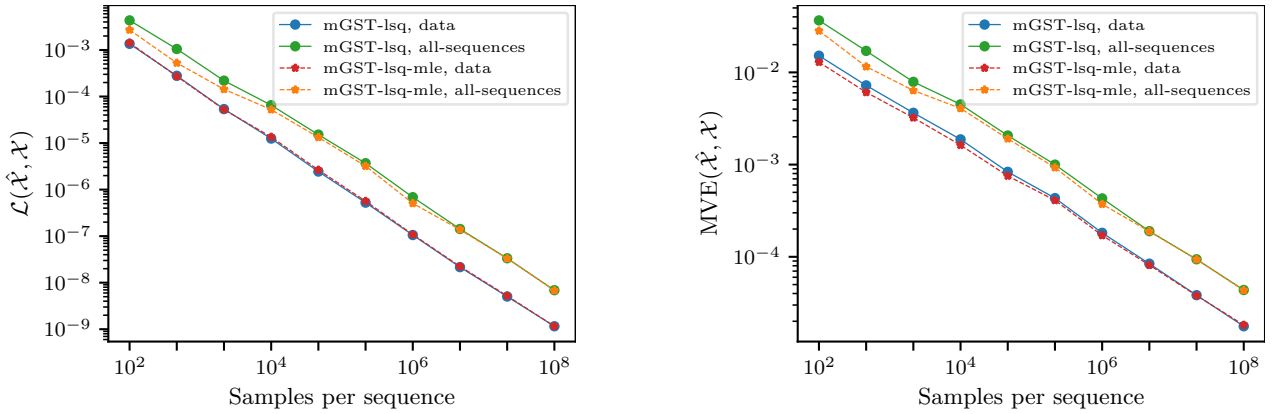


Figure 9. Effect of optimizing the log-likelihood function after the least squares objective function. The results of only the least squares optimization are denoted by **mGST-lsq** and those of additional log-likelihood optimization by **mGST-lsq-mle**. The **left** plot shows the least squares objective function $\mathcal{L}(\hat{\mathcal{X}}, \mathcal{X})$ on the data sequences (used for estimation), as well as on all sequences of a given length l (here $l = 7$). The **right** plot uses the MVE, again on data sequences or all sequences.

The underlying gate set is given by the XYI model with depolarizing noise of strength $p = 0.01$ on each gate and $p = 0.01$ on the initial state, as well as random unitary rotations $e^{i\gamma H}$ with $H \sim \text{GUE}$ and $\gamma = 0.01$ on each gate. To ensure **mGST-lsq** is fully converged, we set the desired relative precision to $\epsilon = 10^{-5}$ in the convergence criterion, cp. Eq. (30).

A well motivated alternative to the least squares objective function defined in Eq. (20) is the likelihood function

$$L_I(A, \mathcal{K}, B|y) := \prod_{i \in I} \prod_{j \in [n_E]} p_{j|i}(A, \mathcal{K}, B)^{k_{j|i}}, \quad (95)$$

where $y_{j|i} = k_{j|i}/m$ denotes again the relative number out of m times outcome j was measured for sequence i , see also Eq. (28). The likelihood function at a given model parametrization $(A, \mathcal{K}, B)$ and for measurement results y is precisely the probability of observing y , given the model probabilities $p(A, \mathcal{K}, B)$. To simplify the optimization, often the logarithm of the likelihood function is chosen, since it shares the same maxima. This log-likelihood function, which is also used in the final optimization procedure of **pyGSTi** [18] is then given by

$$\log L_I(A, \mathcal{K}, B|y) := m \sum_{i \in I} \sum_j y_{j|i} \log [p_{j|i}(A, \mathcal{K}, B)]. \quad (96)$$

In Figure 9 we show the effects of augmenting mGST (which is by default run on the least squares objective function for numerical reasons) with the log-likelihood function after a least squares estimate was found. To do this, we use Algorithm 2 with the negative log-likelihood function (96) as objective function.

We observe that for the XYI-gate set, which we use to compare mGST and pyGSTi in the main text, optimizing the log likelihood function decreases the mean variation error on the data sequences as well as on all sequences of the same length. The improvement is stronger for fewer samples and becomes negligible at around 10^6 samples. Interestingly, log-likelihood optimization also improves the least squares error on all sequences, at the cost of slightly increasing it on the data sequences. This indicates that in the sample count range of $10^2 - 10^6$, optimizing the log-likelihood function leads to less overfitting.

E. Noise-mitigation of shadow estimation with GST characterization

In Section V we numerically demonstrated how the results of a low-rank GST experiment can be used to correct estimation protocols based on inverting an informationally complete POVM. Such protocols are recently referred to as shadow estimation [46]. We now give a short mathematical description of the method and explain how low rank GST estimates can be included in a scalable way.

In the following, we use bra-ket notation also for the space of linear operators $L(\mathcal{H})$ and its dual space as defined by the canonical isomorphism induced by the Hilbert Schmidt inner product $(O|\rho) = \text{Tr}(O^\dagger \rho)$. The quantum channel of a unitary U is written by the corresponding calligraphic letter, e.g. $\mathcal{U}|\rho) \equiv |U\rho U^\dagger)$.

We consider the task of estimating the expectation value of multiple observables in an unknown quantum state that we can repeatedly prepare on a quantum device. Being able to measure an informationally complete POVM $\{\Pi_x\}$ on the state, one can construct an estimator for an observable O . *Informationally completeness* is equivalent to, in mathematical terms, the POVM constituting a *frame* for $L(\mathcal{H})$, and the associated frame operator $\mathcal{M} = \sum_x |\Pi_x)(\Pi_x|$ being invertible, see e.g. Ref. [104]. We can calculate the canonical dual frame to the POVM as $|\tilde{\Pi}_x) = \mathcal{M}^{-1} |\Pi_x)$. By construction we have the frame duality relation

$$\sum_x |\tilde{\Pi}_x)(\Pi_x| = \text{Id}_{L(\mathcal{H})}. \quad (97)$$

Thus, for any state ρ ,

$$(O|\rho) = \sum_x (O|\tilde{\Pi}_x)(\Pi_x|\rho). \quad (98)$$

By Born's rule, repeated measurements of the POVM yield i.i.d. samples $\Omega = (x_1, \dots, x_m)$ from the distribution with density $p_\rho(x) = (\Pi_x|\rho)$. Given Ω we can calculate the empirical mean estimator

$$\hat{o} = \frac{1}{|\Omega|} \sum_{x \in \Omega} (O|\tilde{\Pi}_x), \quad (99)$$

and by (98), $\mathbb{E}[\hat{o}] = (O|\rho)$.

The sequence of dual frame elements $(\tilde{\Pi}_{x_1}, \dots, \tilde{\Pi}_{x_m})$ given by the measured samples Ω has been called the *classical shadow* of ρ in Ref. [46].

A practical implementation of an informationally complete POVM on a digital quantum computer can be realized with measurements in randomly selected bases from a sufficiently large group. To be explicit, we will consider the simplest and perhaps most well-known example: the measurement in a randomly chosen multi-qubit Pauli-basis. The POVM can be implemented by applying a random (different) local Clifford rotation on every qubit and measuring in the computational basis. For informational-completeness it is sufficient to choose the rotations uniformly from the set $C = \{\text{Id}, H, HS\}$, where H is the Hadamard gate and S the phase gate. In our notation, we consider POVM effects $\Pi_x = \Pi_{\mathbf{g}, \mathbf{b}} = \otimes_l \Pi_{g_l, b_l}$ indexed by $C^n \times \{0, 1\}^n$ that are the tensor products of the local POVM effects $\Pi_{g_l, b_l} = \frac{1}{3} g_l^\dagger |b_l\rangle\langle b_l| g_l$ with $g_l \in C$. Let $\{\hat{\sigma}_k \mid k \in \{0, 1, 2, 3\}\}$ denote the Pauli matrices normalized in Frobenius

norm. The frame operator is given by

$$3^n \mathcal{M} = \frac{1}{3^n} \left(\sum_{U \in \{\mathbb{1}, H, HS\}} \mathcal{U}^\dagger (|\hat{\sigma}_0\rangle \langle \hat{\sigma}_0| + |\hat{\sigma}_3\rangle \langle \hat{\sigma}_3|) \mathcal{U} \right)^{\otimes n} \quad (100)$$

$$= \frac{1}{3^n} \left(3 |\hat{\sigma}_0\rangle \langle \hat{\sigma}_0| + \sum_i |\hat{\sigma}_i\rangle \langle \hat{\sigma}_i| \right)^{\otimes n} \quad (101)$$

$$= \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \frac{1}{3} & 0 & 0 \\ 0 & 0 & \frac{1}{3} & 0 \\ 0 & 0 & 0 & \frac{1}{3} \end{pmatrix}^{\otimes n}, \quad (102)$$

where the matrix in the last line is represented in the Pauli-basis. Since $\mathcal{M}^{-1}$ acts on qubit l as $\mathcal{M}_l^{-1}(X) = 3X - \text{Tr}(X)\mathbb{1}$ for any X , we find $\tilde{\Pi}_{g_l, b_l} = \bigotimes_{i=1}^n \left(3g_l^\dagger |b_l\rangle \langle b_l| g_l - \mathbb{1} \right)$ [46].

Ref. [46] showed that when using random Pauli basis measurements, the variance of the mean estimator for estimating local observables does not scale with the system size. Using a median-of-means estimator to boost the confidence, Ref. [46] further establishes that the expectation value of M different k -local observables can be estimated to ϵ -additive precision from $\mathcal{O}(\log(M)4^k/\epsilon^2)$ state copies.

Experimental implementations of the POVM are prone to errors, effectively implementing a noisy POVM with effects $\Pi_x^\natural$. In the envisioned implementation here, noise sources effect the implementation of the gates C and the noise induces a bias in the estimator for the observables. However, if the noise is characterized to some extent we can correct the estimators for this bias. To this end, let $\mathcal{M}^\natural$ be the (half-sided) noisy frame operator $\mathcal{M}^\natural = \sum_x |\Pi_x^\natural\rangle \langle \Pi_x^\natural|$. If we know $\mathcal{M}^\natural$ in the classical post-processing, we can calculate a dual frame to $(\Pi_x^\natural|$ by $|\tilde{\Pi}_x^\natural\rangle = \mathcal{M}^{\natural^{-1}} |\Pi_x^\natural\rangle$. Note that using the ‘half-sided noisy’ frame operator instead of the frame operator of the noisy POVM yields an expression of a dual frame in terms of the ideal POVM and not the noisy POVM. Using $\{\tilde{\Pi}_x^\natural\}$ instead of the ideal dual-frame in (99) yields unbiased estimators of observables even in the presence of noise, thus, effectively mitigating the noise.

This motivates our approach to noise mitigated shadow estimation. Having extracted a noise model via gate set tomography, we can numerically estimate $\mathcal{M}^\natural$ and, thus, construct (approximately) unbiased estimators. Our method is summarized in Protocol 1 below.

Protocol 1: GST-mitigated shadow estimation

input: Target observable O , native local gate set $\mathcal{G}$ with $C \subseteq \mathcal{G}$	
1 Perform 2-qubit-mGST on implementation of $\mathcal{G}$ for qubit pairs $(1, 2), \dots, (n-1, n)$	} GST
2 Gauge optimize mGST estimators to unitary target gates	
3 for $i \in [N]$ do	} Cl. Shadows data acquisition
4 Select setting $g \in C^{\otimes n}$ uniformly at random	
5 Measure $\mathcal{U}_g \rho\rangle$ in the standard basis	
6 Save setting g and outcome b	
7 end	
8 Construct $\mathcal{M}^\natural$ from mGST gate estimates	} Combined post processing
9 Compute single shot estimators $\{\hat{o}_i = (O \mathcal{M}^{\natural^{-1}} \Pi_{g, b})\}_{i=1}^N$	
10 return $\hat{O} = \text{mean or median-of-means}(\{\hat{o}_i\})$	

The results in Section V demonstrate our scheme numerically in simple but already practically relevant settings that we describe in the following. When the multi-qubit unitaries and computational basis measurement implementing the POVM factorize into local tensor products, so does the ideal frame operator $\mathcal{M}$ and the dual frame (shadow). But due to correlated gate-dependent noise, $\mathcal{M}^\natural$ might not exhibit this computationally tractable structure. Furthermore, characterizing the implementation of exponentially many multi-qubit unitaries and the basis measurements without additionally assumption is infeasible. In practice, however, noise-induced correlations and crosstalk might still predominantly affect a limited number of qubits simultaneously. For example, when noise predominantly affects neighboring qubits, we can use the implementation $\mathcal{X}$ of a gate set including $C \times C$ on neighboring qubits extracted via mGST to calculate $\mathcal{M}^\natural$. To this end, let $\mathcal{G}_{g_1, g_2}^{(i, i+1)}$ denote the two-qubit process implementing the gate $g_1 \times g_2$ on qubit i and $i+1$. For simplicity we ignore errors in the computational basis measurement. We set

$\Pi_{(g_1, g_2), (b_1, b_2)}^{(i, i+1)} = \frac{1}{9} (\mathcal{G}_{g_1, g_2}^{(i, i+1)})^\dagger |b_1, b_2\rangle \langle b_1, b_2| (\mathcal{G}_{g_1, g_2}^{(i, i+1)})^\dagger$ and numerically calculate

$$\mathcal{M}_{i, i+1}^\natural = \sum_{g_1, g_2 \in C, b_1, b_2 \in \{0, 1\}} |\Pi_{g_1, b_1}\rangle |\Pi_{g_2, b_2}\rangle (\Pi_{(g_1, g_2), (b_1, b_2)}^{(i, i+1)})^\dagger. \quad (103)$$

This amounts to calculating a 16×16 matrix in the Pauli-basis that can be easily inverted. The noise-mitigated single-shot estimators, thus, read

$$(O|\tilde{\Pi}_{\mathbf{g}, \mathbf{b}}^\natural) = (O|\bigotimes_{i=1}^{n/2} (M_{2i, 2i+1}^\natural)^{-1} |\Pi_{g_{2i}, b_{2i}}\rangle |\Pi_{g_{2i+1}, b_{2i+1}}\rangle). \quad (104)$$

With this expression at hand, the rest of the protocol consists of computing the mean or median-of-means from a collection of single shot estimators, following the standard method of shadow estimation [46].

Ref. [101] proposes a complimentary approach for robust shadow-estimation, also inferring an approximation of the noisy frame operator from a separate calibration experiment. Under the assumption of gate-independent noise, the authors derive a 2^n parameter expression for $\mathcal{M}^\natural$ as a Pauli-noise channel and devise (SPAM-robust) RB-style experiments to learn arbitrarily many of its parameters, where each parameter corresponds to an irreducible representation of the local Clifford group. We find, in the gate-dependent noise model used here, that the frame operators significantly deviate from being a Pauli-noise channel. For this reason, this particular setting is more amenable to GST-mitigated shadows than to the protocol of Ref. [101]. The plots on the left in Figure 7 show that already a typical $n = 2$ frame operator with gate-dependent noise does not adhere to being diagonal in Pauli basis.

Ultimately, we envision that different robust and self-consistent noise and error characterization protocols, such as mGST for local coherent errors and RB for incoherent noise strength in different irreducible representations, can be combined to arrive at accurate and scalable estimates of the effective frame operator in the presence of noise.

ACRONYMS

RB	randomized benchmarking	1
CPT	completely positive and trace preserving	2
GST	gate set tomography	2
GUE	Gaussian unitary ensemble	8
MPS	matrix product state	4
MVE	mean variation error	6
MSE	mean squared error	6
POVM	positive operator valued measure	3
SFN	saddle free Newton	7
SPAM	state preparation and measurement	1
XEB	cross-entropy benchmarking	3

-
- [1] J. Eisert, D. Hangleiter, N. Walk, I. Roth, D. Markham, R. Parekh, U. Chabaud, and E. Kashefi, *Quantum certification and benchmarking*, *Nature Reviews Physics* **2**, 382 (2020), [arXiv:1910.06343 \[quant-ph\]](#).
- [2] M. Kliesch and I. Roth, *Theory of quantum system certification*, *PRX Quantum* **2**, 010201 (2021), tutorial, [arXiv:2010.05925 \[quant-ph\]](#).
- [3] J. Emerson, R. Alicki, and K. Życzkowski, *Scalable noise estimation with random unitary operators*, *J. Opt. B* **7**, S347 (2005), [arXiv:quant-ph/0503243](#).
- [4] E. Knill, D. Leibfried, R. Reichle, J. Britton, R. B. Blakestad, J. D. Jost, C. Langer, R. Ozeri, S. Seidelin, and D. J. Wineland, *Randomized benchmarking of quantum gates*, *Phys. Rev. A* **77**, 012307 (2008), [arXiv:0707.0963 \[quant-ph\]](#).
- [5] E. Magesan, J. M. Gambetta, and J. Emerson, *Characterizing quantum gates via randomized benchmarking*, *Phys. Rev. A* **85**, 042311 (2012), [arXiv:1109.6887](#).
- [6] J. Helsen, I. Roth, E. Onorati, A. H. Werner, and J. Eisert, *A general framework for randomized benchmarking*, [arXiv:2010.07974 \[quant-ph\]](#).
- [7] S. Kimmel, M. P. da Silva, C. A. Ryan, B. R. Johnson, and T. Ohki, *Robust extraction of tomographic information via randomized benchmarking*, *Phys. Rev. X* **4**, 011050 (2014), [arXiv:1306.2348 \[quant-ph\]](#).
- [8] S. Kimmel and Y. K. Liu, *Phase retrieval using unitary 2-designs*, in *2017 International Conference on Sampling Theory and Applications (SampTA)* (2017) pp. 345–349, [arXiv:1510.08887 \[quant-ph\]](#).
- [9] I. Roth, R. Kueng, S. Kimmel, Y. K. Liu, D. Gross, J. Eisert, and M. Kliesch, *Recovering quantum gates from few average gate fidelities*, *Phys. Rev. Lett.* **121**, 170502 (2018), [arXiv:1803.00572 \[quant-ph\]](#).
- [10] S. T. Flammia and J. J. Wallman, *Efficient estimation of Pauli channels*, *ACM Transactions on Quantum Computing* **1**, 1 (2020), [arXiv:1907.12976 \[quant-ph\]](#).
- [11] S. T. Merkel, J. M. Gambetta, J. A. Smolin, S. Poletto, A. D. Córcoles, B. R. Johnson, C. A. Ryan, and M. Steffen, *Self-consistent quantum process tomography*, *Phys. Rev. A* **87**, 062119 (2013), [arXiv:1211.0322 \[quant-ph\]](#).
- [12] C. Stark, *Self-consistent tomography of the state-measurement Gram matrix*, *Phys. Rev. A* **89**, 052109 (2014), [arXiv:1209.5737 \[quant-ph\]](#).
- [13] C. Stark, *Simultaneous estimation of dimension, states and measurements: Computation of representative density matrices and POVMs*, [arXiv:1210.1105 \[quant-ph\]](#).
- [14] R. Blume-Kohout, J. King Gamble, E. Nielsen, J. Mizrahi, J. D. Sterk, and P. Maunz, *Robust, self-consistent, closed-form tomography of quantum logic gates on a trapped ion qubit*, [arXiv:1310.4492 \[quant-ph\]](#).
- [15] R. Blume-Kohout, J. K. Gamble, E. Nielsen, K. Rudinger, J. Mizrahi, K. Fortier, and P. Maunz, *Demonstration of qubit operations below a rigorous fault tolerance threshold with gate set tomography*, *Nat. Commun.* **8**, 14485 (2017), [arXiv:1605.07674 \[quant-ph\]](#).
- [16] D. Greenbaum, *Introduction to quantum gate set tomography*, [arXiv:1509.02921 \[quant-ph\]](#) (2015).
- [17] E. Nielsen, K. Rudinger, T. Proctor, A. Russo, K. Young, and R. Blume-Kohout, *Probing quantum processor performance with pyGSTi*, *Quantum Sci. Technol.* **5**, 044002 (2020), [arXiv:2002.12476 \[quant-ph\]](#).
- [18] E. Nielsen, J. K. Gamble, K. Rudinger, T. Scholten, K. Young, and R. Blume-Kohout, *Gate set tomography*, *Quantum* **5**, 557 (2021), [arXiv:2009.07301 \[quant-ph\]](#).
- [19] J. P. Dehollain, J. T. Muhonen, R. Blume-Kohout, K. M. Rudinger, J. K. Gamble, E. Nielsen, A. Laucht, S. Simmons, R. Kalra, A. S. Dzurak, *et al.*, *Optimization of a solid-state electron spin qubit using gate set tomography*, *New Journal of Physics* **18**, 103018 (2016).
- [20] K. Rudinger, C. W. Hogle, R. K. Naik, A. Hashim, D. Lobser, D. I. Santiago, M. D. Grace, E. Nielsen, T. Proctor, S. Seritan, S. M. Clark, R. Blume-Kohout, I. Siddiqi, and K. C. Young, *Experimental Characterization of Crosstalk Errors with Simultaneous Gate Set Tomography*, *PRX Quantum* **2**, 040338 (2021), [arXiv:2103.09890 \[quant-ph\]](#).
- [21] C. Song, J. Cui, H. Wang, J. Hao, H. Feng, and Y. Li, *Quantum computation with universal error mitigation on a superconducting quantum processor*, *Science Advances* **5**, eaaw5686 (2019), [arXiv:1812.10903 \[quant-ph\]](#).
- [22] S. Zhang, Y. Lu, K. Zhang, W. Chen, Y. Li, J.-N. Zhang, and K. Kim, *Error-mitigated quantum gates exceeding physical fidelities in a trapped-ion system*, *Nature Communications* **11**, 587 (2020), [arXiv:1905.10135 \[quant-ph\]](#).
- [23] L. Cincio, K. Rudinger, M. Sarovar, and P. J. Coles, *Machine learning of noise-resilient quantum circuits*, [arXiv:2007.01210 \[quant-ph\]](#) (2020).
- [24] E. Knill, *Quantum computing with realistically noisy devices*, *Nature (London)* **434**, 39 (2005), [arXiv:quant-ph/0410199](#).
- [25] P. Aliferis, D. Gottesman, and J. Preskill, *Quantum accuracy threshold for concatenated distance-3 codes*, *Quant. Inf. Comput.* **6**, 97 (2006), [arXiv:quant-ph/0504218](#).
- [26] A. W. Cross, D. P. DiVincenzo, and B. M. Terhal,

- A comparative code study for quantum fault-tolerance*, Quant. Inf. Comp. **9**, 0541 (2009), [arXiv:0711.1556 \[quant-ph\]](#).
- [27] P. Aliferis and J. Preskill, *Fibonacci scheme for fault-tolerant quantum computation*, Phys. Rev. A **79**, 012332 (2009), [arXiv:0809.5063 \[quant-ph\]](#).
- [28] R. Kueng, D. M. Long, A. C. Doherty, and S. T. Flammia, *Comparing experiments to the fault-tolerance threshold*, Phys. Rev. Lett. **117**, 170502 (2016), [arXiv:1510.05653 \[quant-ph\]](#).
- [29] Y. R. Sanders, J. J. Wallman, and B. C. Sanders, *Bounding quantum gate error rate based on reported average fidelity*, New J. Phys. **18**, 012002 (2016), [arXiv:1501.04932 \[quant-ph\]](#).
- [30] J. J. Wallman, *Bounding experimental quantum error rates relative to fault-tolerant thresholds*, [arXiv:1511.00727 \[quant-ph\]](#) (2015).
- [31] A. Shabani, R. L. Kosut, M. Mohseni, H. Rabitz, M. A. Broome, M. P. Almeida, A. Fedrizzi, and A. G. White, *Efficient measurement of quantum dynamics via compressive sensing*, Phys. Rev. Lett. **106**, 100401 (2011), [arXiv:0910.5498 \[quant-ph\]](#).
- [32] S. T. Flammia, D. Gross, Y.-K. Liu, and J. Eisert, *Quantum tomography via compressed sensing: error bounds, sample complexity and efficient estimators*, New J. Phys. **14**, 095022 (2012), [arXiv:1205.2300 \[quant-ph\]](#).
- [33] M. Kliesch, R. Kueng, J. Eisert, and D. Gross, *Improving compressed sensing with the diamond norm*, IEEE Trans. Inf. Th. **62**, 7445 (2016), [arXiv:1511.01513 \[cs.IT\]](#).
- [34] M. Kliesch, R. Kueng, J. Eisert, and D. Gross, *Guaranteed recovery of quantum processes from few measurements*, Quantum **3**, 171 (2019), [arXiv:1701.03135 \[quant-ph\]](#).
- [35] C. H. Baldwin, A. Kalev, and I. H. Deutsch, *Quantum process tomography of unitary and near-unitary maps*, Phys. Rev. A **90**, 012110 (2014), [arXiv:1404.2877 \[quant-ph\]](#).
- [36] R. Kueng, H. Rauhut, and U. Terstiege, *Low rank matrix recovery from rank one measurements*, Appl. Comp. Harm. Anal. **10.1016/j.acha.2015.07.007** (2015), [arXiv:1410.6913 \[cs.IT\]](#).
- [37] A. V. Rodionov, A. Veitia, R. Barends, J. Kelly, D. Sank, J. Wenner, J. M. Martinis, R. L. Kosut, and A. N. Korotkov, *Compressed sensing quantum process tomography for superconducting quantum gates*, Phys. Rev. B **90**, 144504 (2014), [arXiv:1407.0761 \[quant-ph\]](#).
- [38] C. A. Riofrio, D. Gross, S. T. Flammia, T. Monz, D. Nigg, R. Blatt, and J. Eisert, *Experimental quantum compressed sensing for a seven-qubit system*, Nat. Commun. **8**, 15305 (2017), [arXiv:1608.02263 \[quant-ph\]](#).
- [39] A. Steffens, C. A. Riofrio, W. McCutcheon, I. Roth, B. A. Bell, A. McMillan, M. S. Tame, J. G. Rarity, and J. Eisert, *Experimentally exploring compressed sensing quantum tomography*, Quantum Sci. Technol. **2**, 025005 (2017), [arXiv:1611.01189 \[quant-ph\]](#).
- [40] A. Kyrillidis, A. Kalev, D. Park, S. Bhojanapalli, C. Caramanis, and S. Sanghavi, *Provable compressed sensing quantum state tomography via non-convex methods*, npj Quant. Inf. **4**, 36 (2018), [arXiv:1711.02524 \[quant-ph\]](#).
- [41] I. Roth, J. Wilkens, D. Hangleiter, and J. Eisert, *Semi-device-dependent blind quantum tomography*, [arXiv:2006.03069 \[quant-ph\]](#) (2020).
- [42] D. Gross, Y.-K. Liu, S. T. Flammia, S. Becker, and J. Eisert, *Quantum state tomography via compressed sensing*, Phys. Rev. Lett. **105**, 150401 (2010), [arXiv:0909.3304 \[quant-ph\]](#).
- [43] D. Gross, *Recovering low-rank matrices from few coefficients in any basis*, IEEE Trans. Inf. Th. **57**, 1548 (2011), [arXiv:0910.1879 \[cs.IT\]](#).
- [44] S. Foucart and H. Rauhut, *A mathematical introduction to compressive sensing* (Springer, 2013).
- [45] E. Nielsen, R. Blume-Kohout, L. Saldyt, J. Gross, T. L. Scholten, K. Rudinger, T. Proctor, J. K. Gamble, and A. Russo, *pygstio/pygsti: Version 0.9.9.3* (2020).
- [46] H.-Y. Huang, R. Kueng, and J. Preskill, *Predicting many properties of a quantum system from very few measurements*, Nature Physics **16**, 1050–1057 (2020), [arXiv:2002.08953 \[quant-ph\]](#).
- [47] T. Abrudan, J. Eriksson, and V. Koivunen, *Conjugate gradient algorithm for optimization under unitary matrix constraint*, Signal Processing **89**, 1704 (2009).
- [48] A. Edelman, T. A. Arias, and S. T. Smith, *The geometry of algorithms with orthogonality constraints*, SIAM Journal on Matrix Analysis and Applications **20**, 303 (1998), [arXiv:physics/9806030 \[physics.comp-ph\]](#).
- [49] J. H. Manton, *Optimization algorithms exploiting unitary constraints*, IEEE Trans. Signal Process. **50**, 635 (2002).
- [50] Y. Sun, N. Flammarion, and M. Fazel, *Escaping from saddle points on Riemannian manifolds*, in *Advances in Neural Information Processing Systems*, Vol. 32 (2019) [arXiv:1906.07355 \[math.OC\]](#).
- [51] M. A. d. A. Bortoloti, T. A. Fernandes, O. P. Ferreira, and J. Yuan, *Damped Newton's method on Riemannian manifolds*, Journal of Global Optimization **77**, 643 (2020), [arXiv:1803.05126 \[math.OC\]](#).
- [52] S. Wisdom, T. Powers, J. R. Hershey, J. Le Roux, and L. Atlas, *Full-capacity unitary recurrent neural networks*, in *Adv. Neural Inf. Process. Syst.*, Vol. 29 (2016) [arXiv:1611.00035 \[stat.ML\]](#).
- [53] N. Boumal, P.-A. Absil, and C. Cartis, *Global rates of convergence for nonconvex optimization on manifolds*, IMA Journal of Numerical Analysis **39**, 1 (2019), [arXiv:1605.08101 \[math.OC\]](#).
- [54] U. Helmke and J. B. Moore, *Optimization and dynamical systems* (Springer Science & Business Media, 2012).
- [55] J. Manton, R. Mahony, and Y. Hua, *The geometry of weighted low-rank approximations*, IEEE Transactions on Signal Processing **51**, 500 (2003).
- [56] L. Eldén and H. Park, *A procrustes problem on the stiefel manifold*, Numerische Mathematik **82**, 599 (1999).
- [57] T. J. Bridges and S. Reich, *Computing lyapunov exponents on a stiefel manifold*, Physica D: Nonlinear Phenomena **156**, 219 (2001).
- [58] F. Lotte, L. Bougrain, A. Cichocki, M. Clerc, M. Congedo, A. Rakotomamonjy, and F. Yger, *A review of classification algorithms for EEG-based brain-computer interfaces: a 10 year update*, Journal of Neural Engineering **15**, 031005 (2018).
- [59] E. Chiumiento and M. Melgaard, *Stiefel and grassmann manifolds in quantum chemistry*, Journal of Geometry and Physics **62**, 1866 (2012).
- [60] X. Yu, J.-C. Shen, J. Zhang, and K. B. Letaief, *Alternating minimization algorithms for hybrid precoding in millimeter wave mimo systems*, IEEE Journal of Se-

- lected Topics in Signal Processing **10**, 485 (2016).
- [61] F. Liu, C. Masouros, A. Li, H. Sun, and L. Hanzo, *Mu-mimo communications with mimo radar: From co-existence to joint transmission*, *IEEE Transactions on Wireless Communications* **17**, 2755 (2018).
 - [62] I. A. Luchnikov, M. E. Krechetov, and S. N. Filippov, *Riemannian geometry and automatic differentiation for optimization problems of quantum physics and quantum technologies*, *New J. Phys.* **23**, 073006 (2021), [arXiv:2007.01287 \[quant-ph\]](#).
 - [63] D. Hangleiter, I. Roth, J. Eisert, and P. Roushan, *Precise Hamiltonian identification of a superconducting quantum processor*, [arXiv:2108.08319 \[quant-ph\]](#).
 - [64] D. Hangleiter, I. Roth, D. Nagaj, and J. Eisert, *Easing the Monte Carlo sign problem*, *Science Advances* **6**, eabb8341 (2020), [arXiv:1906.02309 \[quant-ph\]](#).
 - [65] Y. N. Dauphin, R. Pascanu, C. Gulcehre, K. Cho, S. Ganguli, and Y. Bengio, *Identifying and attacking the saddle point problem in high-dimensional non-convex optimization*, in *Advances in Neural Information Processing Systems*, Vol. 27 (2014) [arXiv:1406.2572 \[cs.LG\]](#).
 - [66] D. Süß, *Due to, or in spite of? The effect of constraints on efficiency in quantum estimation problems*, Ph.D. thesis, University of Cologne (2018).
 - [67] M. Imaizumi, T. Maehara, and K. Hayashi, *On tensor train rank minimization : Statistical efficiency and scalable algorithm*, *Adv. Neural Inf. Process. Syst.* **30** (2017), [arXiv:1708.00132 \[stat.ML\]](#).
 - [68] H. Rauhut, R. Schneider, and Z. Stojanac, *Low rank tensor recovery via iterative hard thresholding*, *Linear Algebra and its Applications* **523**, 220 (2017), [arXiv:1602.05217 \[cs.IT\]](#).
 - [69] N. Ghadermarzy, Y. Plan, and Ö. Yilmaz, *Near-optimal sample complexity for convex tensor completion*, *Information and Inference: A Journal of the IMA* (2018), [arXiv:1711.04965 \[cs.LG\]](#).
 - [70] M. Ashraphijuo and X. Wang, *Characterization of deterministic and probabilistic sampling patterns for finite completability of low tensor-train rank tensor*, [arXiv:1703.07698 \[cs.LG\]](#).
 - [71] H. Rauhut and Ž. Stojanac, *Tensor theta norms and low rank recovery*, *Numer. Algor.* **88**, 25 (2021), [arXiv:1505.05175 \[cs.IT\]](#).
 - [72] B. Huang, C. Mu, D. Goldfarb, and J. Wright, *Provable low-rank tensor recovery*, *Optimization-Online* **4252**, 455 (2014).
 - [73] C. Liu, H. Shan, and C. Chen, *Tensor p -shrinkage nuclear norm for low-rank tensor completion*, *Neurocomputing* **387**, 255 (2020), [arXiv:1907.04092 \[cs.LG\]](#).
 - [74] S. Boixo, S. V. Isakov, V. N. Smelyanskiy, R. Babush, N. Ding, Z. Jiang, M. J. Bremner, J. M. Martinis, and H. Neven, *Characterizing quantum supremacy in near-term devices*, *Nature Physics* **14**, 595 (2018), [arXiv:1608.00263 \[quant-ph\]](#).
 - [75] J. Helsen, M. Ioannou, I. Roth, J. Kitzinger, E. Onorati, A. H. Werner, and J. Eisert, *Estimating gate-set properties from random sequences*, [arXiv:2110.13178 \[quant-ph\]](#) (2021).
 - [76] Y. Gu, R. Mishra, B.-G. Englert, and H. K. Ng, *Randomized linear gate-set tomography*, *PRX Quantum* **2**, 030328 (2021), [arXiv:2010.12235 \[quant-ph\]](#).
 - [77] T. J. Evans, W. Huang, J. Yoneda, R. Harper, T. Tanttu, K. W. Chan, F. E. Hudson, K. M. Itoh, A. Saraiva, C. H. Yang, A. S. Dzurak, and S. D. Bartlett, *Fast Bayesian tomography of a two-qubit gate set in silicon*, *Phys. Rev. Applied* **17**, 024068 (2022), [arXiv:2107.14473 \[quant-ph\]](#).
 - [78] S. Montangero, *Introduction to Tensor Network Methods* (Springer, 2018).
 - [79] M. Fannes, B. Nachtergaele, and R. Werner, *Finitely correlated states on quantum spin chains*, *Commun. Math. Phys.* **144**, 443 (1992).
 - [80] S. Rommer and S. Östlund, *Class of ansatz wave functions for one-dimensional spin systems and their relation to the density matrix renormalization group*, *Phys. Rev. B* **55**, 2164 (1997).
 - [81] F. Verstraete, D. Porras, and J. I. Cirac, *Density matrix renormalization group and periodic boundary conditions: A quantum information perspective*, *Phys. Rev. Lett.* **93**, 227205 (2004), [arXiv:cond-mat/0404706](#).
 - [82] I. V. Oseledets and E. E. Tyrtyshnikov, *Breaking the curse of dimensionality, or how to use SVD in many dimensions*, *SIAM Journal on Scientific Computing* **31**, 3744 (2009).
 - [83] K. Kraus, *States, effects and operations: fundamental notions of quantum theory*, Lecture notes in physics, Vol. 190 (Springer, Berlin, 1983).
 - [84] J. Lin, B. Buonacorsi, R. Laflamme, and J. J. Wallman, *On the freedom in representing quantum operations*, *New J. Phys.* **21**, 023006 (2019), [arXiv:1810.05631 \[quant-ph\]](#).
 - [85] L. Rudnicki, Z. Puchala, and K. Zyczkowski, *Gauge invariant information concerning quantum channels*, *Quantum* **2**, 60 (2018), [arXiv:1707.06926 \[quant-ph\]](#).
 - [86] T. Proctor, K. Rudinger, K. Young, M. Sarovar, and R. Blume-Kohout, *What randomized benchmarking actually measures*, *Phys. Rev. Lett.* **119**, 130502 (2017), [arXiv:1702.01853 \[quant-ph\]](#).
 - [87] U. Schollwöck, *The density-matrix renormalization group in the age of matrix product states*, *Ann. Phys.* **326**, 96 (2011), [arXiv:1008.3477 \[cond-mat.str-el\]](#).
 - [88] H. Schneider, *Positive operators and an inertia theorem*, *Numerische Mathematik* **7**, 11 (1965).
 - [89] P. Absil, R. Mahony, and R. Sepulchre, *Optimization algorithms on matrix manifolds* (Princeton University Press, 2009).
 - [90] L. Grasedyck, M. Kluge, and S. Krämer, *Variants of alternating least squares tensor completion in the tensor train format*, *SIAM J. Sci. Comput.* **37**, A2424 (2015), [arXiv:1509.00311 \[math.NA\]](#).
 - [91] W. Wang, V. Aggarwal, and S. Aeron, *Tensor completion by alternating minimization under the tensor train (TT) model*, [arXiv:1609.05587 \[cs.NA\]](#).
 - [92] C. Jin, P. Netrapalli, R. Ge, S. M. Kakade, and M. I. Jordan, *On nonconvex optimization for machine learning: Gradients, stochasticity, and saddle points*, *J. ACM* **68**, 1 (2021), [arXiv:1902.04811 \[cs.LG\]](#).
 - [93] R. Brieger, I. Roth, and M. Kliesch, *Python implementation of mGST, a compressive gate set tomography algorithm*, <https://github.com/rabrie/mGST> (2021).
 - [94] G. García-Pérez, M. A. C. Rossi, and S. Maniscalco, *IBM Q Experience as a versatile experimental testbed for simulating open quantum systems*, *npj Quantum Information* **6**, 1 (2020), [arXiv:1906.07099 \[quant-ph\]](#).
 - [95] R. Vershynin, *Introduction to the non-asymptotic analysis of random matrices*, in *Compressed Sensing: Theory and Applications* (Cambridge University Press, 2012) pp. 210–268, [arXiv:1011.3027 \[math.PR\]](#).

- [96] L. Grasedyck and S. Krämer, *Stable als approximation in the TT-format for rank-adaptive tensor completion*, *Numerische Mathematik* **143**, 855 (2019), [arXiv:1701.08045 \[math.NA\]](#).
- [97] A. Goeßmann, M. Götze, I. Roth, R. Sweke, G. Kuntyniok, and J. Eisert, *Tensor network approaches for learning non-linear dynamical laws*, in *NeurIPS 2020 – First Workshop on Quantum Tensor Networks in Machine Learning* (2020) [arXiv:2002.1238 \[math.NA\]](#).
- [98] Z. Cai, R. Babbush, S. C. Benjamin, S. Endo, W. J. Huggins, Y. Li, J. R. McClean, and T. E. O’Brien, *Quantum Error Mitigation*, arXiv e-prints (2022), [arXiv:2210.00921 \[quant-ph\]](#).
- [99] S. Endo, Z. Cai, S. C. Benjamin, and X. Yuan, *Hybrid Quantum-Classical Algorithms and Quantum Error Mitigation*, *Journal of the Physical Society of Japan* **90**, 032001 (2021), [arXiv:2011.01382 \[quant-ph\]](#).
- [100] D. E. Koh and S. Grewal, *Classical Shadows With Noise*, *Quantum* **6**, 776 (2022).
- [101] S. Chen, W. Yu, P. Zeng, and S. T. Flammia, *Robust Shadow Estimation*, *PRX Quantum* **2**, 030348 (2021), [arXiv:2011.09636 \[quant-ph\]](#).
- [102] A. Hjørungnes, *Complex-Valued Matrix Derivatives: With Applications in Signal Processing and Communications*, 1st ed. (Cambridge University Press, USA, 2011).
- [103] A. Van Den Bos, *Complex gradient and Hessian*, *IEE Proceedings-Vision, Image and Signal Processing* **141**, 380 (1994).
- [104] S. F. Waldron, *An introduction to finite tight frames* (Springer, 2018).

B Paper - Stability of classical shadows under gate-dependent noise

Title: Stability of classical shadows under gate-dependent noise

Authors: Raphael Brieger, Markus Heinrich, Ingo Roth, Martin Kliesch

Journal: Unpublished

Publication status: Under revision at Physical Review Letters

Contribution by RB: First author (input approx 65%)

A summary of this publication is presented in Chapter 4.

The idea for the project came through joint discussions between all authors about noise amplification due to dimensional factors in the classical shadows protocol. I derived the main theorems and most of the additional results in the Appendix, with the exception of Proposition 14 and half of Appendix B, which were derived by MH. MH further provided key ideas for the derivation of the variance bounds in Appendix D. I wrote an initial draft for the main text and most of the Appendix. The draft of the main text was then extended by my co-authors and finally rewritten in joint video calls by all authors. IR, MK and MH further gave continuous advice on the clear presentation of results, while MH additionally improved the notation throughout the manuscript.

Stability of classical shadows under gate-dependent noise

Raphael Brieger,^{1,2,*} Markus Heinrich,¹ Ingo Roth,³ and Martin Kliesch²

¹*Heinrich Heine University Düsseldorf, Faculty of Mathematics and Natural Sciences, Germany*

²*Hamburg University of Technology, Institute for Quantum Inspired and Quantum Optimization, Germany*

³*Quantum Research Center, Technology Innovation Institute, Abu Dhabi, UAE*

Expectation values of observables are routinely estimated using so-called classical shadows—the outcomes of randomized bases measurements on a repeatedly prepared quantum state. In order to trust the accuracy of shadow estimation in practice, it is crucial to understand the behavior of the estimators under realistic noise. In this work, we prove that any shadow estimation protocol involving Clifford unitaries is stable under *gate-dependent* noise for observables with bounded *stabilizer norm*—originally introduced in the context of simulating Clifford circuits. For these observables, we also show that the protocol’s sample complexity is essentially identical to the noiseless case. In contrast, we demonstrate that estimation of ‘magic’ observables can suffer from a bias that scales exponentially in the system size. We further find that so-called *robust shadows*, aiming at mitigating noise, can introduce a large bias in the presence of gate-dependent noise compared to unmitigated classical shadows. Nevertheless, we guarantee the functioning of robust shadows for a more general noise setting than in previous works. On a technical level, we identify average noise channels that affect shadow estimators and allow for a more fine-grained control of noise-induced biases.

I. INTRODUCTION

Efficient estimation of observables is crucial for quantum experiments and devices. Classical shadows [1] utilize measurements in randomized bases to perform many relevant estimation tasks on states that are repeatedly prepared in an experiment. A key feature is that one can choose the observables after the collection of the measurement data and only adapt the classical post-processing accordingly. For this reason, the approach is highly flexible and has many applications [2–7].

The central step in the estimation is to invert ‘the overall measurement process’ classically. This inversion can be calculated analytically for measurement bases that are uniformly random, either local or global, Clifford rotations of the computational bases. For these cases, tight sampling complexity bounds in terms of certain norms of the observables have been derived [1]. Most extensions of this paradigm still involve Clifford gates, such as random Clifford circuits [8–10] or Clifford matchgates [11]. In these cases, the inversion is performed using a combination of analytical and numerical techniques.

These calculations crucially rely on the assumption that the unitary gates that perform the basis rotations are perfectly implemented on the quantum device—an assumption that inevitably needs to be relaxed when using classical shadows for precise estimation in practice. This has motivated several works [12–17] studying the noise robustness of classical shadows. Using a restricted

noise model, it has been shown that the effect of noise on the estimator can be either estimated independently, e.g. by a separate calibration experiment [12] or inferred using symmetries of the prepared state [16]. Once the effect is known, it can then be mitigated in post-processing. The derivations of these robust classical shadows assume that the noise in the system is described by the *same* channel acting directly before the measurement in each round. This gate-independent noise model is well-suited to capture the effect of read-out noise affecting the computational basis measurement. However, it is difficult to justify this model for gate noise in realistic experimental setups. Gate-dependent noise models can significantly complicate the mitigation [18]. It is an open question of how stable classical shadows and their robust extensions are under gate-dependent noise. As a matter of concern, the inversion of the effective measurement process typically involves factors that scale exponentially with the system size (or locality of the observable). Thus, even small errors in the gates could in large biases in the estimators. We give an explicit example where this is indeed the case. This raises serious doubts about the accuracy of shadow estimation in practice, especially when the measurement bases require entangling gates.

In this work, we prove that shadow estimation with Clifford circuits is intrinsically stable under gate-dependent noise for observables with bounded *stabilizer norm*. This includes stabilizer states, Pauli observables, and large classes of linear combinations thereof. The *stabilizer norm* [19], also known as $\frac{1}{2}$ -*stabilizer Rényi entropy* [20], is a well-known resource measure in the resource theory of magic states and can be used to bound the runtime of classical stabilizer-based simulation

* raphael.brieger@hhu.de

methods [21]. We formally capture the effect of gate-dependent noise by identifying noise channels averaged over unions of cosets of Clifford subgroups. This allows us to obtain a more fine-grained control of the estimation bias in terms of the stabilizer norm. For uniform sampling from the local and global group, we further show that the sampling complexity of shadow estimation is also stable under gate-dependent noise for the same class of observables.

For the robust classical shadows of Ref. [12], we find that the average noise channels can conspire to cause a larger bias in the robust estimator than in the simple direct estimator. The bias of the robust estimator can even scale exponentially with the system size. Taking a closer look at the necessary structure for this situation to appear, we can further show the stability of the robust shadow estimator for a strictly more general gate-dependent noise model than in Ref. [12], which we call *isotropic Pauli noise*.

II. STABILITY OF SHADOW ESTIMATION

The goal of a general shadow estimation protocol is to estimate expectation values of observables in the same unknown n -qubit quantum state ρ . To this end, the protocol applies a randomly drawn unitary g from a set G with probability $p(g)$ to an unknown quantum state ρ and measures in the computational basis measurement, resulting in some output $x \in \mathbb{F}_2^n$. For an observable O , one can then evaluate a function $\hat{o}(g, x)$ defining an unbiased estimator for the expectation value $\mathbb{E}[\hat{o}] = \text{Tr}(O\rho)$. For this to work for any observable O , the operators $\{g|x\rangle\langle x|g^\dagger\}_{g,x}$ have to form an informationally complete positive operator valued measure (POVM).

To give more details, we first introduce some notation. We define $E_x := |x\rangle\langle x|$ and use $\omega(g)$ to denote the unitary channel $\omega(g)(A) = gAg^\dagger$. Round brackets are used to denote inner and outer products of linear operators in analogy to the usual Dirac notation. In particular, $(A|B) = \text{Tr}(A^\dagger B)$ denotes the Hilbert Schmidt inner product and $|A)(B|$ is the superoperator $C \mapsto (B|C)A$. With this notation, we define the measure-and-prepare channel $M := \sum_{x \in \mathbb{F}_2^n} |E_x)(E_x|$ in the computational basis. Being an informationally complete POVM, the operators $\{\omega(g)(E_x)\}_{g,x}$ are a frame, i.e. a spanning set for the vector space of linear operators, hence the associated *frame operator*

$$S := \mathbb{E}_{g \sim p} [\omega(g)^\dagger M \omega(g)] \quad (1)$$

is invertible. With this notation, we define the estimator $\hat{o}(g, x) := (O|S^{-1}\omega(g)^\dagger|E_x)$ of a given observable O , and a straightforward calculation shows that, indeed, $\mathbb{E}[\hat{o}] = (O|S^{-1}S|\rho) = (O|\rho) = \text{Tr}(O\rho)$.

The frame operator S can often be analytically calculated and inverted for uniform sampling from certain

subgroups $G \subset \text{U}(d)$ where $d = 2^n$. A prominent example for such a subgroup is the *Clifford group* Cl_n , defined as the subgroup of $\text{U}(d)$ that is generated by the Hadamard gate, the phase gate, and the controlled-NOT gate. Since Cl_n is a unitary 2-design, its frame operator is proportional to the identity on the subspace of traceless matrices and has eigenvalue $\frac{1}{d+1}$ [1].

Gates in an actual experiment, however, suffer from noise and imperfections. A fairly general and common noise model replaces $\omega(g)$ by its noisy implementation $\phi(g) = \omega(g)\Lambda(g)$ where $\Lambda(g)$ is an arbitrary noise channel that depends on g . Note that introducing an additional noise channel on the left of ω is equivalent to our model. The existence of such an implementation map ϕ requires the noise to be *Markovian* and *time-stationary*, but it can be otherwise arbitrary. The *noisy frame operator* of a shadow estimation protocol is then given by

$$\tilde{S} := \mathbb{E}_{g \sim p} [\omega(g)^\dagger M \omega(g) \Lambda(g)]. \quad (2)$$

In the presence of noise, the standard shadow estimator is biased. This can be readily seen for a traceless observable O_0 and uniform sampling from the Clifford group Cl_n , for which $S^{-1}(O_0) = (d+1)O_0$. The expected value then reads:

$$\mathbb{E}[\hat{o}_0] = (O_0|S^{-1}\tilde{S}|\rho) = (O_0|\rho) + (d+1)(O_0|S - \tilde{S}|\rho).$$

Due to the exponentially large factor $d+1$ applied in post-processing, one runs the risk of dramatically amplifying errors. In particular, a first straightforward attempt at controlling the noise-induced bias yields a bound of the following form (c.f. Appendix A):

$$|\mathbb{E}[\hat{o}] - \langle O \rangle| \leq (d+1) \max_{g \in G} \|\text{id} - \Lambda(g)\|_\diamond. \quad (3)$$

Here, we quantify the error of the implementation by the maximum diamond distance of the noise channel to the identity channel over all gates. Equation (3) is the first example of a bound controlling the bias of shadow estimation. These bounds take the general form of

$$\text{bias} \leq \kappa(d) \times \text{implementation error}.$$

We say that the estimation is *stable* if scaling function κ is constant, i.e. $\kappa \in \mathcal{O}(1)$. For using shadow estimation in practice noise stability is an essential requirement. However, Eq. (3) suggests that shadow estimation can in fact be *unstable*. Indeed, we can give the following example:

Proposition 1. *Let $O = (|H\rangle\langle H|)^{\otimes n}$ with the magic state $|H\rangle = \frac{1}{\sqrt{2}}(|0\rangle + e^{i\pi/4}|1\rangle)$ and consider shadow estimation with local Clifford unitaries. There exists a state ρ and implementation map $\phi_\epsilon(g) = (1 - \epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g)$ such that $|\mathbb{E}[\hat{o}] - \langle O \rangle| = \kappa\epsilon$ with $\kappa \in \Omega(d^{1/4})$.*

We prove the proposition in Appendix A by an explicit construction. The construction uses noise channels $\Lambda(g)$ that are coherently undoing the basis change of $\omega(g)$.

This result brings us to the central question of our work: *Are there any classes of stable shadow estimation settings?* Perhaps surprisingly, we can answer this question positively for large classes of observables in the case that all unitaries are taken from the Clifford group.

In this setting, a careful analysis of the noisy frame operator (2) allows us to improve the naive estimates significantly. To this end, it is convenient to work in the orthogonal operator basis given by the *Pauli operators* $\sigma_a = \sigma_{a_1} \otimes \cdots \otimes \sigma_{a_n}$, i.e. $(\sigma_a | \sigma_{a'}) = d \delta_{a,a'}$. We choose to label them by binary vectors $a \in \mathbb{F}_2^{2n}$ with the convention $\sigma_{00} = \mathbb{1}$, $\sigma_{01} = X$, $\sigma_{11} = Y$, $\sigma_{10} = Z$. The following technical result on the form of noisy frame operator $\tilde{S}$ then serves as a basis for our main results.

Lemma 2. *Suppose $G \subset \text{Cl}_n$. The noisy frame operator (2) takes the form*

$$\tilde{S} = \frac{1}{d} \sum_{a \in \mathbb{F}_2^{2n}} s_a |\sigma_a)(\sigma_a | \bar{\Lambda}_a, \quad (4)$$

where $s_a \in (0, 1]$ are the eigenvalues of the noise-free frame operator S and $\bar{\Lambda}_a$ are quantum channels depending on Λ . Furthermore, if the $\Lambda(g)$ are Pauli noise channels, then $\bar{\Lambda}_a | \sigma_a) = \bar{\lambda}_a | \sigma_a)$ where $\bar{\lambda}_a \in [-1, 1]$.

A proof of the lemma is given in Appendix A. The channels $\bar{\Lambda}_a$ are averages of gate-dependent noise channels $\Lambda(g)$ for g sampled from specific subsets of G and approach the identity for weak noise. Their form for the global and local Clifford groups is further discussed in Appendix B. In analogy to the channel averages $\bar{\Lambda}_a$, the parameters $\bar{\lambda}_a$ are averaged Pauli eigenvalues of individual Pauli noise channels $\Lambda(g)$. The proofs builds on the simple observation that in Eq. (2), each term $\omega^\dagger(g) M \omega(g)$ is always a sum over Pauli projectors $d^{-1} | \sigma_a)(\sigma_a |$, and noise channels $\Lambda(g)$ associated to the same projector can be grouped together.

Our main stability result uses the *stabilizer norm* of the observable O [19], defined as

$$\|O\|_{\text{st}} := \frac{1}{d} \sum_{a \in \mathbb{F}_2^{2n}} |(\sigma_a | O)|, \quad (5)$$

which is proportional to the ℓ_1 -norm of an operator in the Pauli basis. We will also use the *Hilbert-Schmidt norm* $\|O\|_2 := \sqrt{\langle O | O \rangle}$ of an operator O . We now provide a proof that shadow estimation of observables with constant stabilizer norm is stable against gate-dependent noise.

Theorem 3. *Suppose $G \subset \text{Cl}_n$. The estimation bias is bounded by*

$$\begin{aligned} |\mathbb{E}[\hat{o}] - \langle O \rangle| &\leq \|O\|_{\text{st}} \max_{a \in \mathbb{F}_2^{2n}} \|\text{id} - \bar{\Lambda}_a\|_\diamond \\ &\leq \|O\|_{\text{st}} \max_{g \in \text{Cl}_n} \|\text{id} - \Lambda(g)\|_\diamond \end{aligned} \quad (6)$$

for arbitrary gate-dependent noise and

$$|\mathbb{E}[\hat{o}] - \langle O \rangle| \leq \min \{ \|O\|_2, \|O\|_{\text{st}} \} \max_{a \in \mathbb{F}_2^{2n}} |1 - \bar{\lambda}_a|$$

for gate-dependent Pauli noise.

Proof. From Lemma 2 we immediately obtain the noise-free frame operator $S = d^{-1} \sum_a s_a |\sigma_a)(\sigma_a |$ by setting $\bar{\Lambda}_a = \text{id}$ for all a . It then follows that $S^{-1} \tilde{S} = d^{-1} \sum_a |\sigma_a)(\sigma_a | \bar{\Lambda}_a$. We expand O in the Pauli basis as $|O\rangle = d^{-1} \sum_{a \in \mathbb{F}_2^{2n}} (O | \sigma_a) (\sigma_a |$. In the following, $\|\cdot\|_1$ and $\|\cdot\|_\infty$ denote the nuclear and spectral norm, respectively. The bias $|\langle O \rangle - (O | S^{-1} \tilde{S} | \rho)| = |(O | (\text{id} - S^{-1} \tilde{S}) | \rho)|$ can then be bounded as follows:

$$\begin{aligned} |(O | (\text{id} - S^{-1} \tilde{S}) | \rho)| &= \frac{1}{d} \left| \sum_a (O | \sigma_a) (\sigma_a | (\text{id} - \bar{\Lambda}_a) | \rho) \right| \\ &\leq \frac{1}{d} \sum_a |(O | \sigma_a)| |(\sigma_a | (\text{id} - \bar{\Lambda}_a) | \rho)| \\ &\leq \frac{1}{d} \sum_a |(O | \sigma_a)| \|(\text{id} - \bar{\Lambda}_a) | \rho)\|_1 \\ &\leq \max_a \|\text{id} - \bar{\Lambda}_a\|_\diamond \times \frac{1}{d} \sum_a |(O | \sigma_a)|, \end{aligned}$$

where we used triangle and matrix Hölder inequality, as well as $\|\sigma_a\|_\infty = 1$ and $\|(\text{id} - \bar{\Lambda}_a) | \rho)\|_1 \leq \|\text{id} - \bar{\Lambda}_a\|_\diamond$. The result (6) then follows from the definition of the stabilizer norm. For Pauli noise, Lemma 2 implies that $S^{-1} \tilde{S} = d^{-1} \sum_a |\sigma_a)(\sigma_a | \bar{\lambda}_a$. The bias then becomes

$$\begin{aligned} |\langle O \rangle - (O | S^{-1} \tilde{S} | \rho)| &= \frac{1}{d} |(O | \sum_a |\sigma_a)(\sigma_a | (1 - \bar{\lambda}_a) | \rho)| \\ &\leq \max_a |1 - \bar{\lambda}_a| \frac{1}{d} \sum_{a \neq 0} |(O | \sigma_a) (\sigma_a | \rho)|. \end{aligned}$$

We can bound the last line by either $\max_a |1 - \bar{\lambda}_a| \|O\|_2 \|\rho\|_2 \leq \max_a |1 - \bar{\lambda}_a| \|O\|_2$ via the Cauchy-Schwarz inequality, or by $|(O | \sigma_a | \rho)| \leq 1$ for all a and all states ρ . The latter leads to $d^{-1} \sum_{a \neq 0} |(O | \sigma_a) (\sigma_a | \rho)| \leq d^{-1} \sum_{a \neq 0} (O | \sigma_a) = \|O\|_{\text{st}}$. $\square$

Theorem 3 holds for any informationally complete shadow estimation protocol based on Clifford gates, including uniform sampling from the global or local Clifford group, as well as for alternative proposals such as brickwork circuits [9, 10]. In comparison to Eq. (3), the scaling factor $d + 1$ is replaced by the stabilizer norm $\|O\|_{\text{st}}$. The latter is a magic measure [19], equivalent to the $\frac{1}{2}$ -stabilizer Rényi entropy [20], and can be understood as a quantification of the ‘non-stabilizerness’ of the observable O . In particular, $\|O\|_{\text{st}} = \mathcal{O}(1)$ for many interesting examples such as Pauli observables or stabilizer states. Furthermore, $\|O\|_{\text{st}}$ is also well-behaved for non-stabilizer observables, as long as their Pauli support (and coefficients) is not too large. For instance, consider a k -local Hamiltonian $\sum_{e \in E} h_e$ on a hypergraph (V, E) with maximum degree D , i.e. each $h_e = \sum_a h_e^a \sigma_a$ is a linear combination of at most 3^k Pauli matrices supported on the hyperedge e . Thus, the number of local terms is

$|E| \leq |V|D = nD$, and assuming $|h_e^a| \leq 1$, the stabilizer norm satisfies $\|H\|_{\text{st}} \leq nD3^k$. Finally, the stabilizer norm is multiplicative under tensor products and, hence, simple to compute for product observables, but may be harder to evaluate for more general O . To this end, it might be worth mentioning that $\|O\|_{\text{st}}$ can be upper bounded by higher stabilizer Rényi entropies [20] and the so-called robustness of magic [22].

Interestingly, the classical post-processing of shadow estimation involves the evaluation of $\hat{o}(g, x) = (O|S^{-1}\omega(g)^\dagger|E_x)$. This evaluation can be a computationally hard problem, even for uniform sampling from the (local or global) Clifford group where the difficult bit reduces to evaluating expectation values on stabilizer states. This computational task is well-studied and can be solved in classical runtime $\mathcal{O}(\|O\|_{\text{st}}^2)$ [21, 23]. Therefore, a compelling observation is that *observables with bounded magic allow for both efficient classical post-processing and stable estimation*.

Compared to the arbitrary noise case, the Pauli noise bound in Theorem 3 is both stronger in the error measure and in the dependence on the observable. Concerning the latter, the bound by $\|O\|_2$ is favorable when O is a quantum state or an entanglement witness, since then $\|O\|_2 \leq 1$. The error scaling for Pauli noise is also strictly better which can be seen as follows: $1 - \bar{\lambda}_a$ is an eigenvalue of the diagonal superoperator $\text{id} - \bar{\Lambda}_a$, hence $|1 - \bar{\lambda}_a| \leq \|\text{id} - \bar{\Lambda}_a\|_\infty$. Since $\bar{\Lambda}_a$ is a Pauli channel, we have $\|\text{id} - \bar{\Lambda}_a\|_\infty \leq \|\text{id} - \bar{\Lambda}_a\|_\diamond$ (see Lemma 12 in Appendix C) and thus $\max_a |1 - \bar{\lambda}_a| \leq \max_a \|\text{id} - \bar{\Lambda}_a\|_\diamond$. To ensure that noise is well described by Pauli channels, randomized compiling [24–26] can be used for each gate in the shadow estimation procedure.

Finally, another question that arises in regard to Theorem 3 is whether the given bounds are tight, especially whether the error can really scale with $\|O\|_{\text{st}}$. If so, we would essentially recover our naive bound in Eq. (3) for highly magic observables with $\|O\|_{\text{st}} = 2^{\mathcal{O}(n)}$. In fact, the already discussed example in Proposition 1 also saturates the bound in Theorem 3 since $\|H\rangle\langle H|^{\otimes n}\|_{\text{st}} = \|H\rangle\langle H\|_{\text{st}}^n = \left(\frac{1+\sqrt{2}}{2}\right)^n \geq 2^{n/4}$ (c.f. Appendix F).

A remaining open question is the effect of gate-dependent noise on the protocol’s sample complexity. A potential instability of the latter can also render shadow estimation infeasible in practice. The required number of samples can be bounded using the variance of the estimator $\mathbb{V}[\hat{o}] = \mathbb{E}[\hat{o}^2] - \mathbb{E}[\hat{o}]^2$. In the presence of gate-dependent noise, we show that the second moment $\mathbb{E}[\hat{o}^2]$ can be written in terms of Pauli projectors and average noise channels, akin to Lemma 2. We find the following explicit variance bound for the local and global Clifford group:

Theorem 4. *The variance of shadow estimation for uniform sampling from the global Clifford group under gate-dependent noise is bounded by*

$$\mathbb{V}_{\text{global}}[\hat{o}] \leq \frac{2(d+1)}{(d+2)} \|O_0\|_{\text{st}}^2 + \frac{d+1}{d} \|O_0\|_2^2.$$

For uniform sampling from the local Clifford group, the variance is bounded by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 4^k \|O_{\text{loc}}\|_\infty^2$ for k -local observables $O = O_{\text{loc}} \otimes \mathbf{1}^{n-k}$ and by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 3^{\text{supp}(\sigma_a)}$ for Pauli observables $O = \sigma_a$.

The proof thereof, and more details are given in Appendix D. For global Cliffords, the variance in Theorem 4 is governed by the squared stabilizer norm $\|O_0\|_{\text{st}}^2$ and the squared Hilbert-Schmidt norm $\|O_0\|_2^2$. Let us compare our result to the noise-free result in Ref. [1]. Therein, the variance bound has a similar form, where however the stabilizer norm is replaced by the spectral norm, leading to a stronger bound since $\|O\|_\infty \leq \|O\|_{\text{st}}$. In the presence of noise, we can only recover a similar result for observables with $\|O\|_{\text{st}} = \mathcal{O}(\|O\|_\infty)$. Our findings indicate that the sample complexity of the protocol for highly magical observables with $\|O\|_{\text{st}} = 2^{\mathcal{O}(n)}$ may no longer be controlled. Intriguingly, we exactly recover the noise-free variance bound of Ref. [1] for the local Clifford group, meaning that the sample complexity for many important use cases is not adversely affected by gate-dependent noise. Finally, for more general sampling schemes, we expect that variance results from the noiseless case can be extended to gate-dependent noise, but may again scale with the observable’s stabilizer norm.

III. BIAS MITIGATION

Modifications to the original shadow estimation protocol, coined *robust shadow estimation (RSE)* [12], have already been proposed with the goal of mitigating noise-induced biases. Their performance guarantees rely on the assumption of gate-independent (left) noise and their success under more general noise models is unclear. Intuitively, one may hope that RSE mitigates the ‘dominant contributions’ of an otherwise complicated noise model and therefore generally improves the shadow estimate. We demonstrate in the following that RSE can however *increase the bias rather than reduce it* when the noise model assumption is violated. In the worst case, gate errors can even be amplified exponentially in the number of qubits.

Proposition 5. *Under gate-dependent local noise $\phi_\epsilon(g) = (1 - \epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g)$, robust shadow estimation with the global Clifford group can introduce a bias $|\mathbb{E}[\hat{o}] - \langle O \rangle| \geq |\langle O_0 \rangle|(\frac{1}{2}(1 + \epsilon)^n - 1)|$.*

As we show, this situation generally requires the effect of the noise to be aligned with the support of the observable or the state under scrutiny. More precisely, we prove Proposition 5 in Appendix E by constructing an explicit noise model consisting of local bit-flip errors.

In numerical simulations [12, 16], simple gate-dependent noise models are found to be sufficiently well-behaved in order for RSE to reduce the estimation bias. To explain these findings, we investigate the influence of gate-dependent Pauli noise on RSE in more detail. The

general form of the frame operator given in Lemma 2 allows us to derive more precise conditions on the applicability of RSE and to lay the groundwork for a better understanding of related mitigation approaches, see Appendix F. In particular, we identify a strictly more general noise model than gate-independent left noise [27] for which RSE works as intended and the bias is strongly suppressed in the number of qubits. We call this gate-dependent noise model *isotropic Pauli noise*. Here, the average Pauli eigenvalues $\bar{\lambda}_a$ are allowed to fluctuate randomly around a mean value in a rotation-invariant fashion in the space of eigenvalues. This model is motivated by the intuition that complicated noise processes and the effective averaging introduced by shadow estimation can be well approximated by normally distributed eigenvalues. We formulate this finding as follows.

Proposition 6. *The bias of robust shadow estimation is strongly suppressed in the number of qubits for isotropic Pauli noise.*

If there is reason to believe that the effective noise is not well approximated by isotropic Pauli noise, unmitigated shadow estimation should be trusted over RSE, due to its controlled bias (Theorem 3).

IV. CONCLUSION AND OUTLOOK

In order to trust the accuracy of shadow estimation in practice, it is vital to understand its behavior under realistic noise assumptions. Indeed, due to dimensional factors in the estimators, small gate imperfections can result in large biases of the estimator. Overcoming the restrictions of previous work, we develop a theory for classical shadows based on Clifford unitaries under general gate-dependent, time-stationary, and Markovian noise. In particular, our results apply to global and local Clifford unitaries as well as random Clifford circuits and matchgate 3-designs. We find that shadow estimation is, perhaps surprisingly, robust for large classes of observables, including convex combinations of Pauli observables and stabilizer states. More precisely, we show that the estimation bias scales with the strength of the gate noise, as well as the observable’s *stabilizer norm*. Thus, a bounded stabilizer norm guarantees noise-robust estimates of expectation values. Intriguingly, this is also the class of observables for which one may expect that the classical post-processing in shadow estimation can be done efficiently using stabilizer techniques. In contrast, highly ‘magic’ observables can lead to a strong, even exponential, amplification of noise errors as we illustrate at the example of a magic state. Furthermore, we derive explicit variance bounds for the local and global Clifford groups which guarantee sample efficiency for the same observables as in the noiseless case. For global Cliffords, we again require that the observable’s stabilizer norm is bounded, while the variance for local Cliffords is unchanged under noise. Finally, we show that the intrinsic

stability of classical shadows under gate-dependent noise can even outperform ‘robust classical shadows’; a noise mitigation scheme which relies on a more restricted noise model. Nevertheless, we extend the regime in which robust shadows work reliably to certain gate-dependent noise models, which we call isotropic Pauli noise.

We regard our work as an important step towards understanding the performance of randomized protocols with Clifford unitaries under gate-dependent noise. In particular, our results can provide guidance in devising further approaches to mitigate noise in classical shadows, provide crucial justification and identify caveats.

ACKNOWLEDGMENTS

This work was funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) via the Emmy Noether program (grant no. 441423094), the German Federal Ministry of Education and Research (BMBF) within the funding program “Quantum technologies – from basic research to market” in the joint project MIQRO (grant no. 13N15522), and by the Fujitsu Services GmbH as part of the endowed professorship “Quantum Inspired and Quantum Optimization”.

Appendix A: The frame operator under gate-dependent right noise

We draw gates from a set G and aim for the target implementation $\omega(g)(\rho) = g\rho g^\dagger$ acting as unitaries $g \in \text{U}(d)$ on density operators ρ . We assume that there is a corresponding hardware implementation $\phi : G \rightarrow \text{CPT}(\mathcal{H})$ that takes each gate g to some completely positive and trace preserving (CPT) map $\phi(g)$ on $\mathcal{H}$.

Proposition 7 (Direct worst-case error bound). *Consider an observable O_0 that, w.l.o.g., we assume to be traceless. The bias of shadow estimation with uniform sampling from a unitary 2-design is bounded as*

$$\left| (O_0|S^{-1}\tilde{S}|\rho) - (O_0|\rho) \right| \leq (d+1) \mathbb{E}_{g \in G} \|\phi(g) - \omega(g)\|_\diamond. \quad (\text{A1})$$

Proof. For uniform sampling from a unitary 2-design (such as the Clifford group), the inverse of the ideal frame operator acts as $(O_0|S^{-1} = (d+1)(O_0|$ [1]. Let μ be the uniform probability measure on G , for instance the normalized Haar measure on $\text{U}(d)$ or Cl_n . The error due to noise in the frame operator is then given by

$$\left| (O_0|S^{-1}\tilde{S}|\rho) - (O_0|S^{-1}S|\rho) \right| = (d+1)(O_0| \left[\int \omega(g)^\dagger M \phi(g) d\mu(g) - \int \omega(g)^\dagger M \omega(g) d\mu(g) \right] |\rho) \quad (\text{A2})$$

$$\leq (d+1) \left\| \int \omega(g)^\dagger M \phi(g) d\mu(g) - \int \omega(g)^\dagger M \omega(g) d\mu(g) \right\|_\diamond \quad (\text{A3})$$

$$= (d+1) \left\| \int \omega(g)^\dagger M (\phi(g) - \omega(g)) d\mu(g) \right\|_\diamond \quad (\text{A4})$$

$$\leq (d+1) \int \|\omega(g)^\dagger M\|_\diamond \|\phi(g) - \omega(g)\|_\diamond d\mu(g) \quad (\text{A5})$$

$$= (d+1) \int \|\phi(g) - \omega(g)\|_\diamond d\mu(g), \quad (\text{A6})$$

where we used the submultiplicativity of the diamond norm and that $\|\mathcal{C}\|_\diamond = 1$ for any quantum channel $\mathcal{C}$, in particular for $\mathcal{C} = M$. $\square$

For trace preserving $\phi(g)$, one can quickly verify that $|(O|S^{-1}\tilde{S}|\rho) - (O|\rho)| = |(O_0|S^{-1}\tilde{S}|\rho) - (O_0|\rho)|$, where O_0 is the traceless component of O . This bound suggests an error amplification by a dimensional factor $d+1$. However, the bound cannot be tight since the triangle inequality from Eq. (A4) to Eq. (A5) can only be saturated in the trivial zero-error case $\phi(g) = \omega(g)$. Indeed, in the following, we show that much stronger bounds for sampling measurement unitaries from the Clifford group can be derived. The defining property of Clifford unitaries that they map the Pauli group onto itself plays a central role: even in the presence of noise, the frame operator can be written in a more amenable form when looking at its Pauli transfer matrix.

In the following, we denote by P_n the set of n -qubit Pauli operators. The action of Clifford group elements on M via the representation ω is given by

$$\omega(g)(\sigma_a) = (-1)^{\varphi_a(g)} \Xi_a(g), \quad (\text{A7})$$

with functions $\varphi_a : \text{Cl}_n \rightarrow \mathbb{F}_2$ and $\Xi_a : \text{Cl}_n \rightarrow \text{P}_n$ defined for $a \in \mathbb{F}_2^{2n}$. Moreover, we write $\check{\sigma}_a = \sigma_a/\sqrt{d}$ for the normalized Pauli operators such that $(\check{\sigma}_a|\check{\sigma}_b) = \delta_{a,b}$. In the following, we denote by $Z_z \equiv \bigotimes_{i=1}^n Z^{z_i}$ with $z \in \mathbb{F}_2^n$ the diagonal Pauli operators and set $Z_1 \equiv Z_{e_1}$ with $e_1 = (1, 0, \dots, 0)$. We write again $\check{Z}_z = Z_z/\sqrt{d}$ for their normalized versions. Moreover, one can easily verify that M is given in terms of the diagonal Paulis as

$$M = \sum_{x \in \mathbb{F}_2^n} |E_x)(E_x| = \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z)(\check{Z}_z|. \quad (\text{A8})$$

Lemma 2. *Suppose $G \subset \text{Cl}_n$. The noisy frame operator (2) takes the form*

$$\tilde{S} = \frac{1}{d} \sum_{a \in \mathbb{F}_2^{2n}} s_a |\sigma_a)(\sigma_a| \bar{\Lambda}_a, \quad (4)$$

where $s_a \in (0, 1]$ are the eigenvalues of the noise-free frame operator S and $\bar{\Lambda}_a$ are quantum channels depending on Λ . Furthermore, if the $\Lambda(g)$ are Pauli noise channels, then $\bar{\Lambda}_a|\sigma_a) = \bar{\lambda}_a|\sigma_a)$ where $\bar{\lambda}_a \in [-1, 1]$.

Proof. We show the statement by a direct calculation:

$$\tilde{S} = \sum_{g \in G} p(g) \omega(g)^\dagger \sum_{z \in \mathbb{F}_2^n} |\tilde{Z}_z\rangle \langle \tilde{Z}_z| \omega(g) \Lambda(g) \quad (\text{A9})$$

$$= \sum_{g \in G} \sum_{z \in \mathbb{F}_2^n} p(g) (-1)^{\varphi_a(g)} |\tilde{\Xi}_z(g)\rangle \langle \tilde{\Xi}_z(g)| (-1)^{\varphi_a(g)} \Lambda(g) \quad (\text{A10})$$

$$= \sum_{g \in G} \sum_{z \in \mathbb{F}_2^n} p(g) |\tilde{\Xi}_z(g)\rangle \langle \tilde{\Xi}_z(g)| \Lambda(g) \quad (\text{A11})$$

$$= \sum_{a \in \mathbb{F}_2^{2n}} |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a| \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} p(g) \Lambda(g) \quad (\text{A12})$$

$$= \sum_{a \in \mathbb{F}_2^{2n}} s_a |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a| \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} \frac{p(g)}{s_a} \Lambda(g) \quad (\text{A13})$$

$$= \sum_{a \in \mathbb{F}_2^{2n}} s_a |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a| \bar{\Lambda}_a, \quad (\text{A14})$$

where $s_a := \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} p(g)$ and $\bar{\Lambda}_a := \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} \frac{p(g)}{s_a} \Lambda(g)$.

For Pauli noise parameterized by $\Lambda(g) = |\mathbb{1}\rangle \langle \mathbb{1}| + \sum_{a \neq 0} \lambda_a(g) |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a|$, the term $(\tilde{\sigma}_a | \bar{\Lambda}_a$ simplifies to

$$(\tilde{\sigma}_a | \bar{\Lambda}_a = \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} \frac{p(g)}{s_a} (\tilde{\sigma}_a | \Lambda(g) = \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} \frac{p(g)}{s_a} \lambda_a(g) (\tilde{\sigma}_a | =: \bar{\lambda}_a (\tilde{\sigma}_a |. \quad (\text{A15})$$

The frame operator for Pauli noise then becomes $\tilde{S} = \sum_{a \in \mathbb{F}_2^{2n}} s_a \bar{\lambda}_a |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a|$. One can quickly verify that the set $\{\frac{p(g)}{s_a} | g \in \Xi_z^{-1}(\sigma_a), z \in \mathbb{F}_2^n\}$ corresponds to a normalized probability distribution by definition of s_a . Thus, each $\bar{\Lambda}_a$ is an average over noise channels and thereby, a quantum channel itself. The condition $\bar{\lambda}_a \in [-1, 1]$ follows from the general property of Pauli eigenvalues $\lambda_a(g) \in [-1, 1]$. $\square$

Crucially, these average channels depend on a and are taken over preimages of σ_a . The most immediate example where Lemma 2 can be applied is shadow estimation with the global Clifford group for uniform sampling, i.e. $p(g) = 1/|\text{Cl}_n|$. In this case, we have $s_a = 1/(d+1)$ for $a \neq 0$ and $s_0 = 1$ [1]. Furthermore, each $\bar{\Lambda}_a$ is an average over $|\text{Cl}_n|/(d+1)$ right noise channels $\Lambda(g)$, and a characterization of the sets over which averages are taken can be found in Appendix B.

For trace-preserving noise, the identity component in $\tilde{S}$ can be treated differently since, for any channel $\Lambda(g)$, the trace-preservation condition is equivalent to $(\mathbb{1} | \Lambda(g) = (\mathbb{1} |$ and, thus, $(\mathbb{1} | \bar{\Lambda}_a = (\mathbb{1} |$. Moreover, since the adjoint unitary action $\omega(g)$ for each gate g is trace-preserving, it holds that $(\mathbb{1} | \omega(g) | Z_z) = (\mathbb{1} | Z_z) = d \delta_{0,z}$. As $\Xi_z^{-1}(\mathbb{1})$ is defined to be precisely the set of all gates that map Z_z to $\mathbb{1}$, it turns out that $\Xi_z^{-1}(\mathbb{1}) = \emptyset$ for $z \neq 0$ and $\Xi_z^{-1}(\mathbb{1}) = \text{Cl}_n$ for $z = 0$. Thus, for any probability distribution $p(g)$ over Cl_n , it holds true that $s_0 = \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\mathbb{1})} p(g) = \sum_{g \in \Xi_0^{-1}(\mathbb{1})} p(g) = \sum_{g \in \text{Cl}_n} p(g) = 1$.

For clarity of presentation, we restate the main Theorem and its proof here.

Theorem 3. Suppose $G \subset \text{Cl}_n$. The estimation bias is bounded by

$$\begin{aligned} |\mathbb{E}[\hat{o}] - \langle O \rangle| &\leq \|O\|_{\text{st}} \max_{a \in \mathbb{F}_2^{2n}} \|\text{id} - \bar{\Lambda}_a\|_\diamond \\ &\leq \|O\|_{\text{st}} \max_{g \in \text{Cl}_n} \|\text{id} - \Lambda(g)\|_\diamond \end{aligned} \quad (6)$$

for arbitrary gate-dependent noise and

$$|\mathbb{E}[\hat{o}] - \langle O \rangle| \leq \min\{\|O\|_2, \|O\|_{\text{st}}\} \max_{a \in \mathbb{F}_2^{2n}} |1 - \bar{\lambda}_a|$$

for gate-dependent Pauli noise.

Proof. From Lemma 2 we see that the ideal frame operator with $\Lambda(g) = \text{id}$ for any $g \in G$ is given by

$$S = \sum_{a \in \mathbb{F}_2^{2n}} s_a |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a| \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(\sigma_a)} \frac{p(g)}{s_a} \text{id} = \sum_{a \in \mathbb{F}_2^{2n}} s_a |\tilde{\sigma}_a\rangle \langle \tilde{\sigma}_a|, \quad (\text{A16})$$

and its inverse is $S^{-1} = \sum_{a \in \mathbb{F}_2^{2n}} \frac{1}{s_a} |\check{\sigma}_a\rangle\langle\check{\sigma}_a|$. Therefore the effective operation performed by the shadow estimation protocol is $S^{-1}\tilde{S} = \sum_{a \in \mathbb{F}_2^{2n}} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \bar{\Lambda}_a$. We can now bound the absolute error as follows:

$$|(O|\rho) - (O|S^{-1}\tilde{S}|\rho)| = \left| (O|(\text{id} - \sum_a |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \bar{\Lambda}_a)|\rho) \right| \quad (\text{A17})$$

$$= \left| \sum_a (O|\check{\sigma}_a\rangle\langle\check{\sigma}_a|(\text{id} - \bar{\Lambda}_a)|\rho) \right| \quad (\text{A18})$$

$$\leq \sum_a |(O|\check{\sigma}_a)| |(\check{\sigma}_a|(\text{id} - \bar{\Lambda}_a)|\rho)| \quad (\text{A19})$$

$$\leq \sum_a |(O|\check{\sigma}_a)| \|\check{\sigma}_a\|_\infty \|(\text{id} - \bar{\Lambda}_a)(\rho)\|_1 \quad (\text{A20})$$

$$\leq \sum_a |(O|\check{\sigma}_a)| \|\check{\sigma}_a\|_\infty \|\text{id} - \bar{\Lambda}_a\|_\diamond \quad (\text{A21})$$

$$\leq \max_a \|\text{id} - \bar{\Lambda}_a\|_\diamond \sum_a \frac{|(O|\check{\sigma}_a)|}{\sqrt{d}} \quad (\text{A22})$$

$$= \max_a \|\text{id} - \bar{\Lambda}_a\|_\diamond \|O\|_{\text{st}}, \quad (\text{A23})$$

where from (A19) to (A20) we used the matrix Hölder inequality. From (A21) to (A22) we used $\|\check{\sigma}_a\|_\infty = \|\sigma_a\|_\infty/\sqrt{d} = 1/\sqrt{d}$ and $\|(\text{id} - \bar{\Lambda}_a)(\rho)\|_1 \leq \|\text{id} - \bar{\Lambda}_a\|_\diamond$ [28]. Finally, we used $(O|\check{\sigma}_a) = (O|\sigma_a)/\sqrt{d}$ and the definition of the stabilizer norm in Eq. (5).

It remains to prove the bound for Pauli noise, where we use that $(\check{\sigma}_a|\bar{\Lambda}_a = \bar{\lambda}_a(\check{\sigma}_a|$ as seen in Eq. (A15). This leads to $S^{-1}\tilde{S} = \sum_{a \in \mathbb{F}_2^{2n}} \bar{\lambda}_a |\check{\sigma}_a\rangle\langle\check{\sigma}_a|$ and $\text{id} - S^{-1}\tilde{S} = \sum_{a \in \mathbb{F}_2^{2n}} (1 - \bar{\lambda}_a) |\check{\sigma}_a\rangle\langle\check{\sigma}_a|$. Consequently, we can bound the bias for Pauli noise as

$$|(O|\rho) - (O|S^{-1}\tilde{S}|\rho)| = \left| (O| \sum_{a \in \mathbb{F}_2^{2n}} (1 - \bar{\lambda}_a) |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \rho) \right| \quad (\text{A24})$$

$$\leq \max_a |1 - \lambda_a| \sum_{a \neq 0} |(O|\check{\sigma}_a\rangle\langle\check{\sigma}_a|\rho)| \quad (\text{A25})$$

$$\leq \max_a |1 - \lambda_a| \|O\|_2 \|\rho\|_2 \quad (\text{A26})$$

$$\leq \max_a |1 - \lambda_a| \|O\|_2. \quad (\text{A27})$$

From Eq. (A25) to Eq. (A26), the Cauchy-Schwarz inequality has been used, and the last line follows from the fact that $\|\rho\|_2 \leq 1$. We can alternatively bound Eq. (A25) by using $|(O|\check{\sigma}_a|\rho)| \leq 1/\sqrt{d}$ to come by $\max_a |1 - \lambda_a| \sum_a |(O|\check{\sigma}_a\rangle\langle\check{\sigma}_a|\rho)| \leq \max_a |1 - \lambda_a| \sum_a |(O|\check{\sigma}_a)|/\sqrt{d} = \max_a |1 - \lambda_a| \|O\|_{\text{st}}$. $\square$

As mentioned in the discussion of Theorem 3 in the main text, the error bound for Pauli noise is as strong as one could hope for in the case where O has unit rank. The maximum error on Pauli eigenvalues is bounded by the diamond norm for Pauli channels and $\|O\|_2$ relates to the largest expectation value over all input states via $\|O\|_2 \leq \sqrt{\text{rank } O} \|O\|_\infty$. In general we have $\|O\|_{\text{st}} \leq \sqrt{d} \|O\|_2$, meaning that the bound in terms of $\|O\|_2$ can be stronger by a factor of $\sqrt{d}$ compared to the bound in terms of $\|O\|_{\text{st}}$. We will now give an explicit example showing that, at least for local Clifford shadow estimation, the bias can scale exponentially in the system size.

Proposition 1. *Let $O = (|H\rangle\langle H|)^{\otimes n}$ with the magic state $|H\rangle = \frac{1}{\sqrt{2}}(|0\rangle + e^{i\pi/4}|1\rangle)$ and consider shadow estimation with local Clifford unitaries. There exists a state ρ and implementation map $\phi_\epsilon(g) = (1 - \epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g)$ such that $|\mathbb{E}[\hat{o}] - \langle O \rangle| = \kappa\epsilon$ with $\kappa \in \Omega(d^{1/4})$.*

Proof. In the following we use a noise model $\Lambda(g) = \bigotimes_{i=1}^n \Lambda_i(g_i)$. Recall from Eq. (A7) that the local Cliffords g_i act as $\omega(g_i)(Z) = (-1)^{\varphi_Z(g_i)} \Xi_Z(g_i)$. Then, we define

$$\Lambda_i(g_i) = \begin{cases} \omega(g_i)^\dagger \omega(X) & \text{if } \varphi_Z(g_i) = 1 \\ \omega(g_i)^\dagger & \text{if } \varphi_Z(g_i) = 0, \end{cases} \quad (\text{A28})$$

where $\omega(X)(\rho) = X\rho X$. The frame operator inherits the ϵ -dependence from ϕ_ϵ , and we write $\tilde{S}_\epsilon \equiv \tilde{S}(\phi_\epsilon)$.

Let us first consider the case $\epsilon = 1$. Since we can write the Z-basis measurement operator as $M = \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z\rangle\langle\check{Z}_z| = (|\check{1}\rangle\langle\check{1}| + |\check{Z}\rangle\langle\check{Z}|)^{\otimes n}$, the frame operator factorizes for the above local noise model:

$$\tilde{S}_1 = \frac{1}{|\text{Cl}_1^{\times n}|} \sum_{g \in \text{Cl}_1^{\times n}} \omega(g)^\dagger \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z\rangle\langle\check{Z}_z| \omega(g) \Lambda(g) \quad (\text{A29})$$

$$= \bigotimes_{i=1}^n \left(\frac{1}{|\text{Cl}_1|} \sum_{g_i \in \text{Cl}_1} \omega(g_i)^\dagger (|\check{1}\rangle\langle\check{1}| + |\check{Z}\rangle\langle\check{Z}|) \omega(g_i) \Lambda_i(g_i) \right). \quad (\text{A30})$$

The action of a local Clifford operation g_i can be written as $\omega(g_i)^\dagger |Z\rangle = (-1)^{\varphi_Z(g_i)} |\Xi_Z(g_i)\rangle$ with $\varphi_Z(g_i) \in \mathbb{F}_2$. By inserting $\Lambda_i(g_i)$ we get

$$\tilde{S}_1 = \bigotimes_{i=1}^n \left(|\check{1}\rangle\langle\check{1}| + \frac{1}{|\text{Cl}_1|} \sum_{g_i \in \text{Cl}_1} (-1)^{\varphi_Z(g_i)} |\check{\Xi}_Z(g_i)\rangle\langle\check{Z}| (-1)^{\varphi_Z(g_i)} \right) = \left(|\check{1}\rangle\langle\check{1}| + \frac{1}{3} \sum_{a \in \mathbb{F}_2^2 \setminus \{0\}} |\check{\sigma}_a\rangle\langle\check{Z}| \right)^{\otimes n}. \quad (\text{A31})$$

With this error model, each measurement is done in the Z-basis, and the sign cancels with the sign acquired in post-processing from the action of $\omega(g_i)^\dagger$. Since the ideal frame operator is given by $S = \tilde{S}_0 = \bigotimes_{i=1}^n \left(|\check{1}\rangle\langle\check{1}| + \frac{1}{3} \sum_{a \neq 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \right)$, we find that $S^{-1} \tilde{S} = \left(|\check{1}\rangle\langle\check{1}| + \sum_{a \neq 0} |\check{\sigma}_a\rangle\langle\check{Z}| \right)^{\otimes n}$. The expectation value of an observable O on the initial state $E_0 = \left[\frac{1}{\sqrt{2}} (\check{1} + \check{Z}) \right]^{\otimes n}$ is consequently given by

$$(O | S^{-1} \tilde{S}_1 | E_0) = \frac{1}{\sqrt{d}} (O | \left(|\check{1}\rangle + \sum_{a \in \mathbb{F}_2^2 \setminus \{0\}} |\check{\sigma}_a\rangle \right)^{\otimes n}) \quad (\text{A32})$$

$$= \frac{1}{\sqrt{d}} \sum_{a \in \mathbb{F}_2^{2n}} (O | \check{\sigma}_a). \quad (\text{A33})$$

It remains now to compute $\tilde{S}_\epsilon$ for the implementation map $\phi_\epsilon(g) = (1 - \epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g)$. For this we note that

$$\tilde{S}_\epsilon = \mathbb{E}_g [\omega^\dagger(g) M ((1 - \epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g))] \quad (\text{A34})$$

$$= (1 - \epsilon) \mathbb{E}_g [\omega^\dagger(g) M \omega(g)] + \epsilon \mathbb{E}_g [\omega^\dagger(g) M \omega(g) \Lambda(g)] \quad (\text{A35})$$

$$= (1 - \epsilon) \tilde{S}_0 + \epsilon \tilde{S}_1. \quad (\text{A36})$$

Therefore the bias is given by $|\mathbb{E}[\hat{o}] - \langle O \rangle| = |(1 - \epsilon)\langle O \rangle + \epsilon(O | S^{-1} \tilde{S}_1 | E_0) - \langle O \rangle| = \epsilon |\sum_{a \in \mathbb{F}_2^{2n}} (O | \check{\sigma}_a) / \sqrt{d} - \langle O \rangle|$. So far the calculation was independent of the observable, and we now consider the magic state $O = (|H\rangle\langle H|)^{\otimes n}$. From the definition of $|H\rangle$ we can read off $\langle O \rangle = |\langle H | E_0 \rangle|^{2n} = \frac{1}{d}$. Furthermore, one can show that $|H\rangle\langle H| = \frac{1}{\sqrt{2}} \left[\check{1} + \frac{1}{\sqrt{2}} (\check{X} + \check{Y}) \right]$, and we get

$$(O | S^{-1} \tilde{S}_1 | E_0) = \left[\frac{1}{2} \sum_{a \in \mathbb{F}_2^2} ((\check{1} + \frac{1}{\sqrt{2}} (\check{X} + \check{Y})) | \check{\sigma}_a) \right]^n = \left[\frac{1 + \sqrt{2}}{2} \right]^n. \quad (\text{A37})$$

Since the entries of $O = (|H\rangle\langle H|)^{\otimes n}$ in the Pauli basis are positive, the stabilizer norm of O is by Eq. (A33) exactly the just derived expression, $\|O\|_{\text{st}} = \left[\frac{1 + \sqrt{2}}{2} \right]^n \geq 2^{n/4}$. By putting everything together, we arrive at the desired result $|\mathbb{E}[\hat{o}] - \langle O \rangle| = \|O\|_{\text{st}} - 1/d | \epsilon = \kappa \epsilon$. $\square$

Whether the same error scaling with $\|O\|_{\text{st}}$ can also occur for Clifford-based shadow estimation protocols other than uniform sampling from the local Clifford group remains open. To understand for instance the difference between local and global Cliffords, we can look at the above example in terms of average error channels. What allows the errors to accumulate in Proposition 1 is the fact that for local noise, the local noise averages $\bar{\Lambda}_X, \bar{\Lambda}_Y, \bar{\Lambda}_Z$ are each taken over disjoint subsets over the local Clifford group and, thus, an error model exists such that they can be independently chosen. For uniform sampling from the global Clifford group this is not the case anymore. As shown in the next section, each of the $d^2 - 1$ average noise channels $\bar{\Lambda}_a$ is an average over $|\text{Cl}_n|/(d + 1)$ channels $\Lambda(g)$, therefore the $\bar{\Lambda}_a$ cannot all be independent.

Appendix B: Explicit calculation of the noise average for the Clifford group

In the following, we write $\text{CNOT}_{i,j}$ for the controlled-NOT gate with control qubit i and target qubit j . More generally, given a 2-qubit unitary U , we write $U_{i,j}$ for its application on the ordered qubit pair (i, j) . Moreover, SWAP is the 2-qubit SWAP gate and $\mathcal{P}_n = \langle \mathcal{P}_n \rangle$ denotes the n -qubit Pauli group.

Lemma 8. *Let $\mathcal{N} := \bigcup_{i=1}^{n-1} \left\{ \prod_{j=1}^i U_{j+1,j} \mid U \in \{\text{SWAP}, \text{CNOT}\} \right\} \cup \mathbb{1}$. Then the following holds:*

$$M = \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z\rangle \langle \check{Z}_z| = |\check{\mathbb{1}}\rangle \langle \check{\mathbb{1}}| + \sum_{g \in \mathcal{N}} \omega(g)^\dagger |\check{Z}_1\rangle \langle \check{Z}_1| \omega(g). \quad (\text{B1})$$

Proof. To see this we first note that SWAP and $\text{CNOT}_{2,1}$ are given as

$$\text{CNOT}_{2,1} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix}, \quad \text{SWAP} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Since $Z_{(1,0)} = \text{diag}(1, 1, -1, -1)$, $Z_{(0,1)} = \text{diag}(1, -1, 1, -1)$ and $Z_{(1,1)} = \text{diag}(1, -1, -1, 1)$ are all diagonal and the adjoint actions of SWAP and $\text{CNOT}_{2,1}$ permute diagonal elements, one can quickly verify that $\omega(\text{SWAP})|Z_{(1,0)}\rangle = |Z_{(0,1)}\rangle$ and $\omega(\text{CNOT}_{2,1})|Z_{(1,0)}\rangle = |Z_{(1,1)}\rangle$. Consecutively applying either $\text{CNOT}_{2,1}$ or SWAP along the qubit chain to Z_1 generates all Z_z with $z \in \mathbb{F}_2^n \setminus \{0, e_1\}$. One can also see that $|\mathcal{N}| = 1 + \sum_{i=1}^{n-1} 2^i = 2^n - 1$, which is the number of non-identity terms in Eq. (B1). $\square$

Lemma 9. *The stabilizer of $|Z_1\rangle \langle Z_1|$ under the action $g \mapsto \omega^\dagger(g)(\cdot)\omega(g)$ with $g \in \text{Cl}_n$ is given by $\text{St}_{e_1} \equiv \text{St}_{Z_1} = \text{Cl}_{n-1} \cdot \text{HW}_n \cdot \langle S_1, \{\text{CZ}_{1,i}, \text{CNOT}_{1,i}\}_{i \in \{2, \dots, n\}} \rangle$.*

Proof. Since Z_1 is just the identity on qubits $2, \dots, n$, any adjoint action by a unitary acting only on qubits $2, \dots, n$ leaves Z_1 invariant. This holds, in particular, for $\text{id} \otimes \omega(g)$ with $g \in \text{Cl}_{n-1}$. For $g \in \mathcal{P}_n$ we know that $\omega(g)^\dagger |\sigma_a\rangle = \pm |\sigma_a\rangle$ for any $a \in \mathbb{F}_2^n$ and, thus, $\omega^\dagger(g)|\sigma_a\rangle \langle \sigma_a| \omega(g) = |\sigma_a\rangle \langle \sigma_a|$. The remaining Clifford group elements that can stabilize $|Z_1\rangle \langle Z_1|$ act only on qubit 1 or between qubit 1 and qubits $2, \dots, n$. Since Z_1 is diagonal in the computational basis, this includes all diagonal Cliffords which we did not already count in Cl_{n-1} , namely $\langle S_1, \text{CZ}_{1,i} \rangle$. In addition, identical diagonal elements of Z_1 can be permuted and one can easily verify that, for instance, $\omega(\text{CNOT})|Z_{(1,0)}\rangle = |Z_{(1,0)}\rangle$.

To show that this is indeed the complete stabilizer, note that the orbit of $|Z_1\rangle \langle Z_1|$ corresponds to all non-identity Pauli strings and thus has size $4^n - 1$. Hence, we have $|\text{Cl}_n|/|\text{St}_{e_1}| = 4^n - 1$, where

$$|\text{Cl}_n| = 2^{2n+3} 2^{n^2} \prod_{i=1}^n (4^i - 1), \quad |\mathcal{P}_n| = 2^{2n+2}. \quad (\text{B2})$$

The order of the above defined group $\text{Cl}_{n-1} \cdot \mathcal{P}_n \cdot \langle S_1, \{\text{CZ}_{1,i}, \text{CNOT}_{1,i}\}_{i \in \{2, \dots, n\}} \rangle$ can be computed by observing that we can simply compute the cardinalities of the first and last factor up to Pauli operators, and multiply those by $|\mathcal{P}_n|$. Since CNOT normalizes diagonal Clifford unitaries, $\langle S_1, \text{CZ}_{1,i}, \text{CNOT}_{1,i} \rangle = \langle S_1, \text{CZ}_{1,i} \rangle \rtimes \langle \text{CNOT}_{1,i} \rangle$ is a semidirect product and $\langle S_1, \text{CZ}_{1,i} \rangle$ is Abelian. We have 2^{n-1} possibilities of applying $\text{CZ}_{1,i}$ and $\text{CNOT}_{1,i}$ and the S-gate has order 4, hence $|\langle S_1, \text{CZ}_{1,i}, \text{CNOT}_{1,i} \rangle| = 2 \cdot 2^{n-1} 2^{n-1} = 2^{2n-1}$. Putting everything together, we find

$$|\text{Cl}_{n-1} \cdot \mathcal{P}_n \cdot \langle S_1, \{\text{CZ}_{1,i}, \text{CNOT}_{1,i}\}_{i \in \{2, \dots, n\}} \rangle| = |\text{Cl}_{n-1}/\mathcal{P}_{n-1}| \times |\mathcal{P}_n| \times |\langle S_1, \text{CZ}_{1,i}, \text{CNOT}_{1,i} \rangle|/|\langle Z_1 \rangle| \quad (\text{B3})$$

$$= 2^{(n-1)^2+1} \prod_{i=1}^{n-1} (4^i - 1) \times 2^{2n+2} \times 2^{2n-1} \quad (\text{B4})$$

$$= 2^{2n+3} 2^{n^2} \frac{1}{4^n - 1} \prod_{i=1}^n (4^i - 1) \quad (\text{B5})$$

$$= \frac{|\text{Cl}_n|}{4^n - 1}, \quad (\text{B6})$$

which shows the claim. $\square$

Let g_a be any element of Cl_n that satisfies $\omega^\dagger(g_a)|Z_1\rangle \langle Z_1| \omega(g_a) = |\sigma_a\rangle \langle \sigma_a|$ for $a \neq 1$. The stabilizer of $|\sigma_a\rangle \langle \sigma_a|$ can then simply be written as $\text{St}_a = g_a \text{St}_{e_1}$. Moreover, the stabilizers $g_a \text{St}_{e_1}$ are exactly the left cosets of St_{e_1} and, thus, $\text{Cl}_n = \bigcup_{a \neq 0} g_a \text{St}_{e_1}$, where $g_a \text{St}_{e_1}$ and $g_{a'} \text{St}_{e_1}$ are pairwise disjoint for $a \neq a'$.

Proposition 10. *The frame operator for shadow estimation with uniform sampling from the n -qubit Clifford group is given by*

$$\tilde{S} = |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{d+1} \sum_{a \neq 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \mathbb{E}_{h \in \mathcal{N}} \mathbb{E}_{h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g_a), \quad (\text{B7})$$

where g_a can be any element of Cl_n satisfying $\omega^\dagger(g_a)|Z_1\rangle\langle Z_1|\omega(g_a) = |\sigma_a\rangle\langle\sigma_a|$.

Proof. With the use of Lemma 8 and Lemma 9 we can successively rewrite the frame operator as follows:

$$\tilde{S} = \frac{1}{|\text{Cl}_n|} \sum_{g \in \text{Cl}_n} \omega^\dagger(g) \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z\rangle\langle\check{Z}_z| \omega(g) \Lambda(g) \quad (\text{B8})$$

$$= |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{|\text{Cl}_n|} \sum_{h \in \mathcal{N}} \sum_{g \in \text{Cl}_n} \omega(hg)^\dagger |\check{Z}_1\rangle\langle\check{Z}_1| \omega(hg) \Lambda(g) \quad (\text{B9})$$

$$= |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{|\text{Cl}_n|} \sum_{g \in \text{Cl}_n} \omega(g)^\dagger |\check{Z}_1\rangle\langle\check{Z}_1| \omega(g) \sum_{h \in \mathcal{N}} \Lambda(h^{-1}g) \quad (\text{B10})$$

$$= |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{|\text{Cl}_n|} \sum_{g \in \text{Cl}_n / \text{St}_{e_1}} \omega(g)^\dagger |\check{Z}_1\rangle\langle\check{Z}_1| \omega(g) \sum_{h \in \mathcal{N}, h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g) \quad (\text{B11})$$

$$= |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{|\text{Cl}_n|} \sum_{a \neq \mathbb{I}} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \sum_{h \in \mathcal{N}, h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g_a). \quad (\text{B12})$$

Since $|\text{St}_{e_1}| = \frac{|\text{Cl}_n|}{4^n - 1}$ and $|\mathcal{N}| = d - 1$, we can rewrite the last line in terms of the channel averages as

$$\frac{1}{|\text{Cl}_n|} \sum_{h \in \mathcal{N}, h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g_a) = \frac{1}{(d+1)(d-1)|\text{St}_{e_1}|} \sum_{h \in \mathcal{N}, h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g_a) \quad (\text{B13})$$

$$= \frac{1}{d+1} \mathbb{E}_{h \in \mathcal{N}} \mathbb{E}_{h' \in \text{St}_{e_1}} \Lambda(h^{-1}h'g_a). \quad (\text{B14})$$

□

This result ties back to the general form derived in Lemma 2 via $\Xi_b^{-1}(a) = h_b^{-1} \text{St}_{e_1} g_a$ and $s_a = \frac{1}{d+1}$. We will now turn to the local Clifford group and determine the factors s_a , as well as the compositions of $\bar{\Lambda}_a$. The result will take a simpler form for *local noise*, which we define as noise that factorizes as $\Lambda(g) = \bigotimes_{i=1}^n \Lambda^{(i)}(g_i)$ on all $g \in \text{Cl}_1^{\times n}$. We also define the support of $a \in \mathbb{F}_2^{2n}$ as $|\text{supp}(a)| = |\{i \in [n] : a_i \neq 0\}|$.

Proposition 11. *Let $G(a) \subseteq \text{Cl}_1$ be defined by*

$$G(a) = \begin{cases} \text{Cl}_1 & a = 0 \\ \text{St}(Z)g_a & a \in (\mathbb{F}_2^{2*}). \end{cases} \quad (\text{B15})$$

Then the frame operator for the n -qubit local Clifford group is given by

$$\tilde{S} = \sum_{a \in \mathbb{F}_2^{2n}} \frac{1}{3^{|\text{supp}(a)|}} |\check{\sigma}_a\rangle\langle\check{\sigma}_a| \bar{\Lambda}_a, \quad (\text{B16})$$

where $\bar{\Lambda}_a = \mathbb{E}_{g_1 \in G(a_1)} \cdots \mathbb{E}_{g_n \in G(a_n)} \Lambda(g_1, \dots, g_n)$ for *global noise* and $\bar{\Lambda}_a = \mathbb{E}_{g_1 \in G(a_1)} \Lambda^{(1)}(g_1) \otimes \cdots \otimes \mathbb{E}_{g_n \in G(a_n)} \Lambda^{(n)}(g_n) = \bigotimes_{i=1}^n \bar{\Lambda}_{a_i}$ for *local noise*.

Proof. Since the local Clifford group also satisfies the conditions in Lemma 2 it remains to show that the averages $\bar{\Lambda}_a$ take the above form and that the coefficients are $s_a = \frac{1}{3^{|\text{supp}(a)|}}$. We will now explicitly proof the case $n = 2$, from which the result for an arbitrary system size can be straightforwardly generalized. In this case $M = (|\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + |\check{Z}\rangle\langle\check{Z}|)^{\otimes 2} = \sum_{z_1, z_2 \in \mathbb{F}_2} |\check{Z}_{z_1}\rangle\langle\check{Z}_{z_1}| \otimes |\check{Z}_{z_2}\rangle\langle\check{Z}_{z_2}|$. For $\text{Cl}_1 \times \text{Cl}_1$, the product representation $\omega(g_1, g_2) = \omega(g_1) \otimes \omega(g_2)$, and implementation map $\phi(g_1, g_2) = (\omega(g_1) \otimes \omega(g_2)) \Lambda(g_1, g_2)$, the frame operator becomes

$$\tilde{S} = \frac{1}{|\text{Cl}_1|^2} \sum_{g_1, g_2 \in \text{Cl}_1} \sum_{z_1, z_2 \in \mathbb{F}_2} (\omega(g_1)^\dagger |\tilde{Z}_{z_1}) (\tilde{Z}_{z_1} |\omega(g_1) \otimes \omega(g_2)^\dagger |\tilde{Z}_{z_2}) (\tilde{Z}_{z_2} |\omega(g_2)) \Lambda(g_1, g_2) \quad (\text{B17})$$

$$= \frac{1}{|\text{Cl}_1|^2} \sum_{g_2 \in \text{Cl}_1} \sum_{a_1 \in \mathbb{Z}_4} \sum_{z_2 \in \mathbb{F}_2} (|\check{\sigma}_{a_1}) (\check{\sigma}_{a_1} | \otimes \omega(g_2)^\dagger |\tilde{Z}_{z_2}) (\tilde{Z}_{z_2} |\omega(g_2)) \sum_{z_1 \in \mathbb{F}_2} \sum_{g_1 \in \Xi_{z_1}^{-1}(a_1)} \Lambda(g_1, g_2) \quad (\text{B18})$$

$$= \frac{1}{|\text{Cl}_1|^2} \sum_{a_1 \in \mathbb{Z}_4} |\check{\sigma}_{a_1}) (\check{\sigma}_{a_1} | \otimes \sum_{a_2 \in \mathbb{Z}_4} |\check{\sigma}_{a_2}) (\check{\sigma}_{a_2} | \sum_{z_1, z_2 \in \mathbb{F}_2} \sum_{g_1 \in \Xi_{z_1}^{-1}(a_1)} \sum_{g_2 \in \Xi_{z_2}^{-1}(a_2)} \Lambda(g_1, g_2). \quad (\text{B19})$$

Since $\omega(g)^\dagger |\mathbb{1})(\mathbb{1} |\omega(g) = \mathbb{1}$ for all $g \in \text{Cl}_1$ and $\omega(g)^\dagger |\sigma_a)(\sigma_a |\omega(g) \neq 1$ for all $a \neq 0$ and $g \in \text{Cl}_1$, we know that $\Xi_0^{-1}(\mathbb{1}) = \text{Cl}_n$, $\Xi_0^{-1}(\sigma_a) = \emptyset$ as well as $\Xi_a^{-1}(\mathbb{1}) = \emptyset$. Moreover, if we apply Proposition 10 to the case $n = 1$, we see that $\mathcal{N} = \{\mathbb{1}\}$ and $\Xi_1^{-1}(\sigma_a) = \text{St}_1 g_a$. The cosets $G(a \neq 0) = \text{St}_1 g_a$ are disjoint for $a \in \{X, Y, Z\}$ and of order $|\text{Cl}_1|/3$. We then get

$$\tilde{S} = \sum_{a_1, a_2 \in \mathbb{Z}_4} \frac{|G(a_1)|}{|\text{Cl}_1|} \frac{|G(a_2)|}{|\text{Cl}_1|} |\check{\sigma}_{a_1}) (\check{\sigma}_{a_1} | \otimes |\check{\sigma}_{a_2}) (\check{\sigma}_{a_2} | \mathbb{E}_{g_1 \in G(a_1)} \mathbb{E}_{g_2 \in G(a_2)} \Lambda(g_1, g_2) \quad (\text{B20})$$

$$= \sum_{a_1, a_2 \in \mathbb{Z}_4} \frac{1}{3^{|\text{supp}(a)|}} |\check{\sigma}_{a_1}) (\check{\sigma}_{a_1} | \otimes |\check{\sigma}_{a_2}) (\check{\sigma}_{a_2} | \mathbb{E}_{g_1 \in G(a_1)} \mathbb{E}_{g_2 \in G(a_2)} \Lambda(g_1, g_2). \quad (\text{B21})$$

The last line follows from

$$\frac{|G(a)|}{|\text{Cl}_1|} = \begin{cases} 1 & a = 0 \\ 1/3 & a \in (\mathbb{F}_2^*)^*, \end{cases} \quad (\text{B22})$$

which we can write as $\frac{|G(a_1)|}{|\text{Cl}_1|} = 3^{-|\text{supp}(a_1)|}$ whereafter $\frac{|G(a_1)||G(a_2)|}{|\text{Cl}_1|^2} = 3^{-|\text{supp}(a)|}$. For local noise, we get

$$\mathbb{E}_{g_1 \in G(a_1)} \mathbb{E}_{g_2 \in G(a_2)} \Lambda(g_1, g_2) = \mathbb{E}_{g_1 \in G(a_1)} \Lambda^{(1)}(g_1) \otimes \mathbb{E}_{g_2 \in G(a_2)} \Lambda^{(2)}(g_2) = \bar{\Lambda}_{a_1} \otimes \bar{\Lambda}_{a_2}. \quad (\text{B23})$$

The n -qubit local Clifford group result then follows. $\square$

A different locality structure of the noise channels $\Lambda(g)$ (other than fully local noise) will lead to the corresponding locality structure on the average noise channels. For instance, if all noise channels factorize along a bipartition of the set of qubits, the average noise channels will inherit this factorization.

Appendix C: A norm inequality for Pauli channels

In the following we consider Pauli channels that act as $\Lambda(\rho) = \sum_{b \in \mathbb{F}_2^{2n}} p_b \sigma_b \rho \sigma_b$, with $p_b \in [0, 1]$ and $\sum_b p_b = 1$. The corresponding superoperators are known to be diagonal in the Pauli basis and to have eigenvalues $\lambda_a = \sum_{b \in \mathbb{F}_2^{2n}} (-1)^{[a, b]} p_b$ where $[a, b] = 0$ if σ_a and σ_b commute and $[a, b] = 1$ otherwise.

Lemma 12. *Let Λ and Λ' be Pauli channels, then it holds that $\|\Lambda - \Lambda'\|_\infty \leq \|\Lambda - \Lambda'\|_\diamond$.*

Proof. Let Λ and Λ' be given by the probability distributions p and p' , respectively, hence their eigenvalues are $\lambda_a = \sum_{b \in \mathbb{F}_2^{2n}} (-1)^{[a, b]} p_b$ and $\lambda'_a = \sum_{b \in \mathbb{F}_2^{2n}} (-1)^{[a, b]} p'_b$. Since Λ and Λ' are both diagonal in the Pauli basis, we have $\|\Lambda - \Lambda'\|_\infty = \max_a |\lambda_a - \lambda'_a|$. It then follows that $\max_a |\lambda_a - \lambda'_a| = \max_a |\sum_b (-1)^{[a, b]} (p_b - p'_b)| \leq \sum_b |p_b - p'_b| = \|\Lambda - \Lambda'\|_\diamond$, where the last step uses a well-known relation for the diamond distance of Pauli channels [29]. $\square$

Since the channel average over Pauli channels is again a Pauli channel, the above statement holds in particular for $\bar{\Lambda}_a$ and the identity operation, $\|\text{id} - \bar{\Lambda}_a\|_\infty \leq \|\text{id} - \bar{\Lambda}_a\|_\diamond$.

Appendix D: Variance bounds for gate-dependent noise

To bound the variance of the estimator, $\mathbb{V}[\hat{o}] = \mathbb{E}(\hat{o}^2) - \mathbb{E}(\hat{o})^2$, we compute the second moment

$$\mathbb{E}(\hat{o}^2) = (O \otimes O | (S^{-1})^{\otimes 2} \sum_{x \in \mathbb{F}_2^n} \sum_{g \in G} p(g) \omega(g)^{\dagger \otimes 2} | E_x \otimes E_x \rangle \langle E_x | \omega(g) \Lambda(g) | \rho \rangle. \quad (\text{D1})$$

It will come in handy again to express the map $M_3 := \sum_{x \in \mathbb{F}_2^n} |E_x \otimes E_x\rangle \langle E_x|$ in the normalized Pauli basis. By using that $\langle E_x | \sigma_a \rangle = 0$ if σ_a is not diagonal and $\langle E_x | Z_z \rangle = (-1)^{x \cdot z}$, one can readily show that

$$M_3 = \frac{1}{\sqrt{d}} \sum_{z, z' \in \mathbb{F}_2^n} |\check{Z}_z \otimes \check{Z}_{z'}\rangle \langle \check{Z}_{z+z'}| = \frac{1}{d^2} \sum_{z, z' \in \mathbb{F}_2^n} |Z_z \otimes Z_{z'}\rangle \langle Z_{z+z'}|. \quad (\text{D2})$$

Here and in the following, addition of binary vector such as $z + z'$ is within the binary field $\mathbb{F}_2$, i.e. to be taken modulo 2. In the absence of noise, one can show that the relevant operator in Eq. (D1) can be written as

$$S_3 := \sum_{g \in G} p(g) \omega(g)^{\dagger \otimes 2} M_3 \omega(g) = \frac{1}{\sqrt{d}} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} s_{a, a'} |\check{\sigma}_a \otimes \check{\sigma}_{a'}\rangle \langle \check{\sigma}_{a+a'}|, \quad (\text{D3})$$

for suitable constants $s_{a, a'} \in \mathbb{R}$. Under noise, we show that the analogous operator

$$\tilde{S}_3 := \sum_{g \in G} p(g) \omega(g)^{\dagger \otimes 2} M_3 \omega(g) \Lambda(g) \quad (\text{D4})$$

can be brought in a similar form to Eq. (D3) where the noise enters linearly and from the right. We then obtain an analogous statement for the variance as for the expectation value in Lemma 2:

Lemma 13. *Consider a shadow estimation protocol with random sampling from the Clifford group according to an arbitrary probability distribution p that ensures informational completeness. Let $s_{a, a'}$ be as in Eq. (D3) and let s_a and $s_{a'}$ be as in Lemma 2. Then, the second moment for an observable O and a state ρ can be written as*

$$\mathbb{E}(\hat{o}^2) = \frac{1}{\sqrt{d}} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} \frac{s_{a, a'}}{s_a s_{a'}} (O | \check{\sigma}_a \rangle \langle O | \check{\sigma}_{a'} \rangle \langle \check{\sigma}_{a+a'} | \bar{\Lambda}_{a, a'} | \rho \rangle, \quad (\text{D5})$$

where $\bar{\Lambda}_{a, a'}$ are suitable averages of the gate noise channels $\Lambda(g)$.

Proof. For the target implementation of Clifford unitaries, we can again write the action of $\omega^\dagger(g) \otimes \omega^\dagger(g) (\cdot) \omega(g)$ on $|Z_z\rangle \otimes |Z_{z'}\rangle \langle Z_{z+z'}|$ in terms of $\varphi_a : \text{Cl}_n \rightarrow \mathbb{F}_2$ and $\Xi_a : \text{Cl}_n \rightarrow \mathbb{P}_n$. It also holds that $\langle Z_{z+z'} | \omega(g) = \langle g Z_z g^\dagger g Z_{z'} g^\dagger | = \langle \Xi_z(g) \Xi_{z'}(g) | (-1)^{\varphi_z(g) + \varphi_{z'}(g)}$, where the Pauli operators $\Xi_z(g)$ and $\Xi_{z'}(g)$ commute. Therefore, we get

$$\begin{aligned} \omega^\dagger(g) \otimes \omega^\dagger(g) |Z_z \otimes Z_{z'}\rangle \langle Z_{z+z'}| \omega(g) &= (-1)^{\varphi_z(g) + \varphi_{z'}(g)} |\Xi_z(g) \otimes \Xi_{z'}(g)\rangle \langle \Xi_z(g) \Xi_{z'}(g)| (-1)^{\varphi_z(g) + \varphi_{z'}(g)} \\ &= |\Xi_z(g) \otimes \Xi_{z'}(g)\rangle \langle \Xi_z(g) \Xi_{z'}(g)|. \end{aligned}$$

Note that if $\sigma_a = \Xi_z(g)$ and $\sigma_{a'} = \Xi_{z'}(g)$ for suitable a, a' with $[a, a'] = 0$, then $\sigma_a \sigma_{a'} = (-1)^{\beta(a, a')} \sigma_{a+a'}$ for a suitable binary function β . Since $\tilde{S}_3$ depends only linearly on the right noise channels $\Lambda(g)$ we can proceed in analogy to Lemma 2 and rewrite $\tilde{S}_3$ as

$$\begin{aligned} \tilde{S}_3 &= \sum_{g \in G} p(g) \omega(g)^{\dagger \otimes 2} M_3 \omega(g) \Lambda(g) \\ &= \frac{1}{d^2} \sum_{g \in G} p(g) \sum_{z, z' \in \mathbb{F}_2^n} |\Xi_z(g)\rangle \otimes |\Xi_{z'}(g)\rangle \langle \Xi_z(g) \Xi_{z'}(g)| \Lambda(g) \\ &= \frac{1}{d^2} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} |\sigma_a \otimes \sigma_{a'}\rangle \langle \sigma_a \sigma_{a'}| \sum_{z, z' \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a) \cap \Xi_{z'}^{-1}(a')} p(g) \Lambda(g) \\ &= \frac{1}{\sqrt{d}} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} s_{a, a'} |\check{\sigma}_a \otimes \check{\sigma}_{a'}\rangle \langle \check{\sigma}_{a+a'}| \sum_{z, z' \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a) \cap \Xi_{z'}^{-1}(a')} \frac{p(g)}{r_{a, a'}} \Lambda(g), \end{aligned}$$

where $s_{a,a'} = (-1)^{\beta(a,a')} r_{a,a'}$ and $r_{a,a'} = \sum_{z,z' \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a) \cap \Xi_{z'}^{-1}(a')} p(g)$. We can now define average channels $\bar{\Lambda}_{a,a'} := \sum_{z,z' \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a) \cap \Xi_{z'}^{-1}(a')} \frac{p(g)}{r_{a,a'}} \Lambda(g)$ and write $\tilde{S}_3$ as

$$\tilde{S}_3 = \frac{1}{\sqrt{d}} \sum_{\substack{a,a' \in \mathbb{F}_2^{2n} \\ [a,a'] = 0}} s_{a,a'} |\check{\sigma}_a \otimes \check{\sigma}_{a'})(\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} . \quad (\text{D6})$$

The second moment of $\hat{o}$ is then

$$\mathbb{E}(\hat{o}^2) = (O \otimes O | S^{-1} \otimes S^{-1} \tilde{S}_3 | \rho) = \frac{1}{\sqrt{d}} \sum_{\substack{a,a' \in \mathbb{F}_2^{2n} \\ [a,a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} (O | \check{\sigma}_a)(O | \check{\sigma}_{a'}) (\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho) \quad (\text{D7})$$

□

In analogy to Theorem 3, we obtain the following bound on the deviation of the second moment from its value in the absence of noise.

Proposition 14. *In the setting of Theorem 3, assume that $|s_{a,a'}|/(s_a s_{a'}) \leq C$ for all $a \neq a'$ with $(O | \sigma_a) \neq 0$ and $(O | \sigma_{a'}) \neq 0$. Then, we have*

$$|\mathbb{E}(\hat{o}^2) - \mathbb{E}(\hat{o}_{\text{noise-free}}^2)| \leq C \|O\|_{\text{st}}^2 \max_{a,b \in \mathbb{F}_2^{2n}} \|\text{id} - \bar{\Lambda}_{a,b}\|_{\diamond} \leq C \|O\|_{\text{st}}^2 \max_{g \in G} \|\text{id} - \Lambda(g)\|_{\diamond}, \quad (\text{D8})$$

where $\hat{o}_{\text{noise-free}}$ is the shadow estimator in the absence of any noise and $\bar{\Lambda}_{a,b}$ are suitably averaged noise channels.

Proof. We have the following expression for $s_{a,a'}$:

$$s_{a,a'} = \sqrt{d} (\check{\sigma}_a \otimes \check{\sigma}_{a'} | S_3 | \check{\sigma}_{a+a'}) = \sqrt{d} \sum_{x \in \mathbb{F}_2^n} \sum_{g \in G} p(g) (\check{\sigma}_a | \omega(g)^{\dagger} | E_x) (\check{\sigma}_{a'} | \omega(g)^{\dagger} | E_x) (E_x | \omega(g) | \check{\sigma}_{a+a'}). \quad (\text{D9})$$

First, note that if $a = a'$, then $\check{\sigma}_{a+a} = \check{\sigma}_0 = \mathbb{1}/\sqrt{d}$ and hence

$$s_{a,a} = \sum_{x \in \mathbb{F}_2^n} \sum_{g \in G} p(g) (\check{\sigma}_a | \omega(g)^{\dagger} | E_x)^2 = \sum_{x \in \mathbb{F}_2^n} \sum_{g \in G} p(g) (\check{\sigma}_a | \omega(g)^{\dagger} | E_x) (E_x | \omega(g) | \check{\sigma}_a) = (\check{\sigma}_a | S | \check{\sigma}_a) = s_a, \quad (\text{D10})$$

using that all matrix coefficients are real. By assumption, $|s_{a,a'}|/(s_a s_{a'}) \leq C$ for all non-zero terms, and thus we find

$$\begin{aligned} |\mathbb{E}(\hat{o}^2) - \mathbb{E}(\hat{o}_{\text{noise-free}}^2)| &= |(O \otimes O | S^{-1} \otimes S^{-1} (\tilde{S}_3 - S_3) | \rho)| \\ &= \frac{1}{\sqrt{d}} \left| \sum_{\substack{a,a' \in \mathbb{F}_2^{2n} \\ [a,a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} (O | \check{\sigma}_a)(O | \check{\sigma}_{a'}) (\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} - \text{id} | \rho) \right| \\ &\leq \frac{C}{d^2} \sum_{a \neq a'} |(O | \sigma_a)| |(O | \sigma_{a'})| \|\bar{\Lambda}_{a,a'} - \text{id}\|_{\diamond} + \frac{1}{d^2} \sum_a \frac{|(O | \sigma_a)|^2}{s_a} \text{Tr}[(\bar{\Lambda}_{a,a'} - \text{id})(\rho)] \\ &\leq C \|O\|_{\text{st}}^2 \max_{a,a'} \|\bar{\Lambda}_{a,a'} - \text{id}\|_{\diamond}, \end{aligned}$$

using that $\bar{\Lambda}_{a,a'}$ is trace preserving. □

An open question is whether we have $|s_{a,a'}|/(s_a s_{a'}) = \mathcal{O}(1)$ for general distributions p on the Clifford group. The best general upper bound we could find is d , but we think that this is too pessimistic for practically relevant cases. For uniform sampling from a Clifford subgroup, we can – in principle – get an analytical handle on the $s_{a,a'}$ (and s_a) using Schur's lemma and information about the irreps of the subgroup. This can help to improve on Proposition 14, as we illustrate in Appendix D1 and Appendix D2 at the case of the local and global Clifford groups.

1. The global Clifford group

For the global Clifford group Cl_n , the representation ω is composed of two irreducible representation and can be decomposed as $\omega = \tau_0 \oplus \tau_1$. The noiseless second moment operator S_3 for uniform sampling from Cl_n then decomposes as (see also Ref. [30, App. C]):

$$S_3 = \frac{1}{|G|} \bigoplus_{i \in \mathbb{F}_2} \bigoplus_{j \in \mathbb{F}_2} \bigoplus_{k \in \mathbb{F}_2} \sum_{g \in G} \tau_i(g)^\dagger \otimes \tau_j(g)^\dagger M_3 \tau_k(g).$$

By Schur's lemma, the operator

$$\Pi_{ijk} = \frac{1}{|G|} \sum_{g \in G} \tau_i(g)^\dagger \otimes \tau_j(g)^\dagger (\cdot) \tau_k(g) \quad (\text{D11})$$

is an orthogonal projector which is only non-zero if the irrep τ_k is contained in $\tau_i \otimes \tau_j$. It is straightforward to see that Π_{ijk} is thus zero for $(ijk) = (001), (100), (010)$. More generally, rank Π_{ijk} is equal to the multiplicity of τ_k in $\tau_i \otimes \tau_j$. Thus, Π_{ijk} is rank one for $(ijk) = (000), (101), (011), (110)$ and rank two (one) for $(ijk) = (111)$ if $n > 1$ ($n = 1$) [30, Eq. (97)]. The rank one cases can be straightforwardly computed by finding a superoperator I_{ijk} in the range of Π_{ijk} and projecting M_3 onto I_{ijk} , i.e.

$$\Pi_{ijk}(M_3) = \frac{\text{Tr}(I_{ijk}^\dagger M_3)}{\text{Tr}(I_{ijk} I_{ijk}^\dagger)} I_{ijk}$$

The rank two case was already computed in Ref. [30, App. C.1]. Evaluating S_3 in this way results in the following Lemma.

Lemma 15. *For uniform sampling from the global Clifford group, the coefficients $s_{a,a'}/(s_a s_{a'})$ in the second moment are given by*

$$\frac{s_{a,a'}}{s_a s_{a'}} = \begin{cases} 1 & a = 0 \vee a' = 0 \\ d+1 & a \neq 0 \wedge a' \neq 0 \wedge a = a' \\ \frac{2(d+1)}{d+2} (-1)^{\beta(a,a')} & a \neq 0 \wedge a' \neq 0 \wedge a \neq a' \wedge [a, a'] = 0 \\ 0 & \text{else.} \end{cases} \quad (\text{D12})$$

Proof. We will first determine the superoperators I_{ijk} and calculate each individual contribution to S_3 .

(i) We can choose $I_{000} = |\check{\mathbb{1}} \otimes \check{\mathbb{1}}\rangle\langle\check{\mathbb{1}}|$ with $\text{Tr}(I_{000}^\dagger I_{000}) = 1$ and

$$\text{Tr}(I_{000}^\dagger M_3) = \frac{1}{\sqrt{d}} \sum_{z, z' \in \mathbb{F}_2^n} (\check{\mathbb{1}} | \check{Z}_z) (\check{\mathbb{1}} | \check{Z}_{z'}) (\check{Z}_{z+z'} | \check{\mathbb{1}}) = \frac{1}{\sqrt{d}}.$$

The first term in S_3 is thus

$$\frac{\text{Tr}(I_{000}^\dagger M_3)}{\text{Tr}(I_{000}^\dagger I_{000})} I_{000} = \frac{1}{\sqrt{d}} |\check{\mathbb{1}} \otimes \check{\mathbb{1}}\rangle\langle\check{\mathbb{1}}|. \quad (\text{D13})$$

(ii) Next we look at $(ijk) = (011), (101)$. Note that $\sum_{a \neq 0} |\check{\sigma}_a\rangle\langle\check{\sigma}_a|$ is the projector onto the traceless subspace and hence left invariant by $\tau_1^\dagger(g)(\cdot)\tau_1(g)$ for all $g \in G$. Consequently $I_{011} = \sum_{a \neq 0} |\check{\mathbb{1}} \otimes \check{\sigma}_a\rangle\langle\check{\sigma}_a|$, $I_{101} = \sum_{a \neq 0} |\check{\sigma}_a \otimes \check{\mathbb{1}}\rangle\langle\check{\sigma}_a|$ are valid choices with $\text{Tr}(I_{011}^\dagger I_{011}) = \text{Tr}(I_{101}^\dagger I_{101}) = d^2 - 1$. For the overlap with M_3 we find

$$\begin{aligned} \text{Tr}(I_{011}^\dagger M_3) &= \frac{1}{\sqrt{d}} \sum_{z, z' \in \mathbb{F}_2^n} \sum_{a \neq 0} (\check{\mathbb{1}} | \check{Z}_z) (\check{\sigma}_a | \check{Z}_{z'}) (\check{Z}_{z+z'} | \check{\sigma}_a) \\ &= \frac{1}{\sqrt{d}} \sum_{z, z' \in \mathbb{F}_2^n} \sum_{a \neq 0} \delta_{z,0} \delta_{z',a} \delta_{z+z',a} \\ &= \frac{1}{\sqrt{d}} \sum_{z' \neq 0} 1 \\ &= \frac{1}{\sqrt{d}} (d-1). \end{aligned}$$

The same result can be obtained for $\text{Tr}(I_{101}^\dagger M_3)$, and we get the contributions

$$\frac{\text{Tr}(I_{011}^\dagger M_3)}{\text{Tr}(I_{011}^\dagger I_{011})} I_{011} = \frac{1}{\sqrt{d}(d+1)} \sum_{a \neq 0} |\check{1} \otimes \check{\sigma}_a)(\check{\sigma}_a|, \quad \frac{\text{Tr}(I_{101}^\dagger M_3)}{\text{Tr}(I_{101}^\dagger I_{101})} I_{101} = \frac{1}{\sqrt{d}(d+1)} \sum_{a \neq 0} |\check{\sigma}_a \otimes \check{1})(\check{\sigma}_a|. \quad (\text{D14})$$

(iii) The (110) case is treated in [30, App. C.1a] using

$$I_{110} = [F - |\check{1} \otimes \check{1}](|\check{1}|,$$

where F is the flip operator on the first two tensor factors. To compute $\text{Tr}(I_{110}^\dagger M_3)$, we first remember the property of the flip operator that $\text{Tr}(F\check{\sigma}_a \otimes \check{\sigma}_b) = (\check{\sigma}_a|\check{\sigma}_b) = \delta_{a,b}$, which leads us to

$$\text{Tr}(I_{110}^\dagger |\check{\sigma}_a \otimes \check{\sigma}_b)(\check{\sigma}_c|) = [\text{Tr}(F\check{\sigma}_a \otimes \check{\sigma}_b) - (\check{1}|\check{\sigma}_a)(\check{1}|\check{\sigma}_b)](\check{1}|\check{\sigma}_c) = (\delta_{a,b} - \delta_{a,0}\delta_{b,0})\delta_{c,0}.$$

Therefore, we can write I_{110} in the Pauli basis as $I_{110} = \sum_{a \neq 0} |\check{\sigma}_a \otimes \check{\sigma}_a)(\check{1}|$, from where one can deduce the normalization $\text{Tr}(I_{110}^\dagger I_{110}) = d^2 - 1$. For the overlap with M_3 we have

$$\text{Tr}(I_{110}^\dagger M_3) = \frac{1}{\sqrt{d}} \sum_{z, z' \in \mathbb{F}_2^n} (\delta_{z, z'} - \delta_{z,0}\delta_{z',0})\delta_{z+z',0} = \frac{d-1}{\sqrt{d}}.$$

The contribution to S_3 is then given as

$$\frac{\text{Tr}(I_{110}^\dagger M_3)}{\text{Tr}(I_{110}^\dagger I_{110})} I_{110} = \frac{1}{\sqrt{d}(d+1)} \sum_{a \neq 0} |\check{\sigma}_a \otimes \check{\sigma}_a)(\check{1}|. \quad (\text{D15})$$

(iv) Lastly we take a look at I_{111} , which was also determined in [30, App. C.1b], where it was shown that

$$\text{Tr}(I_{111} M_3) = \frac{\text{Tr}(I_{\text{ad}}^{(1)} M_3)}{d^3(d^2-1)(d^2-4)} (I_{\text{ad}}^{(1)} + I_{\text{ad}}^{(2)}) \quad (\text{D16})$$

with

$$I_{\text{ad}}^{(1)} = \sum_{\substack{a \neq 0, b \neq 0 \\ a \neq b}} |\sigma_a \otimes \sigma_b)(\sigma_a \sigma_b|, \quad I_{\text{ad}}^{(2)} = \sum_{\substack{a \neq 0, b \neq 0 \\ a \neq b}} |\sigma_a \otimes \sigma_b)(\sigma_b \sigma_a|. \quad (\text{D17})$$

Using $\sigma_a \sigma_b = (-1)^{[a,b]} \sigma_b \sigma_a$ and the above definition we find that

$$I_{\text{ad}}^{(1)} + I_{\text{ad}}^{(2)} = 2 \sum_{\substack{a \neq 0, b \neq 0 \\ a \neq b, [a,b]=0}} |\sigma_a \otimes \sigma_b)(\sigma_a \sigma_b|, \quad (\text{D18})$$

as well as

$$\text{Tr}(I_{\text{ad}}^{(1)\dagger} M_3) = d \sum_{z, z' \in \mathbb{F}_2^n} \sum_{\substack{a \neq 0, b \neq 0 \\ a \neq b, [a,b]=0}} (-1)^{\beta(a,b)} \delta_{z,a} \delta_{z',b} \delta_{z+z',a+b} = d \sum_{\substack{z \neq 0, z' \neq 0 \\ z \neq z'}} 1 = d(d-1)(d-2).$$

The contribution to S_3 is thus

$$\frac{\text{Tr}(I_{111}^\dagger M_3)}{\text{Tr}(I_{111}^\dagger I_{111})} I_{111} = \frac{2}{d^2(d+1)(d+2)} \sum_{\substack{a \neq 0, a' \neq 0 \\ a \neq a', [a,a']=0}} |\sigma_a \otimes \sigma_{a'})(\sigma_a \sigma_{a'}| \quad (\text{D19})$$

$$= \frac{2}{\sqrt{d}(d+1)(d+2)} \sum_{\substack{a \neq 0, a' \neq 0 \\ a \neq a', [a,a']=0}} (-1)^{\beta(a,a')} |\check{\sigma}_a \otimes \check{\sigma}_{a'})(\check{\sigma}_{a+a'}|. \quad (\text{D20})$$

By comparing the prefactors of different terms from S_3 given in Eq. (D13), Eq. (D14), Eq. (D15) and Eq. (D20) with Eq. (D3), we can read off the coefficients $s_{a,a'}$:

$$s_{a,a'} = \begin{cases} 1 & a = 0 \vee a' = 0 \\ \frac{1}{d+1} & a \neq 0 \wedge a' \neq 0 \wedge a = a' \\ \frac{2}{(d+1)(d+2)}(-1)^{\beta(a,a')} & a \neq 0 \wedge a' \neq 0 \wedge a \neq a' \wedge [a, a'] = 0 \\ 0 & \text{else.} \end{cases} \quad (\text{D21})$$

For the second moment $\mathbb{E}(\hat{o})$ we need the fractions $\frac{s_{a,a'}}{s_a s'_a}$, which using $s_a = (d+1)^{-1}$ for $a \neq 0$ and $s_0 = 1$ are found to be

$$\frac{s_{a,a'}}{s_a s'_a} = \begin{cases} 1 & a = 0 \vee a' = 0 \\ d+1 & a \neq 0 \wedge a' \neq 0 \wedge a = a' \\ \frac{2(d+1)}{d+2}(-1)^{\beta(a,a')} & a \neq 0 \wedge a' \neq 0 \wedge a \neq a' \wedge [a, a'] = 0 \\ 0 & \text{else.} \end{cases} \quad (\text{D22})$$

□

2. The local Clifford group

For the local Clifford group, we begin by arguing that for uniform sampling and without the presence of noise, the frame operator factorizes. First note that the Z-basis measurement operator factorizes as

$$M_3 = \frac{1}{\sqrt{d}} \left(\sum_{z, z' \in \mathbb{F}_2} |\check{Z}_z \otimes \check{Z}_{z'}\rangle \langle \check{Z}_{z+z'}| \right)^{\otimes n}. \quad (\text{D23})$$

Consequently, we find:

$$S_3 = \frac{d^{-1/2}}{|\text{Cl}_1^{\otimes n}|} \sum_{g \in \text{Cl}_1^{\otimes n}} \omega^\dagger(g)^{\otimes 2} M_3 \omega(g) \quad (\text{D24})$$

$$= \frac{1}{\sqrt{d}} \left(\frac{1}{|\text{Cl}_1|} \sum_{g \in \text{Cl}_1} \sum_{z, z' \in \mathbb{F}_2} \omega^\dagger(g)^{\otimes 2} |\check{Z}_z \otimes \check{Z}_{z'}\rangle \langle \check{Z}_{z+z'}| \omega(g) \right)^{\otimes n} \quad (\text{D25})$$

$$= \frac{1}{\sqrt{d}} \left(\sum_{a, a' \in \mathbb{F}_2^2} s_{a,a'} |\check{\sigma}_a \otimes \check{\sigma}_{a'}\rangle \langle \check{\sigma}_{a+a'}| \right)^{\otimes n}, \quad (\text{D26})$$

where we used Eq. (D3) for the case $n = 1$ in the last line. Since S also factorizes for the local Clifford group, one can easily verify that $S^{-1} \otimes S^{-1} S_3 = \frac{1}{\sqrt{d}} \left(\sum_{a, a' \in \mathbb{F}_2^2} \frac{s_{a,a'}}{s_a s'_a} |\check{\sigma}_a \otimes \check{\sigma}_{a'}\rangle \langle \check{\sigma}_{a+a'}| \right)^{\otimes n}$. The coefficients were already determined in Eq. (D22) and for $n = 1$ we can simplify them to get (the third case cannot occur):

$$\frac{s_{a,a'}}{s_a s'_a} \stackrel{n=1}{=} \begin{cases} 1 & a = 0 \vee a' = 0 \\ 3 & a \neq 0 \wedge a' \neq 0 \wedge a = a' \\ 0 & \text{else.} \end{cases} \quad (\text{D27})$$

For $n > 1$ qubits, recall that we label Pauli operators as $\sigma_a = \sigma_{a_1} \otimes \cdots \otimes \sigma_{a_n}$ where $a_i \in \mathbb{F}_2^2$. The coefficients are then found as products of the above single qubit coefficients:

$$\frac{s_{a,a'}}{s_a s'_a} = \prod_{i=1}^n \frac{s_{a_i, a'_i}}{s_{a_i} s'_{a'_i}} = \begin{cases} 0 & \exists i \in \text{supp}(a) \cap \text{supp}(a') : a_i \neq a'_i \\ 3^{|\text{supp}(a) \cap \text{supp}(a')|} & \text{else.} \end{cases} \quad (\text{D28})$$

3. Variance bound

Theorem 4. *The variance of shadow estimation for uniform sampling from the global Clifford group under gate-dependent noise is bounded by*

$$\mathbb{V}_{\text{global}}[\hat{o}] \leq \frac{2(d+1)}{(d+2)} \|O_0\|_{\text{st}}^2 + \frac{d+1}{d} \|O_0\|_2^2.$$

For uniform sampling from the local Clifford group, the variance is bounded by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 4^k \|O_{\text{loc}}\|_{\infty}^2$ for k -local observables $O = O_{\text{loc}} \otimes \mathbb{1}^{n-k}$ and by $\mathbb{V}_{\text{loc}}[\hat{o}] \leq 3^{\text{supp}(\sigma_a)}$ for Pauli observables $O = \sigma_a$.

Proof. The variance is given by $\mathbb{V}[\hat{o}] = \mathbb{E}[\hat{o} - \mathbb{E}[\hat{o}]]^2$ and $O = O_0 + \text{Tr}(O)\mathbb{1}/d$ as before. We have $(\mathbb{1} | S^{-1}\omega^\dagger(g) | E_x) = 1$ and $(\mathbb{1} | S^{-1}\tilde{S} | \rho) = (\mathbb{1} | \rho) = 1$ since the involved superoperators and the noise is trace preserving. We then find:

$$\hat{o}(g, x) - \mathbb{E}[\hat{o}] = (O | S^{-1}\omega^\dagger(g) | E_x) - (O | S^{-1}\tilde{S} | \rho) = (O_0 | S^{-1}\omega^\dagger(g) | E_x) - (O_0 | S^{-1}\tilde{S} | \rho) = \hat{o}_0(g, x) - \mathbb{E}[\hat{o}_0] \quad (\text{D29})$$

and the variance only depends on the traceless part $\mathbb{V}(\hat{o}) = \mathbb{V}(\hat{o}_0)$. We will now bound the variance for the global Clifford protocol in terms of the second moment: $\mathbb{V}[\hat{o}_0] = \mathbb{E}[\hat{o}_0^2] - (\mathbb{E}[\hat{o}_0])^2 \leq \mathbb{E}[\hat{o}_0^2]$. To rewrite the second moment given in Eq. (D5), we first note that $\bar{\Lambda}_{a,a'} = \bar{\Lambda}_{a',a}$ which can be seen from its definition and from $s_{a,a'} = s_{a',a}$. The second moment for a traceless observable O_0 can then be written as

$$\mathbb{E}(\hat{o}_0^2) = \frac{1}{\sqrt{d}} \sum_{\substack{a \neq 0, a' \neq 0: \\ [a, a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} (O_0 | \check{\sigma}_a) (O_0 | \check{\sigma}_{a'}) (\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho) \quad (\text{D30})$$

$$= \frac{1}{\sqrt{d}} \frac{2(d+1)}{d+2} \sum_{\substack{a \neq 0, a' \neq 0 \\ a \neq a', [a, a'] = 0}} (-1)^{\beta(a,a')} (O_0 | \check{\sigma}_a) (O_0 | \check{\sigma}_{a'}) (\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho) + \frac{d+1}{\sqrt{d}} \sum_{a \neq 0} (O_0 | \check{\sigma}_a)^2 (\check{\mathbb{1}} | \bar{\Lambda}_{a,a} | \rho). \quad (\text{D31})$$

$$(\text{D32})$$

We can bound the sums by using $|(\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho)| \leq \frac{1}{\sqrt{d}}$ to get

$$\mathbb{E}(\hat{o}_0^2) \leq \frac{2(d+1)}{d(d+2)} \sum_{\substack{a \neq 0, a' \neq 0 \\ a \neq a', [a, a'] = 0}} |(O_0 | \check{\sigma}_a) (O_0 | \check{\sigma}_{a'})| + \frac{d+1}{d} \|O_0\|_2^2 \quad (\text{D33})$$

$$\leq \frac{2(d+1)}{d(d+2)} \sum_{a \neq 0} |(O_0 | \check{\sigma}_a)| \sum_{a' \neq 0} |(O_0 | \check{\sigma}_{a'})| + \frac{d+1}{d} \|O_0\|_2^2 \quad (\text{D34})$$

$$= \frac{2(d+1)}{d+2} \|O_0\|_{\text{st}}^2 + \frac{d+1}{d} \|O_0\|_2^2. \quad (\text{D35})$$

To bound the second moment for the local Clifford group, we start at the general form for Clifford protocols given in Eq. (D5) and use again that $|(\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho)| \leq \frac{1}{\sqrt{d}}$. This leads us to

$$\mathbb{E}(\hat{o}^2) = \frac{1}{\sqrt{d}} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} (O | \check{\sigma}_a) (O | \check{\sigma}_{a'}) (\check{\sigma}_{a+a'} | \bar{\Lambda}_{a,a'} | \rho) \leq \frac{1}{d} \sum_{\substack{a, a' \in \mathbb{F}_2^{2n}: \\ [a, a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} |(O | \check{\sigma}_a) (O | \check{\sigma}_{a'})|. \quad (\text{D36})$$

We first consider k -local Pauli observables $P = \sigma_{b_1} \otimes \cdots \otimes \sigma_{b_n}$, where $b_i = 0$ on $n - k$ sites. Then, Eq. (D36) factorizes and using Eq. (D27) we find that

$$\mathbb{E}(\hat{o}_P^2) \leq \frac{1}{d} \prod_{i=1}^n \sum_{\substack{a_i, a'_i \in \mathbb{F}_2^2: \\ [a_i, a'_i] = 0}} \frac{s_{a_i, a'_i}}{s_{a_i} s_{a'_i}} |(\sigma_{b_i} | \check{\sigma}_{a_i}) (\sigma_{b_i} | \check{\sigma}_{a'_i})| = \frac{1}{d} \prod_{i=1}^n \sum_{\substack{a_i, a'_i \in \mathbb{F}_2^2: \\ [a_i, a'_i] = 0}} \frac{s_{a_i, a'_i}}{s_{a_i} s_{a'_i}} d \delta_{a_i, b_i} \delta_{a'_i, b_i} = 3^k. \quad (\text{D37})$$

The proof for k -local observables $O = O_{\text{loc}} \otimes \mathbb{1}^{\otimes(n-k)}$ is almost identical to the proof without noise given in Huang *et al.* [1], and we repeat the relevant steps here. Akin to Eq. (D37) we find that

$$\mathbb{E}(\hat{o}^2) \leq \frac{1}{2^k} \sum_{\substack{a, a' \in \mathbb{F}_2^{2k}: \\ [a, a'] = 0}} \frac{s_{a,a'}}{s_a s_{a'}} |(O_{\text{loc}} | \check{\sigma}_a) (O_{\text{loc}} | \check{\sigma}_{a'})|. \quad (\text{D38})$$

In the following, we drop the condition $[a, a'] = 0$ in the sum since we have $s_{a,a'} = 0$ for $[a, a'] \neq 0$. Let $\mathbb{F}_2^{2*} = \mathbb{F}_2^2 \setminus \{0\}$ label the single-qubit non-identity Paulis and define a partial order on $\mathbb{F}_2^{2k}$ such that $a \leq b$ iff for all $i \in [k]$ either $a_i = b_i$ or $a_i = 0$ holds. By Eq. (D28), pairs $a, a' \in \mathbb{F}_2^{2k}$ which do not coincide on their common support do not contribute to the sum in Eq. (D38). For the remaining pairs, we can always find a $b \in (\mathbb{F}_2^{2*})^k$ such that $a \leq b$ and $a' \leq b$. We can thus replace the sum in Eq. (D38) as $\sum_{a,a' \in \mathbb{F}_2^{2k}} \rightarrow \sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{a \leq b} \sum_{a' \leq b}$ if we take care of potential over counting. To this end, note that we are free to choose b_i for every index i where $a_i = a'_i = 0$ and there are in total $k - |\text{supp}(a) \cup \text{supp}(a')| = k - |\text{supp}(a)| - |\text{supp}(a')| + |\text{supp}(a) \cap \text{supp}(a')|$ of such indices. Using $\frac{s_{a,a'}}{s_a s_{a'}} = 3^{|\text{supp}(a) \cap \text{supp}(a')|}$, we find:

$$\frac{1}{2^k} \sum_{a,a' \in \mathbb{F}_2^{2k}} \frac{s_{a,a'}}{s_a s_{a'}} |(O_{\text{loc}} | \check{\sigma}_a)(O_{\text{loc}} | \check{\sigma}_{a'})| = \frac{1}{2^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{\substack{a \leq b \\ a' \leq b}} \frac{3^{|\text{supp}(a)| + |\text{supp}(a')|}}{3^k} |(O_{\text{loc}} | \check{\sigma}_a)(O_{\text{loc}} | \check{\sigma}_{a'})| \quad (\text{D39})$$

$$= \frac{1}{2^k 3^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \left(\sum_{a \leq b} 3^{|\text{supp}(a)|} |(O_{\text{loc}} | \check{\sigma}_a)| \right)^2. \quad (\text{D40})$$

With $\sum_{a \leq b} 3^{|\text{supp}(a)|} = 4^k$ for all $b \in (\mathbb{F}_2^{2*})^k$ [1] and the Cauchy-Schwarz inequality, this can be simplified as follows:

$$\frac{1}{2^k 3^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \left(\sum_{a \leq b} 3^{|\text{supp}(a)|} |(O_{\text{loc}} | \check{\sigma}_a)| \right)^2 \leq \frac{1}{2^k 3^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{a \leq b} 3^{|\text{supp}(a)|} \sum_{a' \leq b} 3^{|\text{supp}(a')|} (O_{\text{loc}} | \check{\sigma}_{a'})^2 \quad (\text{D41})$$

$$= \frac{4^k}{2^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{a' \leq b} \frac{3^{|\text{supp}(a')|}}{3^k} (O_{\text{loc}} | \check{\sigma}_{a'})^2. \quad (\text{D42})$$

Note that in the last line, the number of times each $a' \in \mathbb{F}_2^{2k}$ appears in the double sum is given by $3^{k-|\text{supp}(a')|}$ and therefore $\sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{a' \leq b} \frac{3^{|\text{supp}(a')|}}{3^k} f(a') = \sum_{a' \in \mathbb{F}_4^k} f(a')$ for arbitrary summands $f(a')$. This leads us to the final result

$$\frac{4^k}{2^k} \sum_{b \in (\mathbb{F}_2^{2*})^k} \sum_{a' \leq b} \frac{3^{|\text{supp}(a')|}}{3^k} (O_{\text{loc}} | \check{\sigma}_{a'})^2 = \frac{4^k}{2^k} \sum_{a' \in \mathbb{F}_4^k} (O_{\text{loc}} | \check{\sigma}_{a'})^2 = 2^k \|O_{\text{loc}}\|_2^2 \leq 4^k \|O_{\text{loc}}\|_\infty^2. \quad (\text{D43})$$

□

Appendix E: Robust shadow estimation under gate-dependent noise

A prominent noise mitigation technique for shadow estimation are the *robust classical shadows* developed by Chen *et al.* [12]. Robust classical shadows rely on the assumption of gate-independent left-noise, i.e.

$$\phi(g) = \Lambda \omega(g) \quad (\text{E1})$$

for a g -independent channel Λ . In this case, the frame operator can be written as $\tilde{S} = \sum_{\lambda \in \text{Irr}(G)} f_\lambda \Pi_\lambda$, where $\text{Irr}(G)$ is the set of irreducible representations of the group G and Π_λ are projectors onto invariant subspaces. For the global Clifford group and trace preserving noise, this reduces to $\tilde{S} = |\mathbb{1}\rangle(\mathbb{1}| + f \sum_{a \neq 0} |\check{\sigma}_a\rangle\langle \check{\sigma}_a|)$, meaning that a single parameter f needs to be estimated in order to have full knowledge of $\tilde{S}$, which allows for the mitigation with $\tilde{S}^{-1}$ in post-processing. In Chen *et al.* [12], f is determined by the median of means of the single shot estimator

$$\hat{f}(g, x) := (d(E_x | \omega(g) | E_0) - 1) / (d - 1) \quad (\text{E2})$$

where for each round $|0\rangle$ state is prepared, a random operation $g \in G$ is applied and the result x is stored. It is then shown that the expectation value for gate-independent left noise satisfies $\mathbb{E}[\hat{f}(g, x)] = f$.

We will now show how gate-independent left noise can be written in our gate-dependent right noise model and the effect it has on the noisy frame operator. This provides an understanding as to how left- and right noise interrelate and shows consistency of our formalism with previous works on gate independent noise [12, 13]. The assumption (E1)

translates to $\phi(g) = \omega(g)\omega^\dagger(g)\Lambda\omega(g) = \omega(g)\Lambda(g)$ for $\Lambda(g) = \omega^\dagger(g)\Lambda\omega(g)$. The average noise channels $\bar{\Lambda}_a$ of Lemma 2 for uniform sampling from the global Clifford group are given by

$$\bar{\Lambda}_a = \frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a)} \Lambda(g) = \frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a)} \omega^\dagger(g)\Lambda\omega(g). \quad (\text{E3})$$

We know that for every $g \in \Xi_z^{-1}(a)$, $hg \in \Xi_z^{-1}(a)$ with $h \in \mathcal{P}_n$, since

$$\omega^\dagger(hg)|\check{Z}_z\rangle(\check{Z}_z|\omega(hg) = \omega^\dagger(g)\omega^\dagger(h)|\check{Z}_z\rangle(\check{Z}_z|\omega(h)\omega(g) = \omega^\dagger(g)|\check{Z}_z\rangle(\check{Z}_z|\omega(g). \quad (\text{E4})$$

This means that the average in Eq. (E3) contains an average $\frac{1}{|\mathcal{P}_n|} \sum_{h \in \mathcal{P}_n} \omega^\dagger(h)\Lambda\omega(h)$, which is commonly referred to as a Pauli twirl and the result is a Pauli channel (see e.g. [31]). In the following we will restrict ourselves to $a \neq 0$ where we will need the inverse relation

$$\frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a)} \omega(g)|\check{\sigma}_a\rangle(\check{\sigma}_a|\omega(g)^\dagger = \frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} |\Xi_z^{-1}(a)| |\check{Z}_z\rangle(\check{Z}_z| = \frac{1}{d-1} \sum_{z \in \mathbb{F}_2^n} |\check{Z}_z\rangle(\check{Z}_z|, \quad (\text{E5})$$

which can be verified using Proposition 10. The diagonal entries of $\bar{\Lambda}_a$ for $a \neq 0$ can then be determined by

$$\frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a)} (\check{\sigma}_a|\omega(g)^\dagger\Lambda\omega(g)|\check{\sigma}_a) = \frac{(d+1)}{|\text{Cl}_n|} \sum_{z \in \mathbb{F}_2^n} \sum_{g \in \Xi_z^{-1}(a)} \text{Tr} [\Lambda\omega(g)|\check{\sigma}_a\rangle(\check{\sigma}_a|\omega(g)^\dagger] \quad (\text{E6})$$

$$= \frac{1}{d-1} \text{Tr} \left[\Lambda \sum_{z \in \mathbb{F}_2^n \setminus 0} |\check{Z}_z\rangle(\check{Z}_z| \right]. \quad (\text{E7})$$

In summary, the average channels $\bar{\Lambda}_a$ are diagonal since they are Pauli channels, and they all share the same diagonal elements $(\check{\sigma}_a|\bar{\Lambda}_a|\check{\sigma}_a)$, which are independent of a . Note that $(\check{\mathbb{I}}|\Lambda|\check{\mathbb{I}}) = 1$ for trace preserving or unital noise and let $f = \text{Tr} [\Lambda \sum_{z \in \mathbb{F}_2^n \setminus 0} |\check{Z}_z\rangle(\check{Z}_z|] / (d-1) = (\text{Tr}[\Lambda M] - 1) / (d-1)$. Then the noisy frame operator is given by

$$\tilde{S} = |\check{\mathbb{I}}\rangle(\check{\mathbb{I}}| + \frac{1}{d+1} \sum_{a \neq 0} |\check{\sigma}_a\rangle(\check{\sigma}_a|\bar{\Lambda}_a = |\check{\mathbb{I}}\rangle(\check{\mathbb{I}}| + \frac{f}{d+1} \sum_{a \neq 0} |\check{\sigma}_a\rangle(\check{\sigma}_a| \quad (\text{E8})$$

which is the result of [12] for the global Clifford group.

We will now turn our attention back to gate-dependent noise with the following lemma. First, we define $Z^n := \{00, 01\}^n$, the index set for diagonal Paulis.

Lemma 16. *The mitigation parameter in the robust shadow estimation protocol for uniform sampling over the global Clifford group under gate-dependent Pauli noise is given by*

$$\mathbb{E}_{g,x} [\hat{f}(g, x)] = \frac{1}{d+1} \mathbb{E}_{a \in Z^n \setminus 0} \bar{\lambda}_a. \quad (\text{E9})$$

Proof. We need to compute the expectation value of the single shot estimator $\hat{f}(g, x) = \frac{d \cdot (E_x|\omega(g)|E_0) - 1}{d-1}$. In the first step we rewrite the expectation value of $(E_x|\omega(g)|E_0)$.

$$\begin{aligned} \mathbb{E}_{g,x} (E_x|\omega(g)|E_0) &= \frac{1}{|\text{Cl}_n|} \sum_{g \in \text{Cl}_n} \sum_{x \in \mathbb{F}_2^n} (E_x|\omega(g)|E_0)(E_x|\omega(g)\Lambda(g)|E_0) \\ &= \frac{1}{|\text{Cl}_n|} \sum_{g \in \text{Cl}_n} \sum_{x \in \mathbb{F}_2^n} (E_0|\omega(g)^\dagger|E_x)(E_x|\omega(g)\Lambda(g)|E_0) \\ &= (E_0|\tilde{S}|E_0). \end{aligned}$$

By noting that $|E_0\rangle = \left(\frac{1}{\sqrt{2}}(|\check{\mathbb{I}}\rangle + |\check{Z}\rangle) \right)^{\otimes n} = \frac{1}{\sqrt{d}} \sum_{z \in \mathbb{F}_2^n} |\check{Z}(z)\rangle$ we have $(E_0|\tilde{S}|E_0) = \frac{1}{d} \sum_{z, z' \in \mathbb{F}_2^n} (\check{Z}(z)|\tilde{S}|\check{Z}(z'))$. With the use of Lemma 2 and $s_a = 1/(d+1)$ for $a \neq 0$ and $s_0 = 1$, we find:

$$\mathbb{E}_{g,x} (E_x|\omega(g)|E_0) = \frac{1}{d} + \frac{1}{d(d+1)} \sum_{a \in Z^n \setminus 0} \bar{\lambda}_a.$$

It follows that

$$\mathbb{E}_{g,x} \hat{f}(g, x) = \frac{d \mathbb{E}_{g,x} (E_x |\omega(g)| E_0) - 1}{d-1} = \frac{1}{(d+1)(d-1)} \sum_{a \in \mathbb{Z}^n \setminus 0} \bar{\lambda}_a = \frac{1}{d+1} \mathbb{E}_{a \in \mathbb{Z}^n \setminus 0} \bar{\lambda}_a \quad (\text{E10})$$

□

Since Eq. (E9) shows that for Pauli noise $\mathbb{E}[\hat{f}(g, x)]$ always contains the factor $1/(d+1)$, we define the mitigation parameter as $\hat{f}_m := (d+1) \mathbb{E}[\hat{f}(g, x)]$. The final estimate of the robust shadow estimation procedure is then given by

$$\mathbb{E}[\hat{o}_{RS}] := (O | (\tilde{\mathbb{I}})(\tilde{\mathbb{I}} + \frac{d+1}{\hat{f}_m} \sum_{a \neq 0} |\check{\sigma}_a)(\check{\sigma}_a|) \tilde{S} | \rho). \quad (\text{E11})$$

We will now show that whether the robust shadow estimation strategy succeeds for gate-dependent noise is highly dependent on the initial state, the observable and the noise present in the experiment. As can be seen via a simple example, the error can actually be dramatically increased if the noise is not of left-gate-independent form.

Proposition 5. *Under gate-dependent local noise $\phi_\epsilon(g) = (1-\epsilon)\omega(g) + \epsilon\omega(g)\Lambda(g)$, robust shadow estimation with the global Clifford group can introduce a bias $|\mathbb{E}[\hat{o}] - \langle O \rangle| \geq |\langle O_0 \rangle (\frac{1}{2}(1+\epsilon)^n - 1)|$.*

Proof. Consider the local bit flip channel $\Lambda(g) = \mathcal{X}$ for all $g \in \text{Cl}_n$ and let Λ_ϵ be the bit flip channel with error probability ϵ , i.e. $\Lambda_\epsilon = (1-\epsilon)\text{id} + \epsilon\Lambda$. Its process matrix is given by $\Lambda_\epsilon = |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + |\tilde{X}\rangle\langle\tilde{X}| + (1-2\epsilon)(|\tilde{Y}\rangle\langle\tilde{Y}| + |\tilde{Z}\rangle\langle\tilde{Z}|)$. If this Pauli channel acts on all qubits before each Clifford gate, then the global implementation map is $\phi_\epsilon(g) = \omega(g)\Lambda_\epsilon^{\otimes n}$. Since this right noise channel is gate-independent, per Lemma 2 we get $\tilde{S} = (|\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + \frac{1}{d+1} \sum_{a \neq 0} |\check{\sigma}_a)(\check{\sigma}_a|) \cdot \Lambda_\epsilon^{\otimes n}$ and $S^{-1}\tilde{S} = \Lambda_\epsilon^{\otimes n}$. The robust shadow mitigation parameter $\mathbb{E}_{g,x} \hat{f}(g, x)$ is per Lemma 16 given by

$$\mathbb{E}_{g,x} \hat{f}(g, x) = \frac{1}{(d-1)(d+1)} \sum_{a \in \mathbb{Z}^n \setminus 0} (\check{\sigma}_a | \Lambda_\epsilon^{\otimes n} | \check{\sigma}_a) \quad (\text{E12})$$

$$= \frac{1}{(d-1)(d+1)} \left(\sum_{z \in \mathbb{F}_2^n} 1^{n-|z|} (1-2\epsilon)^{|z|} - 1 \right) \quad (\text{E13})$$

$$= \frac{1}{(d-1)(d+1)} \left(\sum_{i=1}^n \binom{n}{i} 1^{n-i} (1-2\epsilon)^i - 1 \right) \quad (\text{E14})$$

$$= \frac{1}{(d+1)(d-1)} (d(1-\epsilon)^n - 1). \quad (\text{E15})$$

In the robust shadow estimation procedure, the inverse frame operator is then given by

$$S_{RS}^{-1} = |\tilde{\mathbb{I}}\rangle\langle\tilde{\mathbb{I}}| + (d+1) \frac{d-1}{d(1-\epsilon)^n - 1} \sum_{a \neq 0} |\check{\sigma}_a)(\check{\sigma}_a|, \quad (\text{E16})$$

and the expected outcome is

$$\mathbb{E}[\hat{o}_{RS}] = (O | S_{RS}^{-1} \tilde{S} | \rho) = \frac{\text{Tr}(O)}{d} + \frac{d-1}{d(1-\epsilon)^n - 1} \sum_{a \neq 0} (O | \check{\sigma}_a)(\check{\sigma}_a | \Lambda_\epsilon^{\otimes n} | \rho). \quad (\text{E17})$$

Consequently, errors on Y-type and Z-type Pauli observables are partially mitigated while previously absent errors on X-Pauli observables are introduced. In particular, for the stabilizer state $O = (|+\rangle\langle+|)^{\otimes n} = \frac{1}{\sqrt{d}} \sum_x |\check{X}^x\rangle$, we find that $\mathbb{E}[\hat{o}] = \langle + | \rho | + \rangle$, meaning the standard shadow estimate is unbiased. Moreover, we have $(O | \Lambda_\epsilon^{\otimes n} = (O |$ and thus

$$\mathbb{E}[\hat{o}_{RS}] - \frac{\text{Tr}(O)}{d} = \frac{d-1}{d(1-\epsilon)^n - 1} \left((O | \Lambda_\epsilon^{\otimes n} | \rho) - \frac{\text{Tr}(O)}{d} \right) = \frac{d-1}{d(1-\epsilon)^n - 1} \langle O_0 \rangle, \quad (\text{E18})$$

using the unitality of $\Lambda_\epsilon^{\otimes n}$ and $\langle O_0 \rangle = (O_0 | \rho)$. Hence,

$$\mathbb{E}[\hat{o}_{RS}] - \frac{\text{Tr}(O)}{d} = \frac{d-1}{d(1-\epsilon)^n - 1} \langle O_0 \rangle > \frac{d-1}{d(1-\epsilon)^n} \langle O_0 \rangle \geq \frac{d-1}{d} (1+\epsilon)^n \langle O_0 \rangle, \quad (\text{E19})$$

where we assumed $d(1-\epsilon)^n - 1 > 1$. Therefore $|\mathbb{E}[\hat{o}_{RS}] - \langle O \rangle| \geq |\langle O_0 \rangle| (\frac{d-1}{d} (1+\epsilon)^n - 1) \geq |\langle O_0 \rangle| (\frac{1}{2} (1+\epsilon)^n - 1)$. □

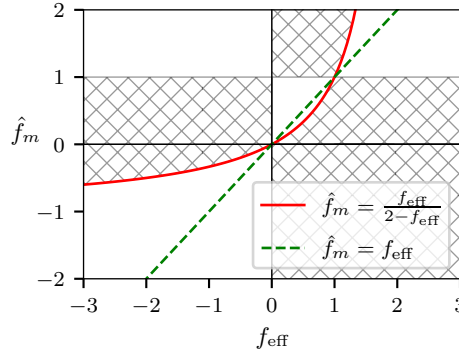


Figure 1. Parameter regions where robust shadow estimation increases the bias (hatched) and regions where it decreases the bias (plain). Perfect mitigation is achieved on the green dashed line.

Appendix F: Bias mitigation conditions

In this section, we derive sufficient conditions for the robust shadow protocol to work even in the presence of gate-dependent Pauli noise. The ultimate aim is to determine under which conditions the estimation bias decreases, i.e., $|\mathbb{E}[\hat{o}_{\text{RS}}] - \langle O \rangle| \leq |\mathbb{E}[\hat{o}] - \langle O \rangle|$. Here, we characterize conditions of overcorrection (Figure 1) and prove that isotropic Pauli noise is well-conditioned in the sense that robust shadow estimation ‘works well’ (Proposition 20). Let $\bar{\lambda} = \sum_{a \neq 0} \bar{\lambda}_a / (d^2 - 1)$ be the mean of $\bar{\lambda}_a$ (excluding λ_0) and O_0, ρ_0 be the traceless parts of O and ρ respectively. We further define $\mathcal{D}(\Delta)$ to be the diagonal channel in Pauli basis with values Δ_a on the diagonal. The following observation provides insight onto how the bias for gate-dependent Pauli noise is determined by a single parameter f_{eff} .

Observation 17. *Let the average Pauli eigenvalues $\{\bar{\lambda}_a\}$ be parameterized as $\bar{\lambda}_a = \bar{\lambda} + \Delta_a$ for $a \neq 0$, and $\lambda_0 = 1$ (trace preserving noise). Then the bias is determined by the quantity $\bar{\lambda} + \frac{(O_0|\mathcal{D}(\Delta)|\rho_0)}{\langle O_0 \rangle} =: f_{\text{eff}}$ as*

$$|\mathbb{E}[\hat{o}_{\text{RS}}] - \langle O \rangle| = |\langle O_0 \rangle| \cdot |1 - \hat{f}_m^{-1} f_{\text{eff}}| \quad \text{and} \quad |\mathbb{E}[\hat{o}] - \langle O \rangle| = |\langle O_0 \rangle| \cdot |1 - f_{\text{eff}}|. \quad (\text{F1})$$

Why this holds is shown in the proof of Proposition 18 below. The difference to the case of left gate-independent noise lies in the additional term $\frac{(O_0|\mathcal{D}(\Delta)|\rho_0)}{\langle O_0 \rangle}$ that quantifies how Pauli noise aligns with the signal $(O_0|\rho_0)$. Perfect bias mitigation is achieved for $\hat{f}_m = f_{\text{eff}}$. If the additional term is small, then a feasible strategy is to just estimate $\bar{\lambda}$ and use it as a mitigation parameter. This is essentially what robust shadow estimation does, albeit by only taking the average over $a \in \mathbb{Z}^n$ since $\hat{f}_m^{-1} = \mathbb{E}_{a \in \mathbb{Z}^n} \bar{\lambda}_a$.

In Figure 1 areas in parameter space corresponding to $|\langle \hat{o}_{\text{RS}} \rangle - \langle O \rangle| \leq |\langle \hat{o} \rangle - \langle O \rangle|$ are visualized. Here, we allow $|\hat{f}_m| > 1$, which is not the case if $\hat{f}_m$ is estimated according to the robust shadow estimation protocol, but still instructive for large $\frac{(O_0|\mathcal{D}(\Delta)|\rho_0)}{\langle O_0 \rangle}$. We can see that for $f_{\text{eff}} \leq 0$ and $f_{\text{eff}} \geq 2$, the error is always reduced by setting $\hat{f}_m \rightarrow \infty$, meaning we return $\hat{o}_{\text{RS}} = 0$.

This is the case for $(O_0|\mathcal{D}(\Delta)|\rho_0) \gg \langle O_0 \rangle$, i.e. large error alignment and small true expectation value $\langle O_0 \rangle$. In the most relevant parameter regime of $0 \leq f_{\text{eff}} \leq 1$, we obtain the condition $\hat{f} \geq \frac{f_{\text{eff}}}{2 - f_{\text{eff}}}$ that safeguards against introducing new errors due to overcorrection. A formal treatment of the success criteria for robust shadow estimation is given in the following proposition. We denote the canonical inner product by $\langle \gamma, \Delta \rangle := \sum_a \gamma_a^* \Delta_a$.

Proposition 18. *Let $\hat{O}_{\text{RS}}$ be the robust shadow estimate on $O = \sum_a (\check{\sigma}_a | O) |\check{\sigma}_a\rangle$, $\rho = \sum_a (\check{\sigma}_a | \rho) |\check{\sigma}_a\rangle$ with mitigation parameter $f_m \in [-1, 1]$ and let $\gamma_a = (O | \check{\sigma}_a) (\check{\sigma}_a | \rho)$ for $a \neq 0$. Under gate-dependent Pauli noise parameterized as $\bar{\lambda}_a = \bar{\lambda} + \Delta_a$ with $\lambda_1 = 1$, it holds that $|\mathbb{E}[\hat{o}_{\text{RS}}] - \langle O \rangle| \leq |\mathbb{E}[\hat{o}] - \langle O \rangle|$ iff*

$$\left(f_{\text{eff}} \geq 0 \wedge f_m \geq \frac{f_{\text{eff}}}{2 - f_{\text{eff}}} \right) \quad \vee \quad \left(f_{\text{eff}} < 0 \wedge f_m \leq \frac{f_{\text{eff}}}{2 - f_{\text{eff}}} \right) \quad (\text{F2})$$

for $f_{\text{eff}} := \bar{\lambda} + \frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle}$.

Proof. The error of non-mitigated shadow estimation under Pauli noise is given in Eq. (A24), from which we gather

that

$$|\mathbb{E}[\hat{\sigma}] - \langle O \rangle| = \left| \sum_{a \neq 0} (O|\tilde{\sigma}_a)(\tilde{\sigma}_a|\rho)|1 - \lambda_a| \right| = \left| \sum_{a \neq 0} \gamma_a |1 - \lambda_a| \right| = |\langle \gamma, \mathbf{1} - \bar{\lambda} \mathbf{1} - \Delta \rangle| \quad (\text{F3})$$

per the assumption that $\lambda_1 = 1$ and $\lambda_{a \neq 0} = \bar{\lambda} + \Delta_a$, where $\mathbf{1}$ is the vector of which all entries are 1. When the mitigation factor is included we obtain $\hat{S}^{-1}\tilde{S} = |\tilde{\mathbf{1}}\rangle(\tilde{\mathbf{1}}| + \sum_{a \neq 0} f_m^{-1} \lambda_a |\tilde{\sigma}_a\rangle(\tilde{\sigma}_a|$, and the mitigated error becomes

$$|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| = \left| \sum_{a \neq 0} \gamma_a (1 - f_m^{-1} \lambda_a) \right| = |\langle \gamma, \mathbf{1} - f_m^{-1} \bar{\lambda} \mathbf{1} - f_m^{-1} \Delta \rangle|. \quad (\text{F4})$$

Since $\langle \gamma, \mathbf{1} \rangle = \sum_{a \neq 0} \gamma_a$ is by definition the expectation value of the traceless part of the observable, $\langle O_0 \rangle$, we can simplify the error terms by looking at the relative errors:

$$\frac{|\mathbb{E}[\hat{\sigma}] - \langle O \rangle|}{|\langle O_0 \rangle|} = \left| 1 - \left(\bar{\lambda} + \frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle} \right) \right| = |1 - f_{\text{eff}}| \quad (\text{F5})$$

and

$$\frac{|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle|}{|\langle O_0 \rangle|} = \left| 1 - f_m^{-1} \left(\bar{\lambda} + \frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle} \right) \right| = |1 - f_m^{-1} f_{\text{eff}}|. \quad (\text{F6})$$

To determine when $|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| \leq \epsilon$ holds, we have to look at four cases corresponding the signs of f_{eff} and f_m . We also assume that $|f_m| \leq 1$ since $|\bar{\lambda}| \leq 1$ for any physical noise model.

(i) f_{eff} and f_m are of opposite sign:

In this case we have that $f_m^{-1} f_{\text{eff}} \leq 0$ and therefore

$$|\mathbb{E}[\hat{\sigma}] - \langle O \rangle| = |1 - f_{\text{eff}}| \leq 1 + |f_{\text{eff}}| \leq 1 + |f_m^{-1} f_{\text{eff}}| = |1 - f_m^{-1} f_{\text{eff}}| = |\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| \quad (\text{F7})$$

and the error mitigation technique always increases the error.

(ii) $f_{\text{eff}} \leq 0$ and $f_m \leq 0$:

Let $f_m \leq \frac{f_{\text{eff}}}{2 - f_{\text{eff}}}$, then $f_m^{-1} \leq \frac{2 - f_{\text{eff}}}{f_{\text{eff}}}$ and $-f_m^{-1} f_{\text{eff}} \leq f_{\text{eff}} - 2$. Therefore, $|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| = |1 - f_m^{-1} f_{\text{eff}}| \leq |-1 + f_{\text{eff}}| = |\mathbb{E}[\hat{\sigma}] - \langle O \rangle|$ and error mitigation is achieved. The condition is tight, since $f_m > \frac{f_{\text{eff}}}{2 - f_{\text{eff}}}$ leads in the same fashion to $|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| \geq |\mathbb{E}[\hat{\sigma}] - \langle O \rangle|$.

(iii) $f_{\text{eff}} > 0$ and $f_m > 0$:

Let now $f_m \geq \frac{f_{\text{eff}}}{2 - f_{\text{eff}}}$. We obtain $f_m^{-1} \geq \frac{2 - f_{\text{eff}}}{f_{\text{eff}}}$ and $-f_m^{-1} f_{\text{eff}} \leq f_{\text{eff}} - 2$, as $-f_{\text{eff}}$ is negative. Thus, again $|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| = |1 - f_m^{-1} f_{\text{eff}}| \leq |-1 + f_{\text{eff}}| = |\mathbb{E}[\hat{\sigma}] - \langle O \rangle|$ and errors are mitigated. The condition is again tight, as $f_m > \frac{f_{\text{eff}}}{2 - f_{\text{eff}}}$ leads to $|\mathbb{E}[\hat{\sigma}_{\text{RS}}] - \langle O \rangle| \geq |\mathbb{E}[\hat{\sigma}] - \langle O \rangle|$.

□

To make a statement about f_{eff} , without assuming explicit knowledge of the noise model and the prepared state, we treat Δ as a random vector. The following lemma gives a concentration inequality for the well known fact that in high dimensions, a uniformly distributed random vector on the unit sphere is almost orthogonal to any given fixed vector with high probability.

Lemma 19 (Adapted from [32]). *For an arbitrary normalized vector $x \in \mathbb{R}^k$ and a random vector g which is uniformly distributed on the unit sphere, holds that $\mathbb{P}(|\langle x, g \rangle| \geq \frac{t}{\sqrt{k-1}}) \leq e^{-t^2/2}$.*

Proof. If g is uniformly distributed, then also the random vectors Og for $O \in \text{SO}(k)$ are distributed uniformly. This implies that $\mathbb{P}(|\langle x, g \rangle| \geq \frac{t}{\sqrt{k-1}})$ is independent of x and w.l.o.g. we choose x to be the first canonical basis vector, $x = e_1$.

The surface area of a k -dimensional unit sphere with radius r is given by $A_S(k, r) = \frac{2\pi^{k/2} r^{k-1}}{\Gamma(\frac{k}{2})}$. Since g is uniformly distributed on the unit sphere, the probability that $|\langle e_1, g \rangle| = |g_1|$ is larger than $\frac{t}{\sqrt{k-1}}$ is given by the ratio of the

surface area of two spherical caps of height $1 - \frac{t}{\sqrt{k-1}}$ to the surface area $A_S(k, 1)$ of the unit sphere. The base of these caps has radius $a = \sqrt{1 - \frac{t^2}{k-1}}$, and we can bound the surface area of the two caps by the surface area of the full sphere of radius a as $2A_{\text{CAP}}(k, a) \leq A_S(k, a)$. This leads us to the bound

$$\mathbb{P} \left[|\langle e_1, g \rangle| \geq \frac{t}{\sqrt{k-1}} \right] = \frac{2A_{\text{CAP}}(k, r)}{A_S(k, 1)} \leq \frac{A_S(k, a)}{A_S(k, 1)} = \left(1 - \frac{t^2}{k-1} \right)^{\frac{k-1}{2}} \leq e^{-t^2/2}. \quad (\text{F8})$$

□

We choose this bound for simplicity, slightly stronger versions can be found in e.g. Dasgupta and Gupta [33] and references therein. An example of uniformly distributed vectors on the sphere is given by normalized standard Gaussian vectors where each entry is drawn from the 0-mean and unit variance normal distribution $\mathcal{N}(0, 1)$.

If the term $\frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle}$ in $f_{\text{eff}} = \bar{\lambda} + \frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle}$ is small, then it is sufficient to estimate $\bar{\lambda}$ to correct the bias. In Proposition 20 below, we show that this is the case under reasonable assumptions.

Proposition 20 (Restatement of Proposition 6). *Let $\langle O_0 \rangle \geq C_1$ and $\|O\|_2 < C_2$. If Δ is a random vector such that $\Delta/\|\Delta\|_{\ell_2}$ is uniformly distributed on the unit sphere and $\|\Delta\|_{\ell_2} \leq \mathcal{O}(d)$ with high probability, then $f_{\text{eff}} = \mathbb{E}[\hat{f}_m] + \mathcal{O}(1/\sqrt{d})$ with high probability.*

Proof. If Δ is uniformly distributed on the unit sphere, then its entries Δ_a are zero mean random variables. Therefore,

$$\mathbb{E}[\hat{f}_m] = \bar{\lambda} + \frac{1}{d-1} \sum_{a \in \mathbb{Z}^n \setminus \mathbf{1}} \mathbb{E}[\bar{\lambda}_a] = \bar{\lambda} + \frac{1}{d-1} \sum_{a \in \mathbb{Z}^n \setminus \mathbf{1}} \mathbb{E}[\Delta^a] = \bar{\lambda} \quad (\text{F9})$$

and it remains to bound $\frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle}$. We write the inner product as $\langle \gamma, \Delta \rangle = \|\gamma\|_{\ell_2} \|\Delta\|_{\ell_2} \langle \tilde{\gamma}, \tilde{\Delta} \rangle$ with $\tilde{\gamma}, \tilde{\Delta}$ normalized. Then we know from Lemma 19 for $k = d^2 - 1$ that

$$\mathbb{P} \left[|\langle \tilde{\gamma}, \tilde{\Delta} \rangle| \geq \frac{t}{\sqrt{d^2 - 2}} \right] \leq e^{-t^2/2}. \quad (\text{F10})$$

From the assumption that $\|\Delta\|_{\ell_2} \leq \mathcal{O}(d)$ and the union bound leads to $\|\Delta\|_{\ell_2} |\langle \tilde{\gamma}, \tilde{\Delta} \rangle| \leq \mathcal{O}(1)$ with high probability. Since $|\langle \rho | \tilde{\sigma}_a \rangle| \leq 1/\sqrt{d}$ for any quantum state ρ , we further have that $\|\gamma\|_{\ell_2} \leq \|O\|_2/\sqrt{d} \leq C/\sqrt{d}$ and thus $\frac{\langle \gamma, \Delta \rangle}{\langle O_0 \rangle} = \mathcal{O}(1/\sqrt{d})$ with high probability. □

Proposition 20 formalizes the intuition that in large systems noise is unlikely to be malicious, i.e. a randomly distributed noise vector Δ is unlikely to align with the signal γ .

The Pauli eigenvalue average $\bar{\lambda}$ can also be related to the average gate fidelity of the frame operator $\tilde{S}$. Using the known relation $F_{\text{Avg}}(\tilde{S}) = (d^{-1} \text{Tr}[\tilde{S}] + 1)/(d+1)$ and $\text{Tr}[\tilde{S}] = 1 + \sum_{a \neq 0} \lambda_a/(d+1)$ [34]. A quick rearrangement yields $\bar{\lambda} = (dF_{\text{Avg}}(\tilde{S}))(d+1)/(d-1)$, suggesting that an estimate of the average gate fidelity $F_{\text{Avg}}(\tilde{S})$ would also give us an estimate of $\bar{\lambda}$.

For a detailed bound in terms of a given error probability δ , we now look to the example of a Gaussian random noise vector. This random vector provides a concrete example for a noise distribution that satisfies the requirements of Proposition 20.

Proposition 21. *Let Δ be a Gaussian random vector of length $k = d^2 - 1$ with i.i.d. entries from $\mathcal{N}(0, \sigma^2)$ and γ be an arbitrary real vector of the same dimension. It holds that*

$$\mathbb{P} \left[|\langle \gamma, \Delta \rangle| \leq \sigma^2 \|\gamma\|_{\ell_2} g(\delta) \sqrt{\frac{k}{k-1}} \left(1 + \frac{g(\delta)}{\sqrt{k}} + \frac{g^2(\delta)}{k} \right) \right] \geq 1 - \delta,$$

where $g(\delta) = \sqrt{\log(2/\delta)}$.

Proof. We begin with a bound for $\|\Delta\|_{\ell_2}^2/\sigma^4 = \|X\|_{\ell_2}^2 = \sum_{i=1}^k x_i^2$, where now $x_i \sim \mathcal{N}(0, 1)$ and therefore $\|X\|_{\ell_2}^2$ is a Chi-squared distributed random variable. We can then use the Laurent-Massart inequality which states that

$\mathbb{P}(\|X\|_{\ell_2}^2 - k \geq 2\sqrt{kt} + 2t) \leq e^{-t}$. We hence obtain $\mathbb{P}(\|X\|_{\ell_2} \geq \sqrt{k}\sqrt{1 + 2\sqrt{t/k} + 2t/k}) \leq e^{-t}$, and we can use $\sqrt{1 + 2\sqrt{t/k} + 2t/k} \leq 1 + \sqrt{t/k} + t/k$. By fixing the failure probability to $e^{-t} = \delta/2$ we arrive at

$$\mathbb{P}\left[\|X\|_{\ell_2} \geq \sqrt{k}\left(1 + \sqrt{\frac{\log(2/\delta)}{k}} + \frac{\log(2/\delta)}{k}\right)\right] \leq \frac{\delta}{2} \quad (\text{F11})$$

and

$$\mathbb{P}\left[\|\Delta\|_{\ell_2} \geq \sigma^2\sqrt{k}\left(1 + \frac{g(\delta)}{\sqrt{k}} + \frac{g^2(\delta)}{k}\right)\right] \leq \frac{\delta}{2}. \quad (\text{F12})$$

Our aim is to bound the inner product $\langle \gamma, \Delta \rangle$, which we will write as $\|\gamma\|_{\ell_2} \|\Delta\|_{\ell_2} \langle \tilde{\gamma}, \tilde{\Delta} \rangle$, with $\tilde{\gamma}, \tilde{\Delta}$ normalized. This allows us to use Lemma 19 for $e^{-t^2/2} = \delta/2$, and we arrive at

$$\mathbb{P}\left[|\langle \gamma, \Delta \rangle| \geq \|\gamma\|_{\ell_2} \|\Delta\|_{\ell_2} \frac{g(\delta)}{\sqrt{k-1}}\right] \leq \delta/2. \quad (\text{F13})$$

Combining this bound with the bound (F12) for $\|\Delta\|_{\ell_2}$ via the union bound completes the proof. $\square$

-
- [1] H.-Y. Huang, R. Kueng, and J. Preskill, *Predicting many properties of a quantum system from very few measurements*, *Nat. Phys.* **16**, 1050–1057 (2020), [arXiv:2002.08953 \[quant-ph\]](#).
 - [2] A. Elben, R. Kueng, H.-Y. R. Huang, R. van Bijnen, C. Kokail, M. Dalmonte, P. Calabrese, B. Kraus, J. Preskill, P. Zoller, and B. Vermersch, *Mixed-state entanglement from local randomized measurements*, *Phys. Rev. Lett.* **125**, 200501 (2020), [arXiv:2007.06305 \[quant-ph\]](#).
 - [3] T. Zhang, J. Sun, X.-X. Fang, X.-M. Zhang, X. Yuan, and H. Lu, *Experimental quantum state measurement with classical shadows*, *Phys. Rev. Lett.* **127**, 200501 (2021), [arXiv:2008.05234 \[quant-ph\]](#).
 - [4] G. Struchalin, Y. A. Zagorovskii, E. Kovlakov, S. Straupe, and S. Kulik, *Experimental estimation of quantum state properties from classical shadows*, *PRX Quantum* **2**, 010307 (2021), [arXiv:2008.05234 \[quant-ph\]](#).
 - [5] H.-Y. Huang, R. Kueng, G. Torlai, V. V. Albert, and J. Preskill, *Provably efficient machine learning for quantum many-body problems*, *Science* **377**, eabk3333 (2022), [arXiv:2106.12627 \[quant-ph\]](#).
 - [6] A. Elben, S. T. Flammia, H.-Y. Huang, R. Kueng, J. Preskill, B. Vermersch, and P. Zoller, *The randomized measurement toolbox*, *Nat. Rev. Phys.* **10.1038/s42254-022-00535-2** (2022), [arXiv:2203.11374](#).
 - [7] W. J. Huggins, B. A. O’Gorman, N. C. Rubin, D. R. Reichman, R. Babbush, and J. Lee, *Unbiasing fermionic quantum Monte Carlo with a quantum computer*, *Nature* **603**, 416–420 (2022), [2106.16235 \[quant-ph\]](#).
 - [8] A. A. Akhtar, H.-Y. Hu, and Y.-Z. You, *Scalable and flexible classical shadow tomography with tensor networks*, *Quantum* **7**, 1026 (2023), [arXiv:2209.02093 \[quant-ph\]](#).
 - [9] C. Bertoni, J. Haferkamp, M. Hinsche, M. Ioannou, J. Eisert, and H. Pashayan, *Shallow shadows: Expectation estimation using low-depth random Clifford circuits*, [arXiv:2209.12924 \[quant-ph\]](#).
 - [10] M. Arienzo, M. Heinrich, I. Roth, and M. Kliesch, *Closed-form analytic expressions for shadow estimation with brickwork circuits*, *Quantum Inf. Comp.* **23**, 961 (2023), [arXiv:2211.09835 \[quant-ph\]](#).
 - [11] K. Wan, W. J. Huggins, J. Lee, and R. Babbush, *Matchgate shadows for fermionic quantum simulation*, [arXiv:2207.13723 \[quant-ph\]](#) (2022).
 - [12] S. Chen, W. Yu, P. Zeng, and S. T. Flammia, *Robust shadow estimation*, *PRX Quantum* **2**, 030348 (2021), [arXiv:2011.09636 \[quant-ph\]](#).
 - [13] D. E. Koh and S. Grewal, *Classical shadows with noise*, *Quantum* **6**, 776 (2022).
 - [14] H. Jnane, J. Steinberg, Z. Cai, H. Chau Nguyen, and B. Koczor, *Quantum Error Mitigated Classical Shadows*, (2023), [arXiv:2305.04956 \[quant-ph\]](#).
 - [15] V. Vitale, A. Rath, P. Jurcevic, A. Elben, C. Branciard, and B. Vermersch, *Estimation of the quantum Fisher information on a quantum processor*, [arXiv:2307.16882 \[quant-ph\]](#) (2023).
 - [16] A. Zhao and A. Miyake, *Group-theoretic error mitigation enabled by classical shadows and symmetries*, (2023), [arXiv:2310.03071 \[quant-ph\]](#).
 - [17] B. Wu and D. Enshan Koh, *Error-mitigated fermionic classical shadows on noisy quantum devices*, (2023), [arXiv:2310.12726 \[quant-ph\]](#).
 - [18] R. Brieger, I. Roth, and M. Kliesch, *Compressive gate set tomography*, *PRX Quantum* **4**, 010325 (2023), [arXiv:2112.05176 \[quant-ph\]](#).
 - [19] E. T. Campbell, *Catalysis and activation of magic states in fault-tolerant architectures*, *Physical Review A* **83**, 10.1103/PhysRevA.83.032317, [arXiv:1010.0104 \[quant-ph\]](#).

- [20] L. Leone, S. F. Oliviero, and A. Hamma, *Stabilizer Rényi entropy*, *Physical Review Letters* **128**, 050402 (2022), [arXiv:2106.12587 \[quant-ph\]](#).
- [21] J. R. Seddon, B. Regula, H. Pashayan, Y. Ouyang, and E. T. Campbell, *Quantifying Quantum Speedups: Improved Classical Simulation From Tighter Magic Monotones*, *PRX Quantum* **2**, 010345 (2021).
- [22] M. Howard and E. Campbell, *Application of a Resource Theory for Magic States to Fault-Tolerant Quantum Computing*, *Phys. Rev. Lett.* **118**, 090501 (2017), [arXiv:1609.07488 \[quant-ph\]](#).
- [23] P. Rall, D. Liang, J. Cook, and W. Kretschmer, *Simulation of qubit quantum circuits via Pauli propagation*, *Phys. Rev. A* **99**, 062337 (2019), [arXiv:1901.09070 \[quant-ph\]](#).
- [24] J. J. Wallman and J. Emerson, *Noise tailoring for scalable quantum computation via randomized compiling*, *Phys. Rev. A* **94**, 052325 (2016), [arXiv:1512.01098 \[quant-ph\]](#).
- [25] A. Hashim, R. K. Naik, A. Morvan, J.-L. Ville, B. Mitchell, J. M. Kreikebaum, M. Davis, E. Smith, C. Iancu, K. P. O'Brien, I. Hincks, J. J. Wallman, J. Emerson, and I. Siddiqi, *Randomized compiling for scalable quantum computing on a noisy superconducting quantum processor*, *Physical Review X* **11**, 041039 (2021), [arXiv:2010.00215 \[quant-ph\]](#).
- [26] M. Ware, G. Ribeill, D. Ristè, C. A. Ryan, B. Johnson, and M. P. da Silva, *Experimental Pauli-frame randomization on a superconducting qubit*, *Phys. Rev. A* **103**, 042604 (2021), [arXiv:1803.01818 \[quant-ph\]](#).
- [27] Note that for (left) Pauli-invariant ensembles, such as local or global Clifford unitaries, gate-independent left noise can be effectively Pauli-twirled and hence the Pauli noise assumption is implicit.
- [28] J. Watrous, *The Theory of Quantum Information* (Cambridge University Press, 2018).
- [29] E. Magesan, J. M. Gambetta, and J. Emerson, *Characterizing quantum gates via randomized benchmarking*, *Phys. Rev. A* **85**, 042311 (2012), [arXiv:1109.6887](#).
- [30] M. Heinrich, M. Kliesch, and I. Roth, *General guarantees for randomized benchmarking with random quantum circuits*, [arXiv:2212.06181 \[quant-ph\]](#) (2022).
- [31] S. T. Flammia and J. J. Wallman, *Efficient estimation of Pauli channels*, *ACM Transactions on Quantum Computing* **1**, 1 (2020), [arXiv:1907.12976 \[quant-ph\]](#).
- [32] U. Vazirani and S. Rao, *Combinatorial algorithms and data structures*, lecture notes (2011).
- [33] S. Dasgupta and A. Gupta, *An elementary proof of a theorem of Johnson and Lindenstrauss*, *Random Structures and Algorithms* **22**, 60 (2003).
- [34] M. A. Nielsen, *A simple formula for the average gate fidelity of a quantum dynamical operation*, *Phys. Lett. A* **303**, 249 (2002), [quant-ph/0205035](#).

C Full reports of compressive GST on a trapped ion system

This section contains GST reports based on results of mGST applied to trapped ion qubits, see Section 3.2. The reports shown here represent only a subset of reports generated throughout the collaboration with quantum optics group at the university of Siegen and supplement the discussion in Section 3.2. Throughout the reports, the number of free parameters up to gauge transformations is listed and should be compared with the number of sequences used. The coloring indicates whether the sequences are expected to be informationally complete, simply based on the parameter counting argument.

Single qubit GST report

(Dated: February 21, 2024)

I. SETUP

- Date of the experiment: 11.01.2014 & 12.01.2024 (DUC).
- Number of sequences: 200.
- Average shots per sequence: 413.
- **Rank: 1 (unitary).**
- Number of free parameters: 36.
- Gate set:

$$\mathcal{X} = (\mathbb{1}_1, \mathbb{1}_2, \sigma_x, \sigma_y, e^{i\frac{\pi}{4}\sigma_x}, e^{i\frac{\pi}{4}\sigma_y}) \quad (1)$$

- Unitary model:

$$U = \exp\left(\frac{i\alpha}{2} \sum_{a \in \{X,Y,Z\}} n_a \sigma_a\right). \quad (2)$$

II. ERROR MEASURES

Table I. Rotation angle and axes tilt with errors corresponding to the 95th percentile over 50 bootstrapping runs.

	Rotation angle $/\pi$	Rotation angle $/\pi$	Axes tilt vs. target (in $^\circ$)	Axes estimation error (in $^\circ$)
Idle-short	0.0105	[0.0067,1.9903]	–	27.2264
Idle-long	0.0110	[0.0095,1.9896]	–	7.3257
Rx(pi):0	0.9994	[0.9949,0.9999]	0.1603	0.0498
Ry(pi):0	0.9990	[0.9965,0.9999]	0.2884	0.0961
Rx(pi/2):1	0.5051	[0.4989,0.5089]	0.2031	0.5884
Ry(pi/2):1	0.5005	[0.4966,0.5024]	0.4722	0.6220

Table II. Normalized rotation axes coefficient. Errors correspond to the 95th percentile over 50 bootstrapping runs.

	Idle-short	Idle-long	Rx(pi):0	Ry(pi):0	Rx(pi/2):1	Ry(pi/2):1
α/π	0.011 [0.007,1.990]	0.011 [0.009,1.990]	0.999 [0.995,1.000]	0.999 [0.996,1.000]	0.505 [0.499,0.509]	0.500 [0.497,0.502]
n_X	-0.434 [-0.552,0.667]	0.584 [-0.587,0.784]	1.000 [-1.000,1.000]	-0.000 [-0.005,0.007]	-1.000 [-1.000,-1.000]	-0.008 [-0.018,-0.000]
n_Y	-0.369 [-0.575,0.699]	-0.517 [-0.629,0.816]	0.002 [-0.008,0.006]	1.000 [-1.000,1.000]	-0.003 [-0.012,0.003]	-1.000 [-1.000,-1.000]
n_Z	0.822 [-0.604,0.976]	0.626 [-0.579,0.824]	-0.001 [-0.004,0.006]	-0.005 [-0.009,0.008]	0.001 [-0.007,0.007]	-0.001 [-0.010,0.008]

Table III. Gate quality measures of the unitary approximation with errors corresponding to the 95th percentile over 50 bootstrapping runs. Disclaimer: This are not the true average gate fidelities/diamond distance of the gates, just of their unitary approximation. The real quality measures can only be deduced from the full rank reconstruction.

	Average gate Fidelity	Diamond distances
Idle-short	0.9998 [0.9996,0.9999]	0.0165 [0.0092,0.0236]
Idle-long	0.9998 [0.9997,0.9999]	0.0173 [0.0149,0.0228]
Rx(pi):0	1.0000 [0.9999,1.0000]	0.0029 [0.0016,0.0116]
Ry(pi):0	1.0000 [0.9999,1.0000]	0.0053 [0.0023,0.0107]
Rx(pi/2):1	1.0000 [0.9999,1.0000]	0.0084 [0.0031,0.0146]
Ry(pi/2):1	1.0000 [0.9999,1.0000]	0.0059 [0.0031,0.0131]

Table IV. State and measurement quality measures with errors corresponding to the 95th percentile over 50 bootstrapping runs.

Final cost	Mean TVD: estimate - data	Mean TVD: target - data	POVM - diamond dist.	State - trace dist.
0.0011 [0.0006,0.0015]	0.0292 [0.0194,0.0329]	0.0332 [0.0338,0.0383]	0.0107 [0.0030,0.0228]	0.0159 [0.0127,0.0391]

III. GATE AND SPAM PLOTS

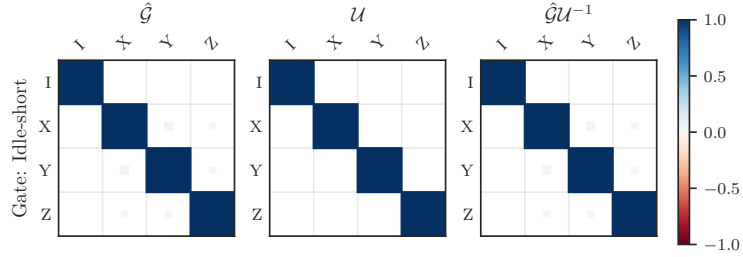


Figure 1. Left: Process matrix of gate 0 in Pauli basis; Right: Target.

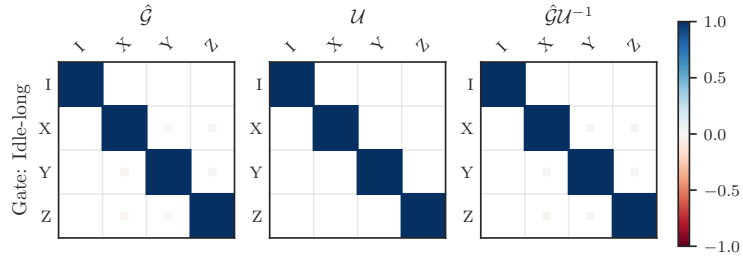


Figure 2. Left: Process matrix of gate 1 in Pauli basis; Right: Target.

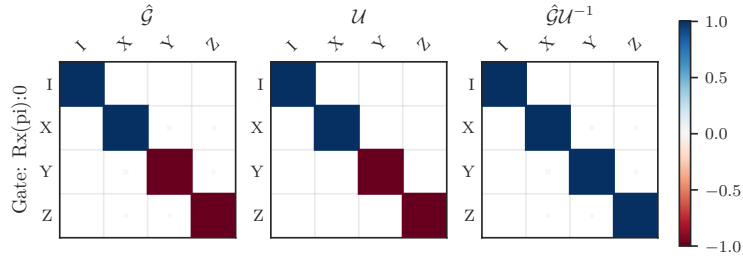


Figure 3. Left: Process matrix of gate 2 in Pauli basis; Right: Target.

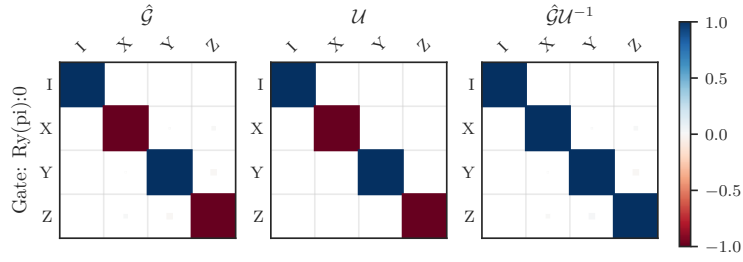


Figure 4. Left: Process matrix of gate 3 in Pauli basis; Right: Target.

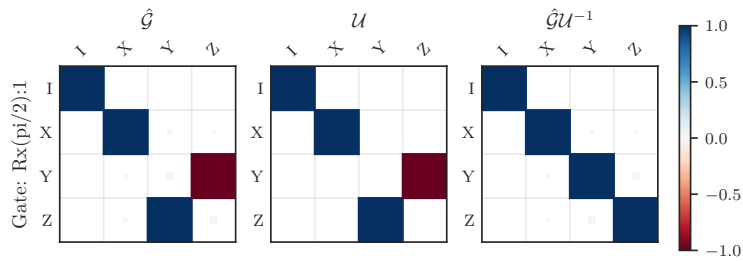


Figure 5. Left: Process matrix of gate 4 in Pauli basis; Right: Target.

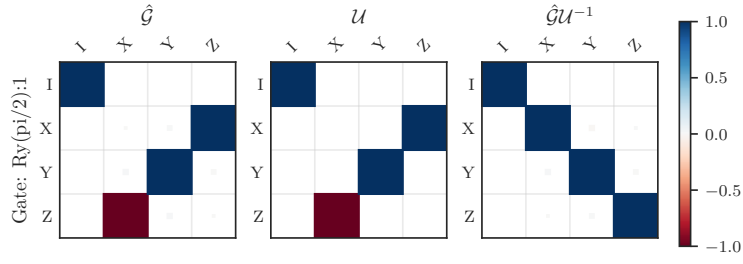
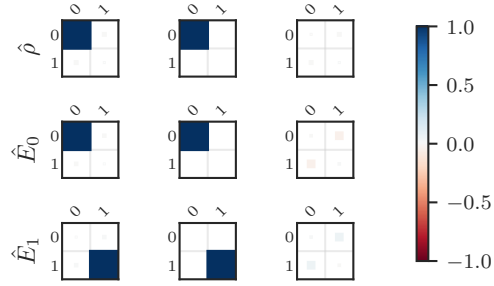
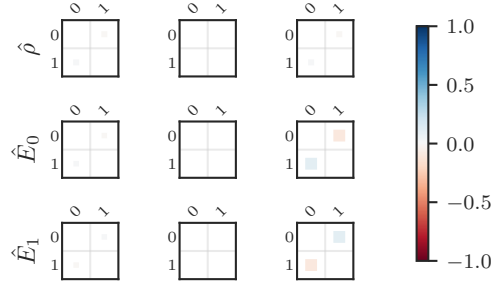


Figure 6. Left: Process matrix of gate 5 in Pauli basis; Right: Target.

Figure 7. Left column: real part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.Figure 8. Left column: imaginary part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.

Single qubit GST report

(Dated: February 21, 2024)

I. SETUP

- Date of the experiment: 11.01.2014 & 12.01.2024 (DUC).
- Number of sequences: 200.
- Average shots per sequence: 413.
- **Rank: 4.**
- Number of free parameters: 180.
- Gate set:

$$\mathfrak{G} = (\mathbb{1}_1, \mathbb{1}_2, \sigma_x, \sigma_y, e^{i\frac{\pi}{4}\sigma_x}, e^{i\frac{\pi}{4}\sigma_y}) \quad (1)$$

- Left noise model: Given the reconstructed channels $\hat{\mathcal{G}}_1, \dots, \hat{\mathcal{G}}_k$ corresponding to unitary target gates $\mathcal{U}_1, \dots, \mathcal{U}_k$, the left noise channel Λ_i for gate i is given by

$$\Lambda_i = \hat{\mathcal{G}}_i \mathcal{U}_i^{-1} \quad (2)$$

such that $\hat{\mathcal{G}}_i = \Lambda_i \mathcal{U}_i$.

- Local dephasing channel: $\Lambda_{\text{dephase}}(\rho) = (1 - p)\rho + p\sigma_z\rho\sigma_z$

II. ERROR MEASURES

Table I. Gate quality measures with errors corresponding to the 95th percentile over 50 bootstrapping runs.

	Average gate fidelity $\mathcal{F}_{\text{avg}}(\mathcal{U}_i, \hat{\mathcal{G}}_i)$	Diamond distance $\frac{1}{2} \ \mathcal{U}_i - \hat{\mathcal{G}}_i\ _{\diamond}$	Unitarity $u(\hat{\mathcal{G}}_i)$
Idle-short	0.9977 [0.9961,0.9990]	0.0207 [0.0150,0.0281]	0.9918 [0.9859,0.9972]
Idle-long	0.9972 [0.9948,0.9979]	0.0221 [0.0185,0.0287]	0.9897 [0.9810,0.9931]
Rx(pi):0	0.9989 [0.9969,0.9999]	0.0048 [0.0029,0.0123]	0.9958 [0.9880,0.9999]
Ry(pi):0	0.9991 [0.9977,1.0000]	0.0060 [0.0039,0.0129]	0.9965 [0.9911,1.0000]
Rx(pi/2):1	0.9994 [0.9980,1.0000]	0.0069 [0.0040,0.0155]	0.9977 [0.9919,1.0000]
Ry(pi/2):1	0.9986 [0.9960,0.9999]	0.0079 [0.0053,0.0158]	0.9944 [0.9843,0.9997]

Table II. State and measurement quality measures with errors corresponding to the 95th percentile over 50 bootstrapping runs.

Final cost	Mean TVD: estimate - data	Mean TVD: target - data	POVM - diamond dist.	State - trace dist.
0.0004 [0.0006,0.0009]	0.0152 [0.0186,0.0223]	0.0332 [0.0346,0.0390]	0.0195 [0.0155,0.0257]	0.0146 [0.0089,0.0192]

III. GATE AND SPAM PLOTS

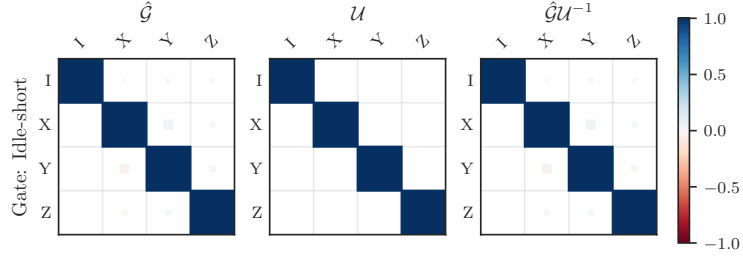


Figure 1. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

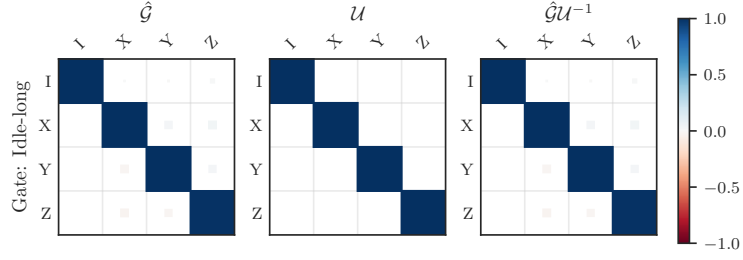


Figure 2. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

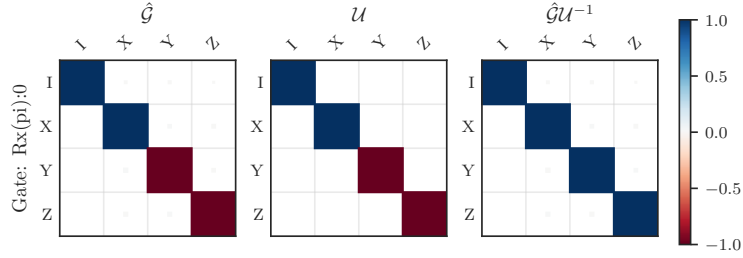


Figure 3. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

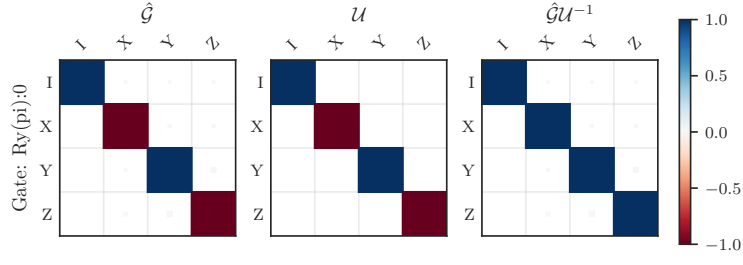


Figure 4. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

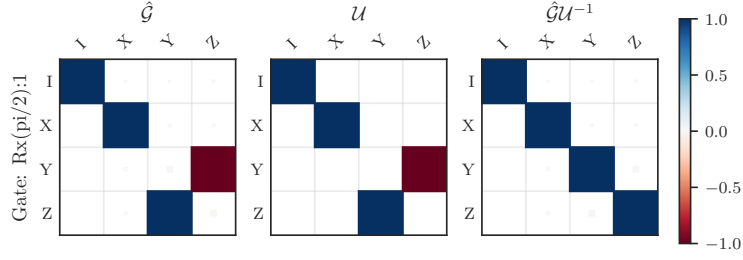


Figure 5. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

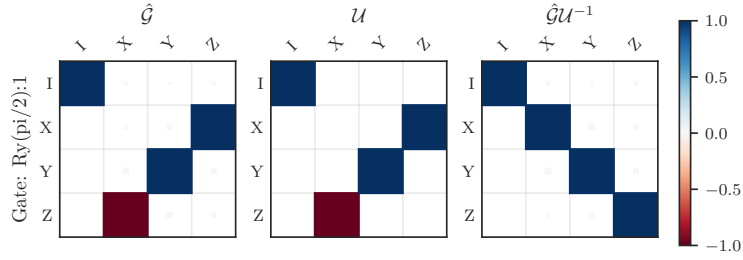


Figure 6. Process matrix in the Pauli basis with entries in $[-1, 1]$. Left side: GST reconstruction, center: ideal gate, right side: error channel (ideally the identity).

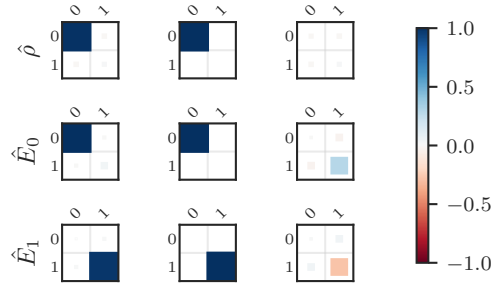


Figure 7. Left column: real part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.

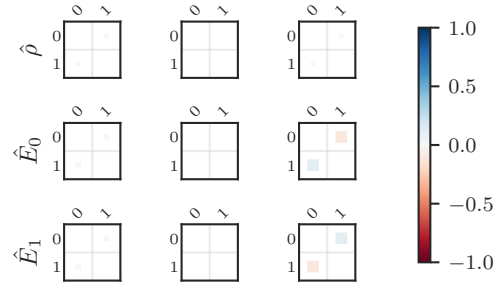


Figure 8. Left column: imaginary part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.

Two qubit GST report

(Dated: February 23, 2024)

I. SETUP

- Date of the experiment: 05.02.2024.
- Number of sequences: 360.
- Average shots per sequence: 116.
- **Rank: 16.**
- Number of free parameters: 1488.
- Target gate set:

$$\mathcal{X} = (\mathbb{1}, e^{-i\frac{\pi}{4}\sigma_x} \otimes \mathbb{1}, e^{-i\frac{\pi}{4}\sigma_y} \otimes \mathbb{1}, \mathbb{1} \otimes e^{-i\frac{\pi}{4}\sigma_x}, \mathbb{1} \otimes e^{-i\frac{\pi}{4}\sigma_y}, e^{-i\frac{\pi}{4}\sigma_z \otimes \sigma_z}). \quad (1)$$

- Left noise model: Given the reconstructed channels $\hat{\mathcal{G}}_1, \dots, \hat{\mathcal{G}}_k$ corresponding to unitary target gates $\mathcal{U}_1, \dots, \mathcal{U}_k$, the left noise channel Λ_i for gate i is given by

$$\Lambda_i = \hat{\mathcal{G}}_i \mathcal{U}_i^{-1} \quad (2)$$

such that $\hat{\mathcal{G}}_i = \Lambda_i \mathcal{U}_i$.

- Local dephasing channel: $\Lambda_{\text{dephase}}(\rho) = (1 - p)\rho + p\sigma_z\rho\sigma_z$

II. ERROR MEASURES

Table I. Gate quality measures

	Average gate fidelity $\mathcal{F}_{\text{avg}}(\mathcal{U}_i, \hat{\mathcal{G}}_i)$	Diamond distance $\frac{1}{2}\ \mathcal{U}_i - \hat{\mathcal{G}}_i\ _{\diamond}$	Unitarity $u(\hat{\mathcal{G}}_i)$	Dephasing probability of Λ_i for Qubit 1	Dephasing probability of Λ_i for Qubit 2
Idle	0.9649	0.1270	0.9140	0.0075	0.0301
Rx(pi/2):0	0.9927	0.0308	0.9815	0.0024	0.0050
Ry(pi/2):0	0.9786	0.0766	0.9466	0.0105	0.0114
Rx(pi/2):1	0.9967	0.0442	0.9934	0.0018	0.0012
Ry(pi/2):1	0.9746	0.0788	0.9371	0.0143	0.0115
Rzz(pi/2)	0.6852	0.5056	0.4164	0.1481	0.2597

Table II. Eigenvalues of the Choi state: The number of nonzero eigenvalues gives the Kraus rank.

	0	1	2	3	4	5	6
Idle	0.95319	0.04671	0.00010	0.00000	0.00000	0.00000	0.00000
Rx(pi/2):0	0.99980	0.00020	0.00000	0.00000	0.00000	0.00000	0.00000
Ry(pi/2):0	0.98114	0.01871	0.00015	0.00000	0.00000	0.00000	0.00000
Rx(pi/2):1	0.99646	0.00354	0.00000	0.00000	0.00000	0.00000	0.00000
Ry(pi/2):0	0.96671	0.03304	0.00026	0.00000	0.00000	0.00000	0.00000
$e^{-i\frac{\pi}{4}Z \otimes Z}$	0.67117	0.22288	0.10595	0.00000	0.00000	0.00000	0.00000

Table III. State and measurement quality measures

Final cost	Mean TVD: estimate - data	Mean TVD: target - data	POVM - diamond dist.	State - trace dist.
0.0017	0.0609	0.1050	0.0293	0.0303

III. GATE AND SPAM PLOTS

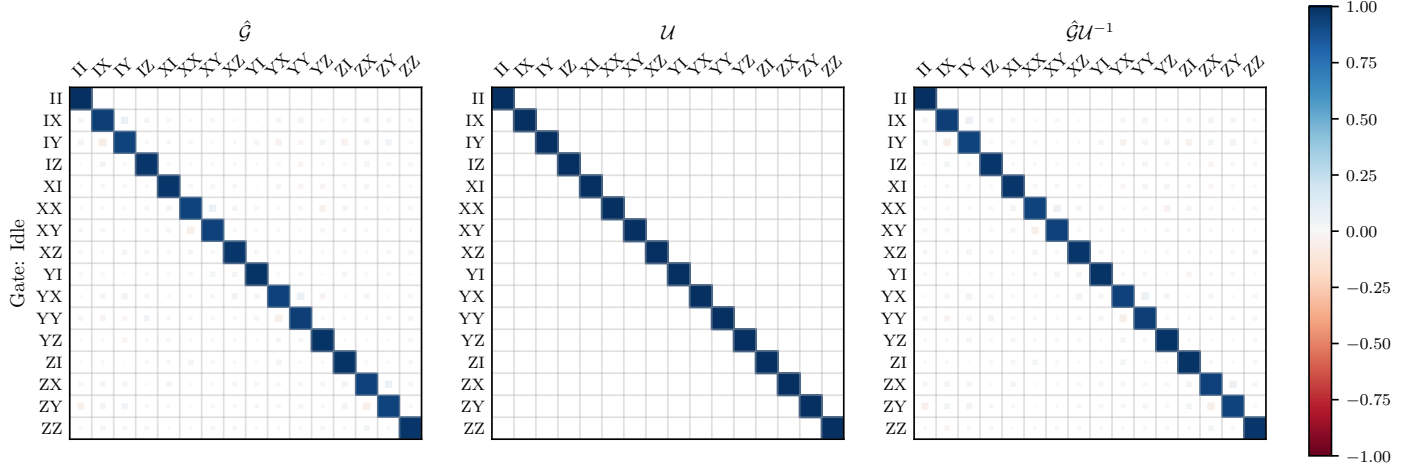


Figure 1. Left: Process matrix of gate 0 in Pauli basis; Right: Target.

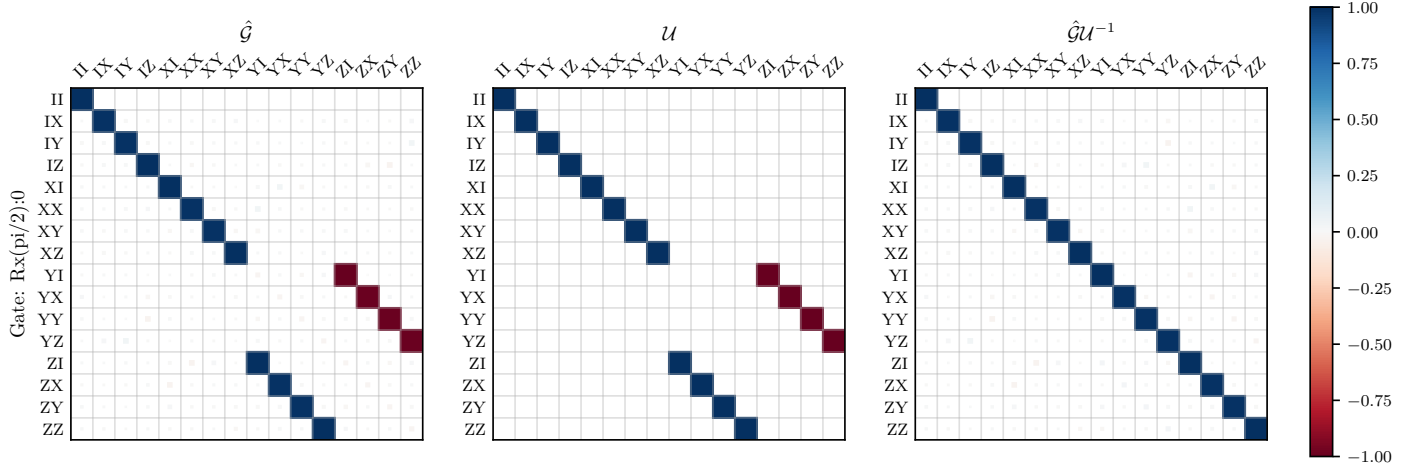


Figure 2. Left: Process matrix of gate 1 in Pauli basis; Right: Target.

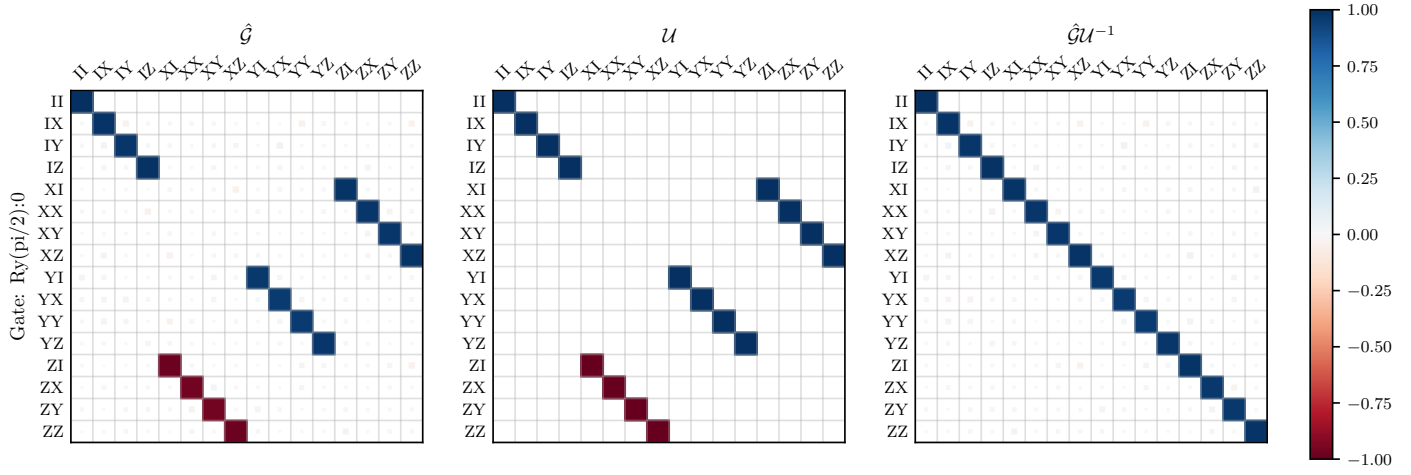


Figure 3. Left: Process matrix of gate 2 in Pauli basis; Right: Target.

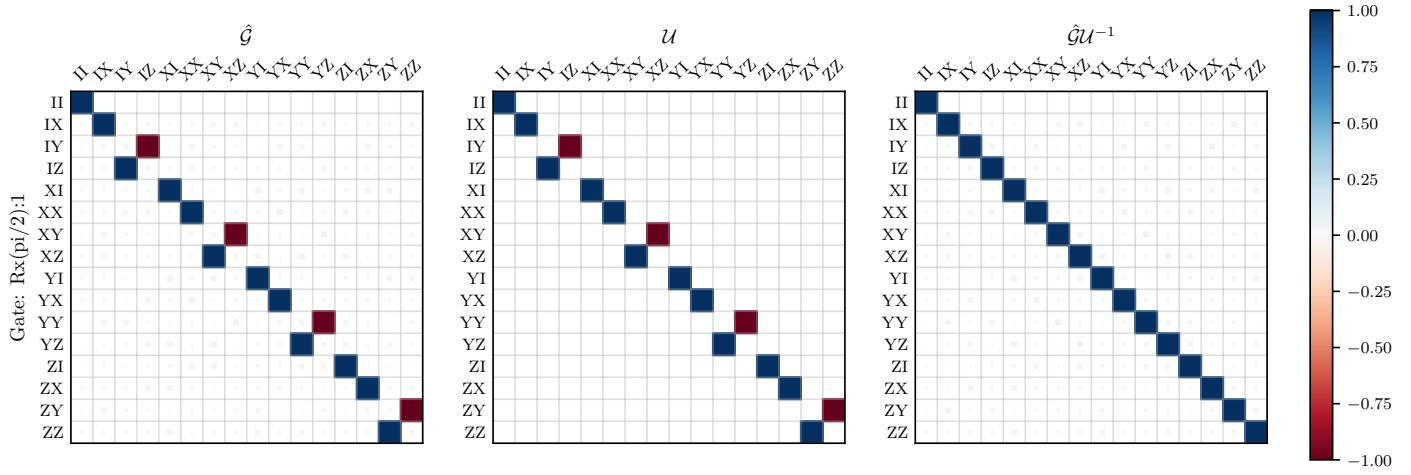


Figure 4. Left: Process matrix of gate 3 in Pauli basis; Right: Target.

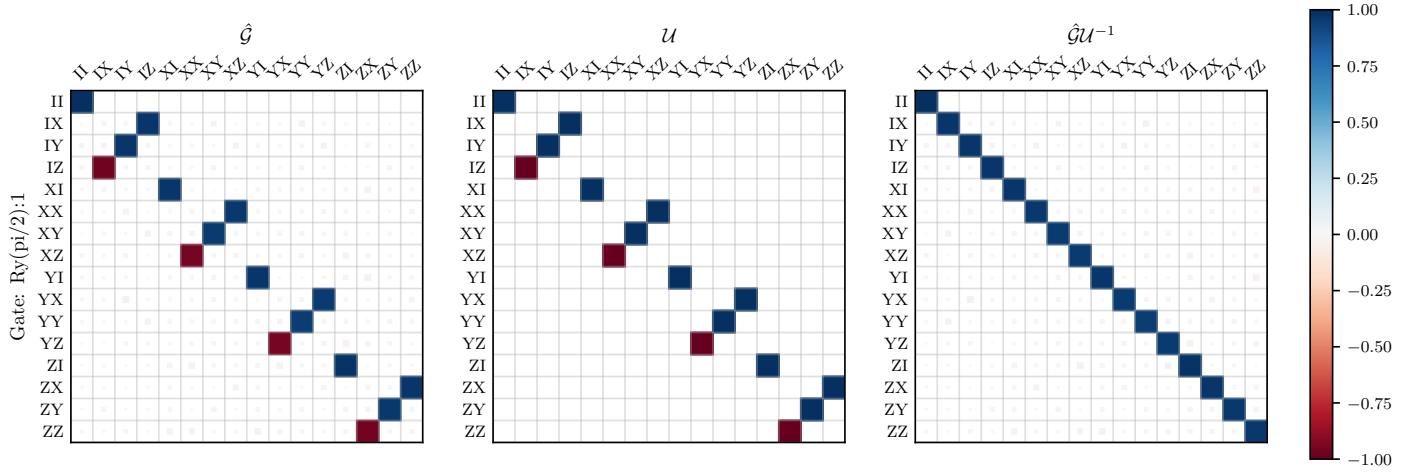


Figure 5. Left: Process matrix of gate 4 in Pauli basis; Right: Target.

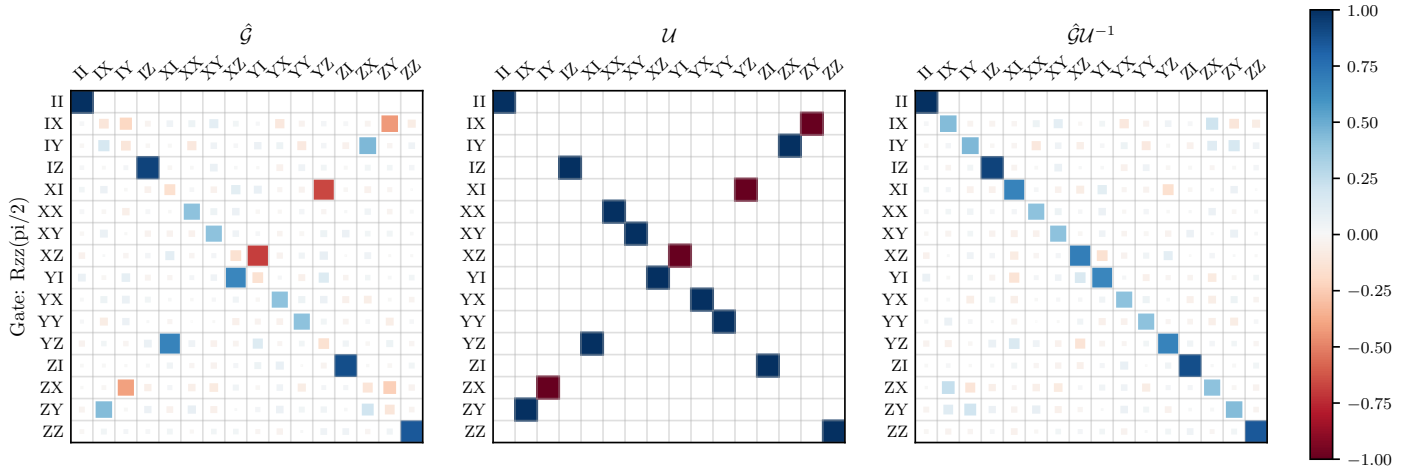


Figure 6. Left: Process matrix of gate 5 in Pauli basis; Right: Target.

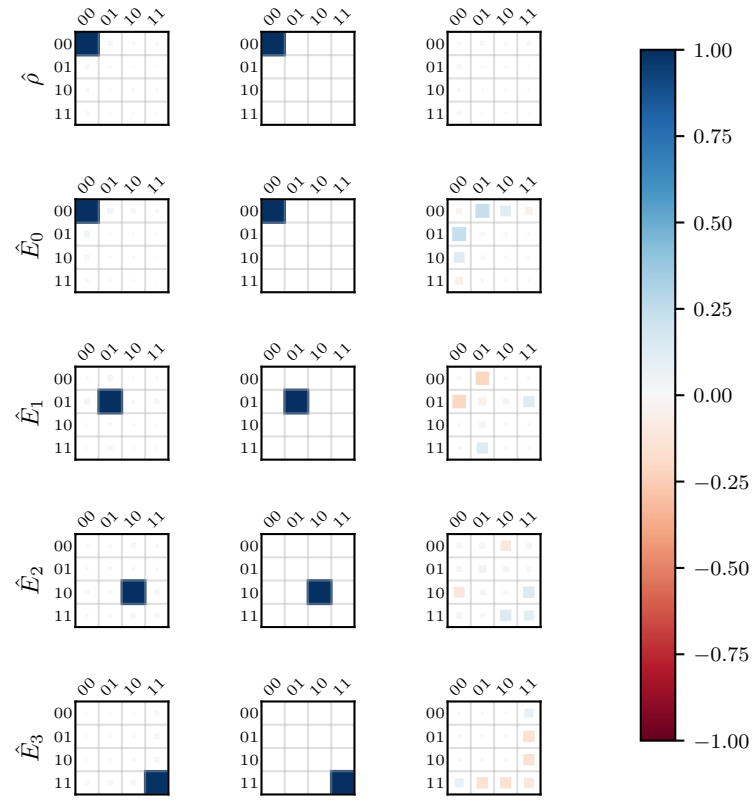


Figure 7. Left column: real part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.

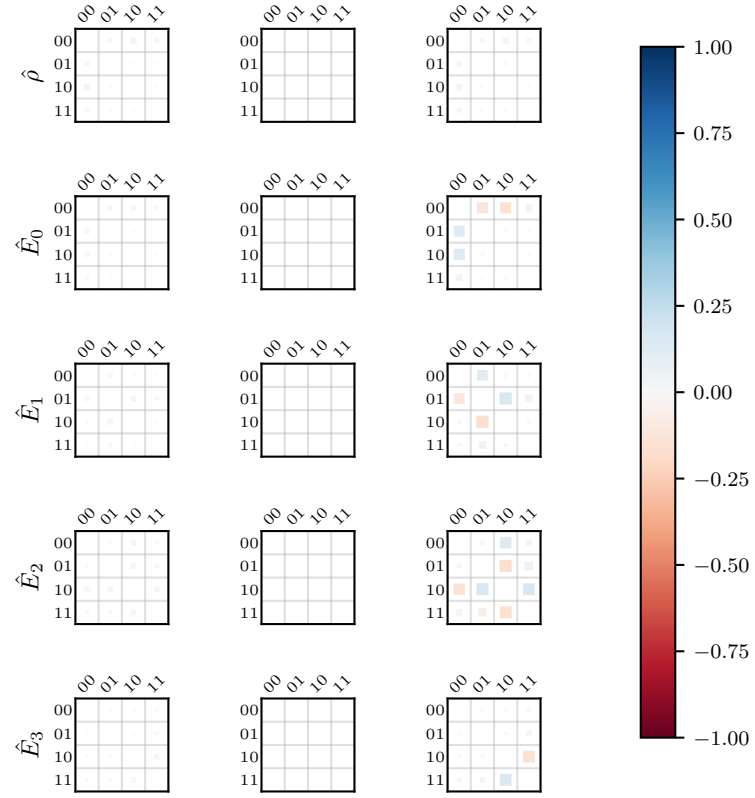


Figure 8. Left column: imaginary part of state and measurement in standard basis, right column: magnified errors to ideal implementation $10 \cdot (\hat{\rho} - \rho_{\text{ideal}})$ and $10 \cdot (\hat{E}_i - E_{i,\text{ideal}})$.

D Tutorial notebook for the mGST python package

Basic examples

March 13, 2024

1 Single qubit/qutrit examples and algorithm usage

In this notebook we use simulated measurements to show how the mGST algorithm is used. We make sure all the necessary functions are available by running the following python scripts.

```
[1]: import numpy as np
import random
import matplotlib.pyplot as plt
import numba
import time
from scipy.linalg import expm
from mGST import additional_fns, low_level_jit, algorithm, compatibility
```

The latest package versions tested are: numpy==1.21.6 pygsti==0.9.10 numba==0.55.1

You can check your versions via the following command:

```
[2]: print('\n'.join(f'{m.__name__}=={m.__version__}' for m in globals(
).values() if getattr(m, '__version__', None)))
```

numpy==1.21.6

numba==0.55.1

1.0.1 First, let's create a random unitary gate set and define the parameters that we need:

```
[3]: pdim = 2    # physical dimension
r = pdim**2    # matrix dimension of the gate superoperators
l = 8    # maximum number of gates in each measurement sequence (sequence length)
d = 4    # number of gates in the gate set
rK_true = 1    # rank of simulated gates used for testing
rK = 1    # rank of the mGST model estimate
n_povm = 2    # number of POVM-elements
```

We can use some of the functions we imported to generate a random gate set. The function `randKrausSet_Haar(d,r,rK_true)` generates `d` sets of Kraus operators, where each set is found by taking a Haar random unitary and using a subset of its columns the generate an isometry of shape $(rK_true \times pdim) \times (pdim)$. For an initial state and a POVM element we can simply take a computational basis elements.

```
[4]: K_true = additional_fns.randKrausSet_Haar(
      d, r, rK_true)    # tensor of random Kraus operators
X_true = np.einsum('ijkl,ijnm -> iknlm', K_true, K_true.conj())
      ).reshape(d, r, r)    # tensor of superoperators

K_depolar = additional_fns.depolar(pd, 0.02) # Kraus-rep of depolarizing channel
G_depolar = np.einsum('jkl,jnm -> knlm', K_depolar, K_depolar.conj()).reshape(r, r)
# |10> initial state with depolarizing noise
rho_true = G_depolar@np.array([[1, 0], [0, 0]]).reshape(-1).astype(np.complex128)

# Computational basis measurement:
E1 = np.array([[1, 0], [0, 0]]).reshape(-1)
E2 = np.array([[0, 0], [0, 1]]).reshape(-1)
E_true = np.array([E1, E2]).astype(np.complex128)    # Full POVM
```

Next up we need some gate sequence instructions and simulated measurements. Each gate is identified by an index between 0 and $d-1$, meaning a gate sequences can most simply be represented by a list of gate indices. We write a full set containing N many sequence instructions as a numpy array J of shape $N \times d$. The resulting state after each sequences is measured $meas_samples$ times with a POVM consisting of n_povm many elements. Therefore we will have n_povm -many estimated probabilities per sequences and we can collect all results in a numpy array y of shape $n_povm \times N$.

```
[5]: N = 100 # Number of sequences
meas_samples = 1e5 # Number of samples per sequences
# generate random numbers between 0 and $d-1$
J_rand = np.array(random.sample(range(d), N))
# turn random numbers into gate instructions
J = np.array([low_level_jit.local_basis(ind, d, 1) for ind in J_rand])
y = np.real(np.array([E_true[i].conj()@low_level_jit.contract(X_true,
    ↪j)@rho_true for j in J]
    for i in range(n_povm)])) # obtain ideal output probabilities
# simulate finite sampling statistics
y_sampled = additional_fns.sampled_measurements(y, meas_samples).copy()
```

For our first test we use an initialization where the state preparation and measurement are random, but the gate are just rotated versions of the ideal gates.

```
[6]: delta = .1 # unitary noise parameter

# Generate noisy version of true gate set
K0 = np.zeros((d, rK, pdim, pdim)).astype(np.complex128)
for i in range(d):
    U_p = expm(delta*1j*additional_fns.randHerm(pdim)
        ).astype(np.complex128) # unitary noise
    K0[i] = np.einsum('jkl,lm', K_true[i], U_p)
X0 = np.einsum('ijkl,ijnm -> iknlm', K0, K0.conj()).reshape(d, r, r)
```

```

rho0 = additional_fns.randpsd(r).copy()    # random initial state
A0 = additional_fns.randKrausSet(1, r, n_povm)[0].conj()    # random POVM
    ↪ decomposition
E0 = np.array([(A0[i].T.conj()@A0[i]).reshape(-1)
               for i in range(n_povm)]).copy()

```

Now it's time to run `mGST` on the data set. Note that if the algorithm is run for the first time on a new machine it can take up to a few minutes to compile the low level functions (such as derivatives). The main function is called `run_mGST`, for information it's variables and outputs we can call the `help(run_mGST)`:

```

[7]: bsize = 50    # The batch size on which the optimization is started
K, X, E, rho, res_list = algorithm.run_mGST(y_sampled, J, l, d, r, rK, n_povm,
    ↪ bsize, meas_samples, method='SFN',
                                max_inits=10, max_iter=30, final_iter=10,
                                target_rel_prec=1e-4, init=[K0, E0, rho0])
plt.semilogy(res_list)    # plot the objective function over the iterations
plt.tight_layout(rect = [.2,.2,.8,.8])
plt.show()
print('Mean variation error:', additional_fns.MVE(X_true, E_true, rho_true, X,
    ↪ E,
        rho, d, l, n_povm)[0])    # output the final mean variation error

```

Starting optimization...

67%| | 20/30 [00:04<00:02, 4.49it/s]

Optimization successful, improving estimate over full data...

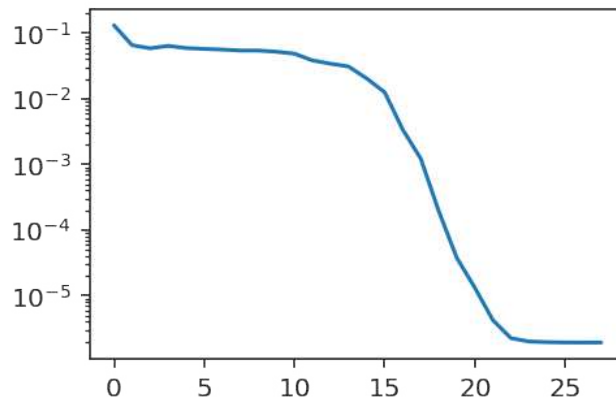
50%| | 5/10 [00:01<00:01, 3.48it/s]

#####

Convergence criterion satisfied

Final objective function value 1.9344087733601114e-06 with # of
initializations: 1

Total runtime: 5.909569263458252



Mean variation error: 0.0005429049138832866

1.0.2 A note on hyperparameters

bsize: This controls the “batch size”, meaning how many of the total sequences are sampled per initialization step. This is done until a good estimate is reached, which is afterwards improved on the full set of sequences. A too small value of bsize leads to chaotic jumping in the parameter space without convergence, while a too large value leads to longer runtimes. Generally, the more free parameters a model has, the higher the batch size needs to be. Good heuristic values are: bsize = 50 for a single quabit gate set of 3 gates, bsize = 80 for a single qutrit gate set of 6 gates. bsize = 120 for a two qubit gate set of 6 gates.

max_inits: Controls the maximum number of reinitializations, this can be increased if the algorithm doesn’t converge. The target Kraus rank rK can have a large influence on the required number of initializations and we find that in many cases rK=2 requires fewer initializations than rK=1 (See also discussion in <https://arxiv.org/abs/2112.05176>).

max_iter: Maximum number of iterations spend on the batch optimization per initialization. Generally values around 150-200 are a good trade off between leaving enough iterations to converge if an optimal value can be reached from the given initialization, and not spending too much time on a bad initialization.

final_iter and **target_rel_prec:** If the convergence criterion is satisfied, a maximum of final_iter - many iterations on the full dataset are performed to fully converge. If the iteration on iteration improvement on the objective function is less than target_rel_prec * delta, where delta is the convergence threshold, then the final iteration loop is terminated and a the resulting gate set estimate is returned. If computation time is not an issue and higher precision of the estimator is desired, then target_rel_prec can be decreased and final_iter increased.

1.1 XYI gate set $\{\text{Id}, e^{i\frac{\alpha}{2}\sigma_y}, e^{i\frac{\alpha}{2}\sigma_x}\}$ from random initialization

The XYI gate set is a minimal gate set that is tomographically complete when applied to the $|0\rangle\langle 0|$ state and constitutes a standard example for gate set tomography. In this example we don’t give an initialization to mGST, resulting in a random initialization to be generated automatically. As a result, more than one initialization attempt might be necessary to converge to a satisfying objective function value. We can tweak the number of allowed initializations with the *max_inits* parameter.

```
[8]: from pygsti.modelpacks import smq1Q_XYI as std
mdl_datagen = std.target_model().depolarize(0.01).randomize_with_unitary(
    0.01) # use pygsti-function to add noise
X_true, E_true, rho_true = compatibility.pygsti_model_to_arrays(
    mdl_datagen, basis='std') # turn pygsti model object into numpy arrays

pdim = 2
r = pdim**2

l = 7
d = 3
```

```
n_povm = 2
rK = 2
```

```
[9]: sequence_count = 100
meas_samples = 1e3
J_rand = np.array(random.sample(range(d**l), sequence_count))
J = np.array([low_level_jit.local_basis(ind, d, l) for ind in J_rand])
y = np.real(np.array([[E_true[i].conj()@low_level_jit.contract(X_true, j)
                      @ rho_true for j in J] for i in range(n_povm)]))
y_sampled = additional_fns.sampled_measurements(y, meas_samples).copy()
```

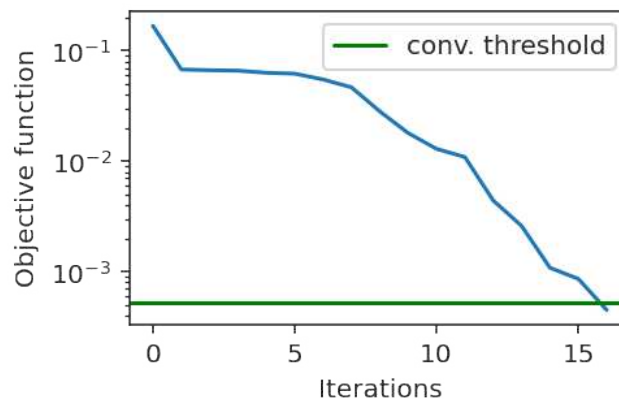
For the following run we set the parameter “testing = True”, which plots the objective function over the number of iterations for every initialization attempt. This helps in finding problems, for instance if max_iter was set too low to allow convergence, the plot will show continued decrease of the objective function up until the iteration limit, without reaching the success threshold. Another problem that can come up is model mismatch, for instance if rK is set to 1, but the actual gate set can not be well approximated by a Rank 1 Channel. Then the default success criterion might be too stringent, since no gate set in the Rank 1 model class can attain a low enough error. This can be fixed by setting the optional variable “threshold_multiplier” to a higher value (the default is threshold_multiplier = 3).

```
[10]: bsize = 50
t0 = time.time()
K, X, E, rho, res_list = algorithm.run_mGST(y_sampled, J, l, d, r, rK, n_povm,
    bsize, meas_samples, method='SFN',
    max_inits=5, max_iter=100, final_iter=10,
    target_rel_prec=1e-4, testing=True)
print('Mean variation error:', additional_fns.MVE(
    X_true, E_true, rho_true, X, E, rho, d, l, n_povm))
```

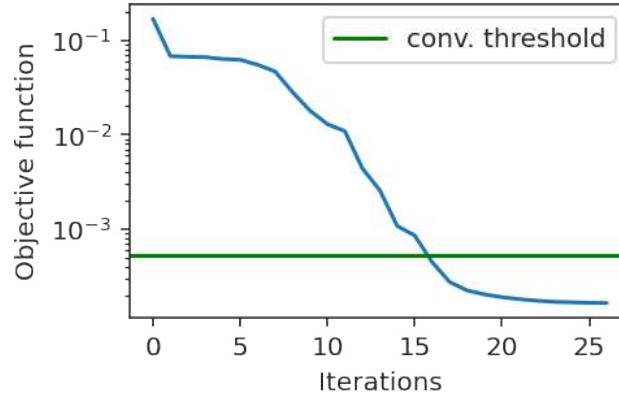
Starting optimization...

15%|

| 15/100 [00:04<00:23, 3.69it/s]



```
Initialization successful, improving estimate over full data...
100%|          | 10/10 [00:03<00:00, 3.00it/s]
```



```
#####
```

```
    Convergence criterion satisfied
    Final objective function value 0.00016605459935186234 with # of
initializations: 1
    Total runtime: 7.849793910980225
Mean variation error: (0.005836489826679743, 0.028832392621800873)
```

1.2 GST on a 3-level system

This example considers a qutrit gate set made up of a qutrit Hadamard gate, as well as X- and Z-gate defined on two-level subspaces. The gate set was previously used for gate set tomography in <https://arxiv.org/pdf/2210.04857.pdf>

```
[11]: pdim = 3
      r = pdim**2
      l = 8
      d = 6
      rK = 2
      n_povm = 3
```

```
[12]: w = np.exp(1j*2*np.pi/3)
      K_true = np.zeros((6, 1, 3, 3)).astype(np.complex128)
      K_true[0] = np.array([[1, 0, 0], [0, 1, 0], [0, 0, 1]]) # Identity
      K_true[1] = np.array([[1, 1, 1], [1, w, w**2], [1, w**2, w]]
                          )/np.sqrt(3) # Hadamard
      # Single qubit X-gate on the {|100>, |010>} subspace
      K_true[2] = np.array([[0, 1, 0], [1, 0, 0], [0, 0, 1]])
      # Single qubit X-gate on the {|010>, |001>} subspace
      K_true[3] = np.array([[1, 0, 0], [0, 0, 1], [0, 1, 0]])
      K_true[4] = np.array([[1, 0, 0], [0, w, 0], [0, 0, 1]])
```

```

        ) # Phase gate on the |010> state
K_true[5] = np.array([[1, 0, 0], [0, 1, 0], [0, 0, w]])
        ) # Phase gate on the |001> state

```

We can also be more creative with the noise model and generate random noise channels that have variable distance to the identity channels:

```

[13]: # Gate set tensor consisting of all gate superoperators
X_ideal = np.einsum('ijkl,ijnm -> iknlm', K_true,
                    K_true.conj()).reshape(d, r, r)
# Generates Kraus representations of random channels; The parameter a controls
# the distance from the identity channel
K_Lambda = additional_fns.randKrausSet(d, r, 9, a=0.2)
X_id = np.array([np.eye(r) for _ in range(d)]) # Identity channels

Lambda = np.einsum('ijkl,ijnm -> iknlm', K_Lambda, K_Lambda.conj())
        ).reshape(d, r, r) # Gate set tensor of noise channels
# Noise channels applied to the ideal gates
X_true = np.einsum('ijk,ikl -> ij1', X_ideal, Lambda)

K_depol = additional_fns.depol(pdimm, 0.02) # Kraus-rep of depolarizing channel
G_depol = np.einsum('jkl,jnm -> knlm', K_depol, K_depol.conj()).reshape(r, r)
# |100> initial state with depolarizing noise
rho_true = G_depol@np.array([[1, 0, 0], [0, 0, 0],
                             [0, 0, 0]]).reshape(-1).astype(np.complex128)

# Computational basis measurement:
E1 = np.array([[1, 0, 0], [0, 0, 0], [0, 0, 0]]).reshape(-1)
E2 = np.array([[0, 0, 0], [0, 1, 0], [0, 0, 0]]).reshape(-1)
E3 = np.array([[0, 0, 0], [0, 0, 0], [0, 0, 1]]).reshape(-1)
E_true = np.array([E1, E2, E3]).astype(np.complex128) # Full POVM

[14]: # A handy function to test whether a gate set satisfies all positivity and
# normalization constraints:
additional_fns.is_positive(X_true, E_true, rho_true)

```

```

Gate 0 positive: True
Gate 0 trace preserving: True
Gate 1 positive: True
Gate 1 trace preserving: True
Gate 2 positive: True
Gate 2 trace preserving: True
Gate 3 positive: True
Gate 3 trace preserving: True
Gate 4 positive: True
Gate 4 trace preserving: True
Gate 5 positive: True

```

```
Gate 5 trace preserving: True
Initial state positive: True
Initial state normalization: (1.0000000000000002+0j)
POVM valid: True
```

```
[15]: N = 300 # Number of sequences
meas_samples = 1e3 # Number of samples per sequences
# generate random numbers between 0 and  $d^l - 1$ 
J_rand = np.array(random.sample(range(d**l), N))
# turn random numbers into gate instructions
J = np.array([low_level_jit.local_basis(ind, d, l) for ind in J_rand])
y = np.real(np.array([[E_true[i].conj()@low_level_jit.contract(X_true,
    ↪j)@rho_true for j in J]
                    for i in range(n_povm)])) # obtain ideal output probabilities
# simulate finite sampling statistics
y_sampled = additional_fns.sampled_measurements(y, meas_samples).copy()

[16]: bsize = 80 # The batch size on which the optimization is started
K, X, E, rho, res_list = algorithm.run_mGST(y_sampled, J, l, d, r, rK, n_povm,
    ↪bsize, meas_samples, method='SFN',
                                max_inits=10, max_iter=300, final_iter=30,
                                target_rel_prec=1e-4)
plt.semilogy(res_list) # plot the objective function over the iterations
plt.tight_layout(rect = [.2,.2,.8,.8])
plt.show()
print('Mean variation error:', additional_fns.MVE(X_true, E_true, rho_true, X,
    ↪E,
    rho, d, l, n_povm)[0]) # output the final mean variation error
```

Starting optimization...

100%| | 300/300 [03:44<00:00, 1.33it/s]

Run 0 failed, trying new initialization...

21%| | 62/300 [00:45<02:55, 1.36it/s]

Initialization successful, improving estimate over full data...

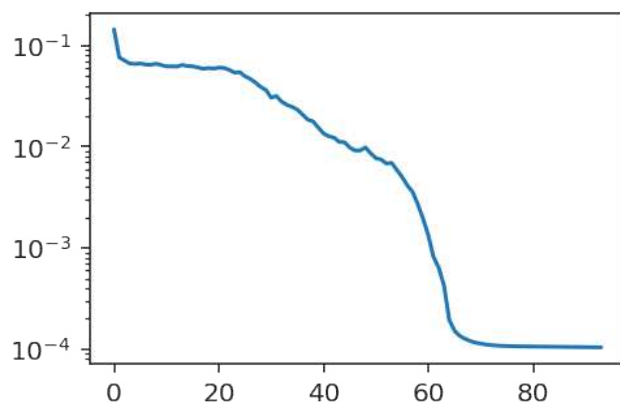
100%| | 30/30 [00:58<00:00, 1.94s/it]

#####

Convergence criterion satisfied

Final objective function value 0.00010438760324681996 with # of
initializations: 2

Total runtime: 328.79967641830444



Mean variation error: 0.008863883369235678