

REPORT DOCUMENTATION PAGE			Form Approved OMB NO. 0704-0188		
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA, 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.					
1. REPORT DATE (DD-MM-YYYY)		2. REPORT TYPE		3. DATES COVERED (From - To)	
		New Reprint		-	
4. TITLE AND SUBTITLE Quantum tomography via compressed sensing: error bounds, sample complexity and efficient estimators			5a. CONTRACT NUMBER		
			W911NF-10-1-0231		
			5b. GRANT NUMBER		
			5c. PROGRAM ELEMENT NUMBER		
			411359		
6. AUTHORS Steven T Flammia, David Gross, Yi-Kai Liu, Jens Eisert			5d. PROJECT NUMBER		
			5e. TASK NUMBER		
			5f. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAMES AND ADDRESSES			8. PERFORMING ORGANIZATION REPORT NUMBER		
Duke University 2200 West Main Street Suite 710 Durham, NC 27705 -4010					
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) U.S. Army Research Office P.O. Box 12211 Research Triangle Park, NC 27709-2211			10. SPONSOR/MONITOR'S ACRONYM(S) ARO		
			11. SPONSOR/MONITOR'S REPORT NUMBER(S) 58103-PH-MQC.18		
12. DISTRIBUTION AVAILABILITY STATEMENT Approved for public release; distribution is unlimited.					
13. SUPPLEMENTARY NOTES The views, opinions and/or findings contained in this report are those of the author(s) and should not be construed as an official Department of the Army position, policy or decision, unless so designated by other documentation.					
14. ABSTRACT Intuitively, if a density operator has small rank, then it should be easier to estimate from experimental data, since in this case only a few eigenvectors need to be learned. We prove two complementary results that confirm this intuition. Firstly, we show that a low-rank density matrix can be estimated using fewer copies of the state, i.e. the					
15. SUBJECT TERMS quantum tomography, compressed sensing					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	15. NUMBER OF PAGES	19a. NAME OF RESPONSIBLE PERSON
a. REPORT	b. ABSTRACT	c. THIS PAGE			Jungsang Kim
UU	UU	UU	UU		19b. TELEPHONE NUMBER
					919-660-5258

Report Title

Quantum tomography via compressed sensing: error bounds, sample complexity and efficient estimators

ABSTRACT

Intuitively, if a density operator has small rank, then it should be easier to estimate from experimental data, since in this case only a few eigenvectors need to be learned. We prove two complementary results that confirm this intuition. Firstly, we show that a low-rank density matrix can be estimated using fewer copies of the state, i.e. the sample complexity of tomography decreases with the rank. Secondly, we show that unknown low-rank states can be reconstructed from an incomplete set of measurements, using techniques from compressed sensing and matrix completion. These techniques use simple Pauli measurements, and their output can be certified without making any assumptions about the unknown state. In this paper, we present a new theoretical analysis of compressed tomography, based on the restricted isometry property for low-rank matrices. Using these tools, we obtain near-optimal error bounds for the realistic situation where the data contain noise due to finite statistics, and the density matrix is full-rank with decaying eigenvalues. We also obtain upper bounds on the sample complexity of compressed tomography, and almost-matching lower bounds on the sample complexity of any procedure using adaptive sequences of Pauli measurements. Using numerical simulations, we compare the performance of two compressed sensing estimators—the matrix Dantzig selector and the matrix Lasso—with standard maximum-likelihood estimation (MLE). We find that, given comparable experimental resources, the compressed sensing estimators consistently produce higher fidelity state reconstructions than MLE. In addition, the use of an incomplete set of measurements leads to faster classical processing with no loss of accuracy. Finally, we show how to certify the accuracy of a low-rank estimate using direct fidelity estimation, and describe a method for compressed quantum process tomography that works for processes with small Kraus rank and requires only Pauli eigenstate preparations and Pauli measurements.

REPORT DOCUMENTATION PAGE (SF298)
(Continuation Sheet)

Continuation for Block 13

ARO Report Number 58103.18-PH-MQC
Quantum tomography via compressed sensing: ...

Block 13: Supplementary Note

© 2012 . Published in New Journal of Physics, Vol. 14 (9) (2012), (4 (9). DoD Components reserve a royalty-free, nonexclusive and irrevocable right to reproduce, publish, or otherwise use the work for Federal purposes, and to authorize others to do so (DODGARS §32.36). The views, opinions and/or findings contained in this report are those of the author(s) and should not be construed as an official Department of the Army position, policy or decision, unless so designated by other documentation.

Approved for public release; distribution is unlimited.

Quantum tomography via compressed sensing: error bounds, sample complexity and efficient estimators

Steven T Flammia^{1,5}, David Gross², Yi-Kai Liu³ and Jens Eisert⁴

¹ Department of Computer Science and Engineering, University of Washington, Seattle, WA, USA

² Institute of Physics, University of Freiburg, 79104 Freiburg, Germany

³ National Institute of Standards and Technology, Gaithersburg, MD, USA

⁴ Dahlem Center for Complex Quantum Systems, Freie Universität Berlin, 14195 Berlin, Germany

E-mail: sflammia@cs.washington.edu

New Journal of Physics **14** (2012) 095022 (28pp)

Received 23 May 2012

Published 27 September 2012

Online at <http://www.njp.org/>

doi:10.1088/1367-2630/14/9/095022

Abstract. Intuitively, if a density operator has small rank, then it should be easier to estimate from experimental data, since in this case only a few eigenvectors need to be learned. We prove two complementary results that confirm this intuition. Firstly, we show that a low-rank density matrix can be estimated using fewer copies of the state, i.e. the *sample complexity* of tomography decreases with the rank. Secondly, we show that unknown low-rank states can be reconstructed from an *incomplete* set of measurements, using techniques from compressed sensing and matrix completion. These techniques use simple Pauli measurements, and their output can be certified without making any assumptions about the unknown state. In this paper, we present a new theoretical analysis of compressed tomography, based on the *restricted isometry property* for low-rank matrices. Using these tools, we obtain near-optimal error bounds for the realistic situation where the data contain noise due to finite statistics, and the density matrix is full-rank with decaying eigenvalues. We also obtain upper bounds on the sample complexity of compressed tomography, and almost-matching lower bounds on the sample complexity of any procedure using adaptive sequences of Pauli measurements. Using numerical simulations, we

⁵ Author to whom any correspondence should be addressed.



Content from this work may be used under the terms of the [Creative Commons Attribution-NonCommercial-ShareAlike 3.0 licence](https://creativecommons.org/licenses/by-nc-sa/3.0/). Any further distribution of this work must maintain attribution to the author(s) and the title of the work, journal citation and DOI.

compare the performance of two compressed sensing estimators—the matrix Dantzig selector and the matrix Lasso—with standard maximum-likelihood estimation (MLE). We find that, given comparable experimental resources, the compressed sensing estimators consistently produce higher fidelity state reconstructions than MLE. In addition, the use of an incomplete set of measurements leads to faster classical processing with no loss of accuracy. Finally, we show how to certify the accuracy of a low-rank estimate using direct fidelity estimation, and describe a method for compressed quantum process tomography that works for processes with small Kraus rank and requires only Pauli eigenstate preparations and Pauli measurements.

Contents

1. Introduction	2
1.1. Related work	5
1.2. Notation and outline	6
2. Theory	7
2.1. Random Pauli measurements	7
2.2. Reconstructing the density matrix	7
2.3. Error bounds	8
2.4. Sample complexity	9
3. Lower bounds	11
4. Certifying the state estimate	15
5. Numerical simulations	18
5.1. Setting the estimator parameters λ and μ	18
5.2. Time needed to switch measurement settings	18
5.3. Other simulation parameters	19
5.4. Results and analysis	20
6. Process tomography	21
6.1. Our method	22
6.2. Related work	23
7. Conclusion	24
Acknowledgments	25
References	25

1. Introduction

In recent years, there has been amazing progress in studying complex quantum mechanical systems under controlled laboratory conditions [1]. Quantum tomography of states and processes is an invaluable tool used in many such experiments, because it enables a complete characterization of the state of a quantum system or process (see, e.g., [2–16]). Unfortunately, tomography is very resource intensive, and scales exponentially with the size of the system. For example, a system of n spin-1/2 particles (qubits) has a Hilbert space with dimension $d = 2^n$, and the state of the system is described by a density matrix with $d^2 = 4^n$ entries.

Recently, a new approach to tomography was proposed: *compressed* quantum tomography, based on techniques from compressed sensing [17, 18]. The basic idea is to concentrate on states that are well approximated by density matrices of rank $r \ll d$. This approach can be applied to many realistic experimental situations, where the ideal state of the system is pure, and physical constraints (e.g. low temperature or the locality of interactions) ensure that the actual (noisy) state still has low entropy.

This approach is convenient because it does not require detailed knowledge about the system. However, note that when such knowledge is available, one can use alternative formulations of compressed tomography, with different notions of sparsity, to further reduce the dimensionality of the problem [19]. We will compare these methods in section 6.2.

The main challenge in compressed tomography is how to exploit this low-rank structure, when one does not know the subspace on which the state is supported. Consider the example of a pure quantum state. Since pure states are specified by only $O(d)$ numbers, it seems plausible that one could be reconstructed after measuring only $O(d)$ observables, compared with $O(d^2)$ for a general mixed state. While this intuition is indeed correct [20–23], it is a challenge to devise a practical tomography scheme that takes advantage of this. In particular, one is restricted to those measurements that can be easily performed in the laboratory; furthermore, one then has to find a pure state consistent with measured data [24], preferably by some procedure that is computationally efficient (note that, in general, finding minimum-rank solutions is NP-hard, hence computationally intractable [25]).

Compressed tomography provides a solution that meets all these practical requirements [17, 18]. It requires measurements of two-outcome Pauli observables, which are feasible in many experimental systems. In total, it uses a random subset of $m = O(rd \log d)$ Pauli observables, which is just slightly more than the $O(rd)$ degrees of freedom that specify an arbitrary rank r state. Then the density matrix ρ is reconstructed by solving a convex program. This can be done efficiently using general-purpose semidefinite program (SDP) solvers [26] or specialized algorithms for larger instances [27–29]. The scheme is robust to noise and continues to perform well when the measurements are imprecise or when the state is only close to a low-rank state.

Here, we follow up on [17, 18] by giving a stronger (and completely different) theoretical analysis, which is based on the *restricted isometry property* (RIP) [30–32]. This answers a number of questions that could not be addressed satisfactorily using earlier techniques based on dual certificates. Firstly, how large is the error in our estimated density matrix when the true state is full-rank with decaying eigenvalues? We show that the error is not much larger than the ‘tail’ of the eigenvalue spectrum of the true state. Secondly, how large is the *sample complexity* of compressed tomography, i.e. how many copies of the unknown state are needed to estimate its density matrix? We show that compressed tomography achieves nearly optimal sample complexity among all procedures using Pauli measurements and, surprisingly, the sample complexity of compressed tomography is nearly independent of the number of measurement settings m as long as $m \geq \Omega(rd \text{ poly } \log d)$.

In addition, we use numerical simulations to investigate the question: given a fixed time T during which an experiment can be run, is it better to do compressed tomography or full tomography, i.e. is it better to use a few measurement settings and repeat them many times or do all possible measurements with fewer repetitions? For the situations we simulate, we find that compressed tomography provides better accuracy at a reduced computational cost compared to standard maximum-likelihood estimation (MLE).

Finally, we provide two useful tools: a procedure for certifying the accuracy of low-rank state estimates and a very simple compressed sensing technique for quantum process tomography.

We now describe these results in more detail.

Theoretical analysis using RIP. We use a fundamental geometric fact: the manifold of rank- r matrices in $\mathbb{C}^{d \times d}$ can be embedded into $O(rd \text{ poly } \log d)$ dimensions, with low distortion in the two-norm. An embedding that does this is said to satisfy the RIP [32]. In [30], it was shown that such an embedding can be realized using the expectation values of a random subset of $O(rd \text{ poly } \log d)$ Pauli matrices. This implies the existence of the so-called ‘universal’ methods for low-rank matrix recovery: there exists a *fixed* set of $O(rd \text{ poly } \log d)$ Pauli measurements, which has the ability to reconstruct *every* rank- r $d \times d$ matrix. Moreover, with high probability, a random choice of Pauli measurements will achieve this. (The earlier results of [17] placed the quantifiers in the opposite order: for every rank- r $d \times d$ matrix ρ , most sets of $O(rd \text{ poly } \log d)$ Pauli measurements can reconstruct that *particular* matrix ρ .)

Intuitively, the RIP says that a set of random Pauli measurements is sensitive to all low-rank errors simultaneously. This is important, because it implies stronger error bounds for low-rank matrix recovery [31]. These bounds show that, when the unknown matrix ρ is *full-rank*, our method returns a (certifiable) rank- r approximation of ρ that is almost as good as the best such approximation (corresponding to the truncated eigenvalue decomposition of ρ).

In [31], these error bounds were used to show the accuracy of certain compressed sensing estimators, for measurements with additive Gaussian noise. Here, we use them to upper-bound the sample complexity of our compressed tomography scheme. (That is, we bound the errors due to estimating each Pauli expectation value from a finite number of experiments.) Roughly speaking, we show that our scheme uses $O(r^2 d^2 \log d)$ copies to characterize a rank- r state (up to constant error in trace norm). When $r = d$, this agrees with the sample complexity of full tomography. Our proof assumes a binomial noise model, but minor modifications could extend this result to other relevant noise models, such as multinomial, Gaussian or Poissonian noise.

Furthermore, we show an information-theoretic lower bound for tomography of rank- r states using adaptive sequences of single-copy Pauli measurements: at least $\Omega(r^2 d^2 / \log d)$ copies are needed to obtain an estimate with constant accuracy in the trace distance. This generalizes a result from [33] for pure states. Therefore, our upper bound on the sample complexity of compressed tomography is nearly tight, and compressed tomography nearly achieves the optimal sample complexity among all possible methods using Pauli measurements.

Our observation that incomplete sets of observables are often sufficient to unambiguously specify a state gives rise to a new degree of freedom when designing experiments: when aiming to reduce statistical noise in the reconstruction, one can either estimate a small set of observables relatively accurately or else a large (e.g. complete) set of observables relatively coarsely. Our bounds (as well as our numerics) show that, remarkably, over a very large range of m the *only* quantity relevant for the reconstruction error is t , the total number of experiments performed. It does not matter over how many observables the repetitions are distributed. Thus, fixing t and varying m , the reduction in the number of observables and the increase in the number of measurements per observable have no net effect with regard to the fidelity of the estimate as long as $m \geq \Omega(rd \text{ poly } \log d)$. This finding holds independently of whether fidelity or trace distance is used to measure the reconstruction error, and we believe that it is plausible that a similar phenomenon occurs for other cost functions as well.

Certification. We generalize the technique of direct fidelity estimation (DFE) [33, 34] to work with low-rank states. Thus, one can use compressed tomography to get an estimated density matrix $\hat{\rho}$, and use DFE to check whether $\hat{\rho}$ agrees with the true state ρ . This check is guaranteed to be sound even if the true state ρ is not approximately low rank. Our extension of DFE may be of more general interest, since it can be used to efficiently certify *any* estimate $\hat{\rho}$ regardless of whether it was obtained using compressed sensing or not, as long as the rank r of the *estimate* is small (and regardless of the ‘true’ rank).

Numerical simulations. We compare the performance of several different estimators (methods for reconstructing the unknown density matrix). They include: constrained trace-minimization (also known as the matrix Dantzig selector), least squares with trace-norm regularization (or the matrix Lasso), as well as a standard MLE [35–37] for comparison.

We observe that our estimators outperform MLE in essentially all aspects, with the matrix Lasso giving the best results. The fidelity of the estimate is consistently higher using the compressed tomography estimators. Also, the accuracy of the compressed sensing estimates is (as mentioned above) fairly insensitive to the number of measurement settings m (assuming that the total time available to run the experiment is fixed). So by choosing $m \ll d^2$, one still obtains accurate estimates, but with much faster classical post-processing, since the size of the data set scales as $O(m)$ rather than $O(d^2)$.

It may be surprising to the reader that we outperform MLE, since it is often remarked (somewhat vaguely) that ‘MLE is optimal’. However, MLE is a general-purpose method that does not specifically exploit the fact that the state is low-rank. Also, the optimality results for MLE only hold asymptotically and for informationally complete measurements [38, 39]; for finite data [40] or for incomplete measurements, MLE can be far from optimal.

From these results, one can extract some lessons about how to use compressed tomography. Compressed tomography involves two separate ideas: (1) measuring an incomplete set of observables (i.e. choosing $m \ll d^2$) and (2) using trace minimization or regularization to reconstruct low-rank solutions. Usually, one does both of these things. Now, suppose that the goal is to reconstruct a low-rank state using as few samples as possible. Our results show that one can achieve this goal by doing (2) *without* (1). At the same time, there is no penalty in the quality of the estimate when doing (1), and there are practical reasons for doing it, such as reducing the size of the data set to speed up the classical post-processing.

Quantum process tomography. Finally, we adapt our method to perform tomography of processes with a small Kraus rank. Our method is easy to implement, since it requires only the ability to prepare eigenstates of Pauli operators and measure Pauli observables. In particular, we require *no* entangling gates or ancillary systems for the procedure. In contrast with [19], our method is not restricted to processes that are element-wise sparse in some known basis, as discussed in section 6.2. This is an important advantage in practice, because while the ideal (or intended) evolution of a system may be sparse in a known basis, it is often the case that the noise processes perturbing the ideal component are not sparse, and knowledge of these noise processes is key to improving the fidelity of a quantum device with some theoretical ideal.

1.1. Related work

While initial work on tomography considered only linear inversion methods [41], most subsequent works have largely focused on maximum likelihood methods and to a lesser extent on Bayesian methods for state reconstruction [35–37, 41–52].

However, there has recently been a flurry of work which seeks to transcend the standard MLE methods and improve on them in various ways. Our contributions can also be seen in this context.

One way in which alternatives to MLE are being pursued is through what we call *full-rank methods*. Here the idea is somewhat antithetical to ours: the goal is to output a full-rank density operator, rather than a rank deficient one. This is desirable in a context where one cannot make the approximation that rare events will never happen. Blume-Kohout's hedged MLE [53] and Bayesian mean estimation [52] are good examples of this type of estimator, as are the minimax estimator of [54] and the so-called Max-Ent estimators [55–58]. The latter are specifically for the setting where the measurement data are *not* informationally complete, and one tries to minimize the bias of the estimate by maximizing the entropy along the directions where one has no knowledge.

By contrast, our *low-rank* methods do not attempt to reconstruct the complete density matrix, but only a rank- r approximation, which is accurate when the true state is close to low-rank. From this perspective, our methods can be seen as a sort of Occam's razor, using as few fit parameters as possible while still agreeing with the data [59]. Furthermore, as we show here and elsewhere [17], informationally incomplete measurements can still provide faithful state reconstructions up to a small truncation error.

One additional feature of our methods is that we are deeply concerned with the *feasibility* of our estimators for a moderately large number of qubits (say, 10–15). In contrast to most of the existing literature, we adopt the perspective that it is not enough for an estimator to be asymptotically efficient in the number of copies for fixed d . We also want the scaling with respect to d to be optimal. We specifically take advantage of the fact that many states and processes are described by low-rank objects to reduce this complexity. In this respect, our methods are similar to tomographic protocols that are tailored to special ansatz classes of states, such as those recently developed for use with permutation-invariant states [60], matrix product states [61] or multi-scale entangled states [62].

Our error bounds are somewhat unique as well. Most previous work on error bounds used either standard resampling techniques or Bayesian methods [42, 43, 45, 48, 50, 52]. Very recently, Christandl and Renner [63] and Blume-Kohout [64] independently derived two closely related approaches for obtaining confidence regions that satisfy or nearly satisfy certain optimality criteria. In particular, the latter approaches can give very tight error bounds on an estimate, but they can be computationally challenging to implement for large systems. The error bounds which most closely resemble ours are of the 'large deviation type'; see, for example, the discussion in [38]. This is true for the new improved error bounds, as well as the original bounds proven in [17, 18]. These types of bounds are much easier to calculate in practice, which agrees with our philosophy on computational complexity, but may be somewhat looser than the optimal error bounds obtainable through other more computationally intensive methods such as those of [63, 64].

1.2. Notation and outline

We denote Pauli operators by P or P_i . We define $[n] = \{1, \dots, n\}$. The norms we use are the standard Euclidean vector norm $\|x\|_2$, the Frobenius norm $\|X\|_F = \sqrt{\text{Tr}(X^\dagger X)}$, the operator norm $\|X\| = \sqrt{\lambda_{\max}(X^\dagger X)}$ and the trace norm $\|X\|_{\text{tr}} = \text{Tr}|X|$, where $|X| = \sqrt{X^\dagger X}$. The unknown 'true' state is denoted by ρ and any estimators for ρ are given a hat: $\hat{\rho}$. The

expectation value of a random variable X is denoted by $\mathbb{E} X$. We denote by $\mathbb{H}^d$ the set of $d \times d$ Hermitian matrices.

The paper is organized as follows. In section 2, we detail the estimators and error bounds and then upper bound the sample complexity. In section 3, we derive lower bounds on the sample complexity. In section 4, we find an efficient method of certifying the state estimate. In section 5, we detail our numerical investigations. We show how our scheme can be applied to quantum channels in section 6 and conclude in section 7.

2. Theory

We describe our compressed tomography scheme in detail. Firstly, we describe the measurement procedure and the method for reconstructing the density matrix. Then we prove error bounds and analyze the sample complexity.

2.1. Random Pauli measurements

Consider a system of n qubits, and let $d = 2^n$. Let $\mathcal{P}$ be the set of all d^2 Pauli operators, i.e. matrices of the form $P = \sigma_1 \otimes \cdots \otimes \sigma_n$ where $\sigma_i \in \{I, \sigma^x, \sigma^y, \sigma^z\}$.

Our tomography scheme works as follows. Firstly, choose m Pauli operators, $P_1, \dots, P_m$, by sampling independently and uniformly at random from $\mathcal{P}$. (Alternatively, one can choose these Pauli operators randomly without replacement [65], but independent sampling greatly simplifies the analysis.) We will use t copies of the unknown quantum state ρ . For each $i \in [m]$, take t/m copies of the state ρ , measure the Pauli observable P_i on each one and average the measurement outcomes to obtain an estimate of the expectation value $\text{Tr}(P_i \rho)$. (We will discuss how to set m and t later. Intuitively, to learn a $d \times d$ density matrix with rank r , we will set $m \sim r d \log^6 d$ and $t \sim r^2 d^2 \log d$.)

To state this more concisely, we introduce some notation. Define the *sampling operator* to be a linear map $\mathcal{A} : \mathbb{H}^d \rightarrow \mathbb{R}^m$ defined for all $i \in [m]$ by

$$(\mathcal{A}(\rho))_i = \sqrt{\frac{d}{m}} \text{Tr}(P_i \rho). \quad (1)$$

(The normalization is chosen so that $\mathbb{E} \mathcal{A}^* \mathcal{A} = \mathcal{I}$, where $\mathcal{I}$ denotes the identity superoperator, and $\mathcal{A}^*$ is the adjoint of $\mathcal{A}$. That is, $\mathcal{A}^* : \mathbb{R}^m \rightarrow \mathbb{H}^d$ is the complex-conjugate transpose of the linear map $\mathcal{A}$.) We can then write the output of our measurement procedure as a vector

$$y = \mathcal{A}(\rho) + z, \quad (2)$$

where z represents statistical noise due to the finite number of samples. (More generally, if the measurements contain systematic or adversarial noise, these can also be represented by z . There are various error bounds for this situation, although they are necessarily more pessimistic than the ones we show here [66].)

2.2. Reconstructing the density matrix

We now show two methods for estimating the density matrix ρ . Both are based on the same intuition: find a matrix $X \in \mathbb{H}^d$ that fits the data y while minimizing the trace norm $\|X\|_{\text{tr}}$,

which serves as a surrogate for minimizing the rank of X . In both cases, this amounts to a convex program, which can be solved efficiently.

(We mention that at this point we do not require that our density operators are normalized to have unit trace. We will return to this point later in section 5.)

The first estimator is obtained by constrained trace-minimization (also known as the matrix Dantzig selector):

$$\hat{\rho}_{\text{DS}} = \arg \min_X \|X\|_{\text{tr}} \quad \text{s.t.} \quad \|\mathcal{A}^*(\mathcal{A}(X) - y)\| \leq \lambda. \quad (3)$$

Note that the constraint bounds the operator norm of $\mathcal{A}^*(\mathcal{A}(X) - y)$. The parameter λ should be set according to the amount of noise in the data y ; we will discuss this later.

The second estimator is obtained by least-squares linear regression with trace-norm regularization (also known as the matrix Lasso):

$$\hat{\rho}_{\text{Lasso}} = \arg \min_X \frac{1}{2} \|\mathcal{A}(X) - y\|_2^2 + \mu \|X\|_{\text{tr}}. \quad (4)$$

Again the regularization parameter μ should be set according to the noise level; we will discuss this later.

One additional point is that we do not require positivity of the output in our definition of the estimators (3) and (4). One can add this constraint (since it is convex) and the conclusions below remain unaltered. We will explicitly add positivity later on when we do numerical simulations, and discuss any trade-offs that result from this.

2.3. Error bounds

In previous works on compressed tomography [17, 18], error bounds were proved by constructing a ‘dual certificate’ [67] (using convex duality to characterize the solution to the trace-minimization convex program). Here we derive stronger bounds using a different tool, known as the RIP. The RIP for low-rank matrices was first introduced in [32], and was recently shown to hold for random Pauli measurements [30].

We say that the sampling operator $\mathcal{A}$ satisfies the rank- r RIP if there exists some constant $0 \leq \delta_r < 1$ such that, for all $X \in \mathbb{C}^{d \times d}$ with rank r ,

$$(1 - \delta_r) \|X\|_{\text{F}} \leq \|\mathcal{A}(X)\|_2 \leq (1 + \delta_r) \|X\|_{\text{F}}. \quad (5)$$

For our purposes, we can assume that X is Hermitian. Note that this notion of RIP is analogous to that used in [19], except that it holds over the set of low-rank matrices, rather than the set of sparse matrices.

The random Pauli sampling operator (1) satisfies RIP with high probability, provided that $m \geq Crd \log^6 d$ (for some absolute constant C); this was recently shown in [30]. We note, however, that this RIP result requires m to be larger by a poly $\log d$ factor compared to previous results based on dual certificates. Although m is slightly larger, the advantage is that when combined with the results of [31], this immediately implies strong error bounds for the matrix Dantzig selector and the matrix Lasso.

To state these improved error bounds precisely, we introduce some definitions. For the rest of section 2, let C, C_0, C_1, C'_0 and C'_1 be fixed absolute constants. For any quantum state ρ , we write $\rho = \rho_r + \rho_c$, where ρ_r is the best rank- r approximation to ρ (consisting of the largest r eigenvalues and eigenvectors), and ρ_c is the residual part. Now we have the following:

Theorem 1. Let $\mathcal{A}$ be the random Pauli sampling operator (1) with $m \geq Crd \log^6 d$. Then, with high probability, the following holds:

Let $\hat{\rho}_{\text{DS}}$ be the matrix Dantzig selector (3), and choose λ so that $\|\mathcal{A}^*(z)\| \leq \lambda$. Then

$$\|\hat{\rho}_{\text{DS}} - \rho\|_{\text{tr}} \leq C_0 r \lambda + C_1 \|\rho_c\|_{\text{tr}}.$$

Alternatively, let $\hat{\rho}_{\text{Lasso}}$ be the matrix Lasso (4), and choose μ so that $\|\mathcal{A}^*(z)\| \leq \mu/2$. Then

$$\|\hat{\rho}_{\text{Lasso}} - \rho\|_{\text{tr}} \leq C'_0 r \mu + C'_1 \|\rho_c\|_{\text{tr}}.$$

In these error bounds, the first term depends on the statistical noise z . This, in turn, depends on the number of copies of the state that are available in the experiment; we will discuss this in the next section. The second term is the rank- r approximation error. It is clearly optimal, up to a constant factor.

Proof. These error bounds follow from the RIP as shown by Theorem 2.1 in [30], and a straightforward modification of Lemma 3.2 in [31] to bound the error in trace norm rather than Frobenius norm (this is similar to the proof of Theorem 5 in [66]).

The modification of Lemma 3.2 in [31] is as follows⁶. (For the remainder of this section, equation numbers of the form (III.x) refer to [31].) In the case of the Dantzig selector, let $H = \hat{\rho}_{\text{DS}} - \rho$ be the difference of the matrix Dantzig selector and the quantum state. Following equation (III.8) and decomposing $H = H_0 + H_c$ exactly as in [31], we can obtain the following bound:

$$\|H\|_{\text{tr}} \leq \|H_0\|_{\text{tr}} + \|H_c\|_{\text{tr}} \leq 2\|H_0\|_{\text{tr}} + 2\|\rho_c\|_{\text{tr}}, \quad (6)$$

where we used the triangle inequality and equation (III.8). Then, at the end of the proof, we write

$$\begin{aligned} \|H_0\|_{\text{tr}} &\leq \sqrt{2r} \|H_0\|_{\text{F}} \leq \sqrt{2r} \|H_0 + H_1\|_{\text{F}} \\ &\leq C_1 4\sqrt{2r} \lambda + C_1 2\sqrt{2} \delta_{4r} \|\rho_c\|_{\text{tr}}, \end{aligned} \quad (7)$$

where we used Cauchy–Schwarz, the fact that H_0 and H_1 are orthogonal, the bound on $\|H_0 + H_1\|_{\text{F}}$ following equation (III.13) and equation (III.7). Combining (6) and (7) gives our desired error bound. The error bound for the Lasso is obtained in a similar way; see section III.G in [31]. $\square$

2.4. Sample complexity

Here we bound the sample complexity of our tomography scheme; that is, we bound the number of copies of the unknown quantum state ρ that are needed to obtain our estimate up to some accuracy. What we show, roughly speaking, is that $t = O((\frac{rd}{\varepsilon})^2 \log d)$ copies are sufficient to reconstruct an estimate of an unknown rank r state up to accuracy ε in the trace distance. For comparison, note that when $r = d$, and one does full tomography, $O(d^4/\varepsilon^2)$ copies are sufficient to estimate a full-rank state with accuracy ε in trace distance⁷.

⁶ Note that [31] contains a typo in Lemma 3.2: on the right-hand side, in the second term, it should be $\|M_c\|_*/\sqrt{r}$ and not $\|M_c\|_*/r$.

⁷ To see this, let $P_1, \dots, P_{d^2}$ denote the Pauli matrices. For each $i \in [d^2]$, measure P_i $O(d^2/\varepsilon^2)$ times, to estimate its expectation value with additive error $\pm O(\varepsilon/d)$. Equivalently, one estimates the expectation value of $P_i/\sqrt{d}$ with additive error $\pm O(\varepsilon/d^{3/2})$. Using linear inversion, one obtains an estimated density matrix with additive error $O(\varepsilon/d^{1/2})$ in Frobenius norm, which implies error $O(\varepsilon)$ in trace norm.

To make this claim precise, we need to specify how we construct our data vector y from the measurement outcomes on the t copies of the state ρ . For the matrix Dantzig selector, suppose that

$$t \geq 2C_4(C_0 r/\varepsilon)^2 d(d+1) \log d \quad (8)$$

for some constants $C_4 > 1$ and $\varepsilon \leq C_0$. (For the matrix Lasso, substitute C'_0 for C_0 in these equations.) We construct an estimate of $\mathcal{A}(\rho)$ as follows: for each $i \in [m]$, we take t/m copies of ρ , measure the random Pauli observable P_i on each of the copies and use this to estimate $\text{Tr}(P_i \rho)$. Then let y be the resulting estimate of $\mathcal{A}(\rho)$, and let $z = y - \mathcal{A}(\rho)$. Everything else is defined exactly as in theorem 1.

Theorem 2. *Given $t = O((\frac{rd}{\varepsilon})^2 \log d)$ copies of ρ as in equation (8) and measured as discussed above, the following holds with high probability over the measurement outcomes.*

Let $\hat{\rho}_{\text{DS}}$ be the matrix Dantzig selector (3), and set $\lambda = \varepsilon/(C_0 r)$ for some $\varepsilon > 0$. Then

$$\|\hat{\rho}_{\text{DS}} - \rho\|_{\text{tr}} \leq \varepsilon + C_1 \|\rho_c\|_{\text{tr}}.$$

Alternatively, let $\hat{\rho}_{\text{Lasso}}$ be the matrix Lasso (4), and set $\mu = \varepsilon/(C'_0 r)$. Then

$$\|\hat{\rho}_{\text{Lasso}} - \rho\|_{\text{tr}} \leq \varepsilon + C'_1 \|\rho_c\|_{\text{tr}}.$$

Proof. Our claim reduces to the following question: if we fix some value of $\lambda > 0$, how many copies of ρ are needed to ensure that the measurement data y satisfy $\|\mathcal{A}^*(y - \mathcal{A}(\rho))\| \leq \lambda$? Then one can apply theorem 1 to obtain an error bound for our estimate of ρ .

Let t be the number of copies of ρ . Say we fix the measurement operator $\mathcal{A}$, i.e. we fix the choice of the Pauli observables $P_1, \dots, P_m$. (The measurement outcomes are still random, however.) For $i \in [m]$ and $j \in [t/m]$, let $B_{ij} \in \{1, -1\}$ be the outcome of the j th repetition of the experiment that measures the i th Pauli observable P_i . Note that $\mathbb{E} B_{ij} = \text{Tr}(P_i \rho)$. Then construct the vector $y \in \mathbb{R}^m$ containing the estimated expectation values (scaled by $\sqrt{d/m}$):

$$y_i = \sqrt{\frac{d}{m}} \cdot \frac{m}{t} \sum_{j=1}^{t/m} B_{ij}, \quad i \in [m]. \quad (9)$$

Note that $\mathbb{E} y = \mathcal{A}(\rho)$.

We will bound the deviation $\|\mathcal{A}^*(y - \mathcal{A}(\rho))\|$, using the matrix Bernstein inequality. First, we write

$$\mathcal{A}^*(y) = \sqrt{\frac{d}{m}} \sum_{i=1}^m P_i y_i = \frac{d}{t} \sum_{i=1}^m \sum_{j=1}^{t/m} P_i B_{ij}, \quad (10)$$

and also

$$\mathcal{A}^* \mathcal{A}(\rho) = \frac{d}{m} \sum_{i=1}^m P_i \text{Tr}(P_i \rho). \quad (11)$$

We can now write $\mathcal{A}^*(y - \mathcal{A}(\rho))$ as a sum of independent (but not identical) matrix-valued random variables:

$$\mathcal{A}^*(y - \mathcal{A}(\rho)) = \sum_{i=1}^m \sum_{j=1}^{t/m} X_{ij}, \quad X_{ij} = \frac{d}{t} P_i [B_{ij} - \text{Tr}(P_i \rho)]. \quad (12)$$

Note that $\mathbb{E} X_{ij} = 0$ and $\|X_{ij}\| \leq 2d/t =: R$. Also, for the second moment we have

$$\mathbb{E}(X_{ij}^2) = \mathbb{E}\left(\frac{d^2}{t^2} I[B_{ij} - \text{Tr}(P_i \rho)]^2\right) = \frac{d^2}{t^2} I[1 - \text{Tr}(P_i \rho)^2]. \quad (13)$$

Then we have

$$\sigma^2 := \left\| \sum_{ij} \mathbb{E}(X_{ij}^2) \right\| = \sum_{ij} \frac{d^2}{t^2} [1 - \text{Tr}(P_i \rho)^2] \leq t \cdot \frac{d^2}{t^2} = \frac{d^2}{t}. \quad (14)$$

Now the matrix Bernstein inequality (Theorem 1.4 in [68]) implies that

$$\Pr[\|\mathcal{A}^*(y - \mathcal{A}(\rho))\| \geq \lambda] \leq d \cdot \exp\left(-\frac{\lambda^2/2}{\sigma^2 + (R\lambda/3)}\right) \leq d \cdot \exp\left(-\frac{t\lambda^2/2}{d(d+1)}\right) \quad (15)$$

(where we assumed that $\lambda \leq 1$).

For the matrix Dantzig selector, we set $\lambda = \varepsilon/(C_0 r)$, and we get that, for any $t \geq 2C_4 \lambda^{-2} d(d+1) \log d = 2C_4 (C_0 r/\varepsilon)^2 d(d+1) \log d$,

$$\Pr\left[\|\mathcal{A}^*(y - \mathcal{A}(\rho))\| \geq \frac{\varepsilon}{C_0 r}\right] \leq d \cdot \exp(-C_4 \log d) = d^{1-C_4}, \quad (16)$$

which is exponentially small in C_4 . Plugging into theorem 2 completes the proof of our claim. A similar argument works for the matrix Lasso. $\square$

3. Lower bounds

How good are the sample complexities of our algorithms? Here we go a long way toward answering this question by proving nearly tight lower bounds on the sample complexity for generic rank r quantum states using single-copy Pauli measurements. Previous work on single-copy lower bounds has only treated the case of pure states [33].

Roughly speaking, we show the following, which we will make precise at the end of the section.

Theorem 3 (Imprecise formulation). *The number of copies t needs to grow at least as fast as $\Omega(r^2 d^2 / \log d)$ to guarantee a constant trace-norm confidence interval for all rank- r states.*

The argument proceeds in three steps. First, we fix our notion of risk to be the minimax risk, meaning we seek to minimize our worst-case error according to some error metric such as the trace distance. We want to know how many copies of the unknown state we need to make this minimax risk an arbitrarily small constant. For a fixed set of two-outcome measurements, say Pauli measurements, we then show that some states require many copies to achieve this. In particular, these states have the property that they are globally distinguishable (their trace distance is bounded from below by a constant) but their (Pauli) measurement statistics look approximately the same (each measurement outcome is close to unbiased). The more such states there are, the more copies we need to distinguish between them all using solely Pauli measurements. Finally, we use a randomized argument to show that, in fact, there are exponentially many such states. This yields the desired lower bound on the sample complexity.

Let Σ be some set of density operators. We want to put lower bounds on the performance of any estimation protocol for states in Σ . (We do not initially restrict ourselves to states with low rank.)

We assume that the protocol has access to t copies of an unknown state $\rho \in \Sigma$ on which it makes measurements one by one. At the i th step, it has to decide which observable to measure. Let us restrict ourselves to two-outcome measurements, which can be described by positive operator-valued measures (POVMs) $\{\Pi_i, \mathbb{1} - \Pi_i\}$, where each Π_i satisfies $0 \leq \Pi_i \leq \mathbb{1}$ and may be chosen from a set $\mathcal{P}$. (We do not initially restrict ourselves to Pauli measurements, either.) We allow the choice of the i th observable to depend on the previous outcomes. We refer to the random variables which jointly describe the choice of the i th measurement and its random outcome as Y_i . At the end, these are mapped to an estimate $\hat{\rho}(Y_1, \dots, Y_t) \in \Sigma$.

In other words, an estimation protocol is specified by a set of functions

$$\begin{aligned}\Pi_i &: Y_1, \dots, Y_{i-1} \mapsto \mathcal{P}, \\ \hat{\rho} &: Y_1, \dots, Y_t \mapsto \Sigma.\end{aligned}$$

Consider a distance measure $\Delta : \Sigma \times \Sigma \rightarrow \mathbb{R}$ on the states in Σ . (For example, this could be the trace distance or the infidelity; we need not specify.) Suppose that we are given t copies of the unknown state ρ , and the maximum deviation we wish to tolerate between our estimate $\hat{\rho}$ and the true state ρ is given by $\epsilon > 0$. Now define the minimax risk

$$M^*(t, \epsilon) = \inf_{\langle \hat{\rho}, \Pi_i \rangle} \sup_{\rho \in \Sigma} \Pr[\Delta(\hat{\rho}(Y), \rho) > \epsilon], \quad (17)$$

where the infimum is over all estimation procedures $\langle \hat{\rho}, \Pi_i \rangle$ on t copies with estimator $\hat{\rho}$ and choice of observables given by Π_i . That is, we are considering the ‘best’ protocol to be the one whose worst-case probability of deviation is the least.

The next lemma shows that if there are a large number of states in Σ which are far apart (by at least ϵ) and whose statistics look nearly random for all measurements in $\mathcal{P}$, then the minimax risk is lower-bounded by a function that decreases linearly with the number of copies t . Hence, to achieve a small minimax risk, the number of copies t must be large.

Lemma 1. Assume that there are states $\rho_1, \dots, \rho_s \in \Sigma$ such that

$$\begin{aligned}\forall i \neq j : \Delta(\rho_i, \rho_j) &\geq \epsilon, \\ \forall i, \forall \Pi \in \mathcal{P} : |\text{tr} \rho_i \Pi - \tfrac{1}{s}| &\leq \alpha.\end{aligned}$$

Then the minimax risk as defined in (17) fulfills

$$M^*(t, \epsilon) > 1 - \frac{4\alpha^2 t + 1}{\log s}.$$

Proof. Let X be a random variable taking values uniformly in $[s]$. Let $Y_1, \dots, Y_t$ be random variables, Y_i describing the outcome of the i th measurement carried out on ρ_X . Define

$$\hat{X}(Y) := \arg \min_i \Delta(\hat{\rho}(Y), \rho_i).$$

Then

$$\Pr[\Delta(\hat{\rho}(Y), \rho_i) > \epsilon] \geq \Pr[\hat{X}(Y) \neq X]. \quad (18)$$

Now combine Fano’s inequality

$$H(X|\hat{X}) \leq 1 + \Pr[\hat{X} \neq X] \log s,$$

the data processing inequality

$$I(X; \hat{X}(Y)) \leq I(X; Y),$$

in terms of the mutual information $I(X; Y) := H(X) - H(X|Y)$, and the fact that $H(X) = \log s$ to obtain

$$\begin{aligned} \Pr[\hat{X}(Y) \neq X] &\geq \frac{H(X|\hat{X}) - 1}{\log s} \\ &= \frac{H(X) - I(X; \hat{X}) - 1}{\log s} \\ &= 1 - \frac{I(X; \hat{X}) + 1}{\log s} \\ &\geq 1 - \frac{I(X; Y) + 1}{\log s} \\ &= 1 - \frac{H(Y) - H(Y|X) + 1}{\log s} \\ &\geq 1 - \frac{t - \frac{1}{s} \sum_{i=1}^s H(Y|X=i) + 1}{\log s}. \end{aligned}$$

Let $h(p)$ be the binary entropy and recall the standard estimate

$$h(1/2 \pm \alpha) \geq (1 - 4\alpha^2).$$

Combine that with the chain rule [69, Theorem 2.5.1]

$$\begin{aligned} H(Y|X=i) &= \sum_{j=1}^t H(Y_j|Y_{j-1}, \dots, Y_1, X=i) \\ &\geq t(1 - 4\alpha^2). \end{aligned}$$

The advertised bound follows. $\square$

We now show how to construct a set of states, each with rank r , that satisfy the conditions of lemma 1 in terms of the trace distance and the set of Pauli measurements. The following lemma shows how, given such a set of states, we can enlarge it by one. We will then apply this lemma repeatedly, to get a total of $s < e^{c^{rd}}$ states.

Lemma 2. For any $0 < \epsilon < 1 - \frac{r}{d}$, let $\rho_1, \dots, \rho_s$ be normalized⁸ rank- r projections on $\mathbb{C}^d$, where $s < e^{c(\epsilon)rd}$ and $c(\epsilon)$ is specified below. Then there exists a normalized rank- r projection ρ such that

$$\forall i \in [s] : \frac{1}{2} \|\rho - \rho_i\|_{\text{tr}} \geq \epsilon, \quad (19)$$

$$\forall P_k \neq \mathbb{1} : |\text{tr}[\frac{1}{2}(\mathbb{1} \pm P_k)\rho] - \frac{1}{2}| \leq \alpha. \quad (20)$$

Here, $\alpha^2 = O(\frac{\log d}{rd})$, the P_k are n -qubit Pauli operators, and

$$c(\epsilon) = \frac{\ln(8/\pi)}{2rd} + \frac{1}{32} \left[\left(1 - \frac{r}{d}\right) - \epsilon \right]^2.$$

⁸ Normalized to have trace 1.

Proof. Let ρ_0 be some normalized rank- r projection and choose ρ according to⁹

$$\rho = O\rho_0 O^T$$

for a Haar-random $O \in \text{SO}(d)$. Here we use the special orthogonal group $\text{SO}(d)$ because the analysis becomes marginally simpler than if we use a unitary group.

To check (19), choose $i \in [s]$ and define R_i to be the projector onto the range of ρ_i . Also define the function

$$f : O \mapsto \|\rho_i - O\rho_0 O^T\|_{\text{tr}}.$$

We can bound the magnitude of f using the pinching inequality:

$$\begin{aligned} f(O) &\geq \|\rho_i - R_i O\rho_0 O^T R_i\|_{\text{tr}} + \|R_i^\perp O\rho_0 O^T R_i^\perp\|_{\text{tr}} \geq \text{tr}(\rho_i) - \text{tr}(R_i O\rho_0 O^T R_i) + \text{tr}(R_i^\perp O\rho_0 O^T R_i^\perp) \\ &= 1 + \text{tr}[(R_i^\perp - R_i)O\rho_0 O^T]. \end{aligned}$$

From this we can bound the expectation value of f over the special orthogonal group:

$$\begin{aligned} \mathbb{E}[f(O)] &\geq 1 + \text{tr}[(R_i^\perp - R_i)(\mathbb{E} O\rho O^T)] \\ &= 1 + \frac{1}{d} \text{tr}[(R_i^\perp - R_i)\mathbb{1}] \\ &= 1 + \frac{d-2r}{d} = 2\left(1 - \frac{r}{d}\right). \end{aligned}$$

Next we get an upper bound of $\frac{4}{\sqrt{r}}$ on the Lipschitz constant of f with respect to the Frobenius norm:

$$\begin{aligned} |f(O + \Delta) - f(O)| &\leq \|(O + \Delta)\rho_0(O + \Delta)^T - O\rho_0 O^T\|_{\text{tr}} \\ &\leq \|O\rho_0 \Delta^T\|_{\text{tr}} + \|\Delta\rho_0 O^T\|_{\text{tr}} + \|\Delta\rho_0 \Delta^T\|_{\text{tr}} \\ &= 2\|\Delta\rho_0\|_{\text{tr}} + \text{tr}(\Delta\rho_0 \Delta^T) \\ &\leq 2\sqrt{r}\|\Delta\rho_0\|_{\text{F}} + \text{tr}(\rho_0 \Delta \Delta^T) \\ &\leq 2\sqrt{r}\|\Delta\|_{\text{F}}\|\rho_0\| + \frac{1}{r}\|\Delta\|\|\Delta^T\|_{\text{tr}} \\ &\leq \frac{2}{\sqrt{r}}\|\Delta\|_{\text{F}} + \frac{2}{r}\sqrt{r}\|\Delta\|_{\text{F}} \leq \frac{4}{\sqrt{r}}\|\Delta\|_{\text{F}}, \end{aligned}$$

where we use $\|\Delta\| \leq 2$ in the last line, which follows from the triangle inequality and the fact that any Δ can be written as a difference $\Delta = O' - O$ for $O' \in \text{SO}(d)$.

From these ingredients, we can invoke Lévy's lemma on the special orthogonal group [70, Theorem 6.5.1] to obtain that, for all $t > 0$,

$$\Pr[\|\rho_i - \rho\|_{\text{tr}} < 2(1 - r/d) - t] \leq e^{c_1} \exp\left(-\frac{c_2 t^2 r d}{16}\right),$$

where the constants are given by $c_1 = \ln(\sqrt{\pi/8})$ and $c_2 = 1/8$. Now choose $t = 2(1 - r/d) - 2\epsilon$ and apply the union bound to obtain

$$\begin{aligned} \Pr[(19) \text{ does not hold}] &< e^{c(\epsilon)rd} \Pr[\|\rho_i - \rho\|_{\text{tr}} < 2(1 - r/d) - t] \\ &\leq \exp\left(rd\left[c(\epsilon) - \frac{c_2 t^2}{16}\right] + c_1\right) = 1. \end{aligned}$$

⁹ For the duration of this proof, the letter O denotes an element of $\text{SO}(d)$ instead of the asymptotic big- O notation.

The upper bound on ϵ follows from the requirement that $t > 0$. This shows that ρ indeed satisfies equation (19).

Now we move on to equation (20). For any non-identity Pauli matrix P_k , define a function

$$f : O \mapsto \text{tr}(P_k O \rho_0 O^T).$$

Clearly, we have $\mathbb{E}[f(O)] = 0$. We again wish to bound the rate of change so that we can use Lévy's lemma, so we compute

$$(\text{d}_O f)(\Delta) = \text{tr}(\rho_0 O^T P_k \Delta) + \text{tr}(P_k O \rho_0 \Delta^T) = \text{tr}[(\rho_0 O^T P_k + \rho_0^T O^T P_k^T) \Delta],$$

which implies that

$$\|\nabla f(O)\|_2 = \|\rho_0 O^T P_k + \rho_0^T O^T P_k^T\|_F \leq \frac{2}{\sqrt{r}}.$$

Lévy's lemma then gives for all $t > 0$

$$\Pr[|\text{tr} P_k \rho| > t] \leq e^{c_1} \exp\left(-\frac{c_2 t^2 r d}{4}\right).$$

Choosing $t = 2\alpha$, and $\alpha^2 = 4 \ln(d^4 \pi / 8) / (r d)$, then the union bound gives us

$$\Pr[(20) \text{ does not hold}] < d^2 \Pr[|\text{tr} P_k \rho| > 2\alpha] \leq \exp\left(2 \ln d - \frac{c_2 (2\alpha)^2 r d}{4} + c_1\right) = 1,$$

from which the lemma follows. $\square$

We remark that a version of lemma 2 continues to hold even if we can adaptively choose from as many as $2^{O(n)}$ additional measurements which are globally unitarily equivalent to Pauli measurements.

We now combine the two previous lemmas, to get a precise version of theorem 3. (To summarize: we use lemma 2 repeatedly to construct a large set of states that are difficult to distinguish using Pauli measurements; then we use lemma 1 to lower-bound the minimax risk for estimating an unknown state using t copies; hence, if an estimation procedure has a small minimax risk, t must be large.)

Theorem 4 (Precise version of theorem 3). Fix $\epsilon \in (0, 1 - \frac{r}{d})$ and $\delta \in [0, 1)$. Then, in order for an estimation procedure to satisfy $M^*(t, \epsilon) \leq \delta$, the number of copies t must grow like

$$t = \Omega\left(\frac{r^2 d^2}{\log d}\right),$$

where the implicit constant depends on δ and ϵ .

4. Certifying the state estimate

Here we sketch how the technique of DFE, introduced in [33, 34] for pure states, can be used to estimate the fidelity between a low-rank estimate $\hat{\rho}$ and the true state ρ . The only assumption we make is that $\hat{\rho}$ is a positive semidefinite matrix with $\text{Tr}(\hat{\rho}) \leq 1$ and $\text{rank}(\hat{\rho}) = r$. No assumption at all is needed on ρ . In fact, we also do not assume that we obtained the estimate $\hat{\rho}$ from any of the estimators which were discussed previously. Our certification procedure works regardless of how one obtains $\hat{\rho}$, and so it even applies to the situation where $\hat{\rho}$ was chosen from a subset of variational ansatz states, as in [61].

Recall that the main idea of DFE is to take a known pure state $|\psi\rangle\langle\psi|$ and from it define a probability distribution $\text{Pr}(i)$ such that by estimating the Pauli expectation values $\text{Tr}(\rho P_i)$ and suitably averaging it over $\text{Pr}(i)$ we can learn an estimate of $\langle\psi|\rho|\psi\rangle$. In fact, one does not need to do a full average; simply sampling from the distribution a few times is sufficient to produce a good estimate.

For non-Hermitian rank-1 matrices $|\phi_j\rangle\langle\phi_k|$ instead of pure states, we need a very slight modification of DFE. Following the notation in [33], we simply redefine the probability distribution as $\text{Pr}(i) = |\langle\phi_j|P_i|\phi_k\rangle|^2/d$ and the random variable $X = \text{Tr}(P_i\rho)/\langle\phi_k|P_i|\phi_j\rangle$. It is easy to check that $\mathbb{E}(X) = \langle\phi_j|\rho|\phi_k\rangle$ and the variance of X is at most one. Then all of the conclusions in [33, 34] hold for estimating the quantity $\langle\phi_j|\rho|\phi_k\rangle$. In particular, we can obtain an estimate to within an error $\pm\epsilon$ with probability $1 - \delta$ by using only single-copy Pauli measurements and $O(\frac{1}{\epsilon^2\delta} + \frac{d \log(1/\delta)}{\epsilon^2})$ copies of ρ .

Our next result shows that by obtaining several such estimates using DFE, we can also infer an estimate of the mixed state fidelity between an unknown state ρ and a low-rank estimate $\hat{\rho}$, given by

$$F(\rho, \hat{\rho}) = [\text{Tr}\sqrt{G}]^2, \quad (21)$$

where for brevity we define $G = \sqrt{\hat{\rho}}\rho\sqrt{\hat{\rho}}$. (Note that some authors use the square root of this quantity as the fidelity. Our definition matches [33].) The asymptotic cost in sample complexity is far less than the cost of initially obtaining the estimate $\hat{\rho}$ when r is sufficiently small compared to d .

Theorem 5. *Given a state estimate $\hat{\rho}$ with $\text{rank}(\hat{\rho}) = r$, the number of copies t of the state ρ required for estimating $F(\rho, \hat{\rho})$ to within $\pm\epsilon$ with probability $1 - \delta$ using single-copy Pauli measurements satisfies*

$$t = O\left(\frac{r^5}{\epsilon^4} [d \log(r^2/\delta) + r^2/\delta]\right). \quad (22)$$

Proof. The result uses the DFE protocol of [33, 34], modified as mentioned above, where the states $|\phi_j\rangle$ are the eigenstates of $\hat{\rho}$.

Expand $\hat{\rho}$ in its eigenbasis as $\hat{\rho} = \sum_{j=1}^r \lambda_j |\phi_j\rangle\langle\phi_j|$. Then we have

$$G = \sum_{j,k=1}^r \sqrt{\lambda_j \lambda_k} \langle\phi_j|\rho|\phi_k\rangle |\phi_j\rangle\langle\phi_k|. \quad (23)$$

For all $1 \leq j \leq k \leq r$, we use DFE to obtain an estimate $\hat{g}_{jk}$ of the matrix element $\langle\phi_j|\rho|\phi_k\rangle$, up to some additive error ϵ_{jk} that is bounded by a constant, $|\epsilon_{jk}| \leq \epsilon_0$.

If each estimate is accurate with probability $1 - 2\delta/(r^2 + r)$, then by the union bound the probability that they are all accurate is at least $1 - \delta$. The total number of copies t required for this is

$$t = O\left(\frac{r^2}{\epsilon_0^2} [d \log(r^2/\delta) + r^2/\delta]\right). \quad (24)$$

Let $\hat{g}_{jk} = \hat{g}_{kj}^*$, and let

$$\hat{G} = \sum_{j,k=1}^r \sqrt{\lambda_j \lambda_k} \hat{g}_{jk} |\phi_j\rangle\langle\phi_k| \quad (25)$$

be our estimate for G . Finally, let $\hat{G}^+ = [\hat{G}]_+$ be the positive part of the Hermitian matrix $\hat{G}$, and let $\hat{F} = [\text{Tr}\sqrt{\hat{G}^+}]^2$ be our estimate of the fidelity $F(\rho, \hat{\rho})$. Note that we may assume that $\hat{F} \leq 1$, since if it were larger, we can only improve our estimate by just truncating it back to 1.

We now bound the error of this fidelity estimate. We can write $\hat{G}$ as a perturbation of G , $\hat{G} = G + E$, where

$$E = \sum_{j,k=1}^r \sqrt{\lambda_j \lambda_k} \epsilon_{jk} |\phi_j\rangle\langle\phi_k|. \quad (26)$$

First note that the Frobenius norm of this perturbation is small,

$$\|E\|_F = \left(\sum_{j,k} \lambda_j \lambda_k |\epsilon_{jk}|^2 \right)^{1/2} \leq \epsilon_0 \left| \sum_j \lambda_j \right| = \epsilon_0. \quad (27)$$

(If $\hat{\rho}$ is subnormalized, then this last equality becomes an inequality.)

Next observe that

$$|F - \hat{F}| = \left| \text{Tr}(\sqrt{G} + \sqrt{\hat{G}^+}) \text{Tr}(\sqrt{G} - \sqrt{\hat{G}^+}) \right| \leq 2\Delta,$$

where we define

$$\Delta = |\text{Tr}(\sqrt{G} - \sqrt{\hat{G}^+})|. \quad (28)$$

Using the reverse triangle inequality, we can bound Δ in terms of the trace norm

$$\Delta \leq \left\| \sqrt{G} - \sqrt{\hat{G}^+} \right\|_{\text{tr}}. \quad (29)$$

Using [71, Theorem X.1.3] in the first step, we find that

$$\left\| \sqrt{G} - \sqrt{\hat{G}^+} \right\|_{\text{tr}} \leq \left\| \sqrt{|G - \hat{G}^+|} \right\|_{\text{tr}} \leq \sqrt{r} \sqrt{\|G - \hat{G}^+\|_{\text{tr}}},$$

where the second bound follows from the Cauchy–Schwarz inequality on the vector of eigenvalues of $|G - \hat{G}^+|$. Using the Jordan decomposition of a Hermitian matrix into a difference of positive matrices $X = [X]_+ - [X]_-$, we can rewrite $G - \hat{G}^+$ as

$$G - \hat{G}^+ = G - [G + E]_+ = -E - [G + E]_-. \quad (30)$$

Then by the triangle inequality and positivity of G ,

$$\|G - \hat{G}^+\|_{\text{tr}} \leq \|E\|_{\text{tr}} + \|[G + E]_-\|_{\text{tr}} \leq 2\|E\|_{\text{tr}}. \quad (31)$$

Using the standard estimate $\|E\|_{\text{tr}} \leq \sqrt{r}\|E\|_F$, we find that

$$\Delta \leq r^{3/4} \sqrt{2\epsilon_0}. \quad (32)$$

This gives our desired error bound, albeit in terms of ϵ_0 instead of the final quantity ε . To get our total error to vanish, we take $\varepsilon = 2r^{3/4} \sqrt{2\epsilon_0}$, which gives us the final scaling in the sample complexity. $\square$

As a final remark, we note that by computing Δ for the special case $\hat{\rho} = \rho = \mathbb{1}/r$, $E = \epsilon_0 \mathbb{1}/r$, we find that

$$\Delta = \sqrt{1 + r\epsilon_0} - 1, \quad (33)$$

so the error bound for this protocol is tight with respect to the scaling in ϵ_0 (and hence ϵ). However, we cannot rule out that there are other protocols that achieve a better scaling. Also, it seems that the upper bound for the current protocol with respect to r could potentially be improved by a factor of r with a more careful analysis.

5. Numerical simulations

We numerically simulated the reconstruction of a quantum state given Pauli measurements and the estimators from equations (3) and (4). Here we compare these estimates to those obtained from a standard maximum likelihood estimate (MLE) as a benchmark.

As mentioned previously, all of our estimators have the advantage that they are convex functions with convex constraints, so the powerful methods of convex programming [72] guarantee that the exact solution can be found efficiently and certified. We take full advantage of this certifiability and do simulations which can be certified by interior-point algorithms. This way we can separate out the performance of the estimators themselves from the (sometimes heuristic) methods used to compute them on very large scale instances.

For our estimators, we will explicitly enforce the positivity condition. That is, we will simulate the following slight modifications of equations (3) and (4) given by

$$\hat{\rho}'_{\text{DS}} = \arg \min_{X \succeq 0} \text{Tr} X \quad \text{s.t.} \quad \|\mathcal{A}^*(\mathcal{A}(X) - y)\| \leq \lambda \quad (34)$$

and

$$\hat{\rho}'_{\text{Lasso}} = \arg \min_{X \succeq 0} \frac{1}{2} \|\mathcal{A}(X) - y\|_2^2 + \mu \text{Tr} X. \quad (35)$$

Moreover, whenever the trace of the resulting estimate is less than 1, we renormalize the state, $\hat{\rho} \leftarrow \hat{\rho} / \text{Tr}(\hat{\rho})$.

5.1. Setting the estimator parameters λ and μ

From theorem 1 we know roughly how we should choose the free parameters λ and μ , modulo the unknown constants C_i , for which we will have to make an empirical guess. We guess that the log factor is an artifact of the analysis and that the state is nearly pure. Then we could choose $\lambda \sim \mu \sim d/\sqrt{t}$. For the matrix Dantzig selector of equation (34), we specifically chose $\lambda = 3d/\sqrt{t}$ for our simulations. However, this did not lead to very good performance for the matrix Lasso of equation (35) when m was large. We chose instead $\mu = 4m/\sqrt{t}$, which agrees with λ when $m \sim d$, as it can be for nearly pure states.

We stress that very little effort was made to optimize these choices for the parameters. We simply picked a few integer values for the constants in front (namely 2, $\dots$, 5), and chose the constant that worked best upon a casual inspection of a few data points. We leave open the problem of determining what the optimal choices are for λ and μ .

5.2. Time needed to switch measurement settings

Real measurements take time. In a laboratory setting, time is a precious commodity for a possibly non-obvious reason. An underlying assumption in the way we typically formulate tomography is that the unknown state ρ can be prepared identically many times. However, the parameters governing the experiment often drift over a certain timescale, and this means that

beyond this characteristic time, it is no longer appropriate to describe the measured ensemble by the symmetric product state $\rho^{\otimes t}$.

To account for this characteristic timescale, we introduce the following simple model. We assume that the entire experiment takes place in a fixed amount of time T . In a real experiment, this is the timescale beyond which we cannot confidently prepare the same state ρ , or it is the amount of time it takes for the student in the laboratory to graduate, or for the grant funding to run out, whichever comes first. Within this time limit, we can either change the measurement *settings*, or we can take more *samples*. We assume that there is a unit cost to taking a single measurement, but that switching to a different measurement configuration takes an amount of time c .

At the one extreme, the switching cost c to change measurement settings could be quite large, so that it is barely possible to make enough independent measurements that tomography is possible within the allotted time T . In this case, we expect that compressed tomography will outperform standard methods, since these methods have not been optimized for the case of incomplete data. At the other extreme, the switching cost c could be set to 0 (or some other small value), in which case we are only accounting for the cost of sampling. Here it is not clear which method is better, for the following reason. Although the standard methods are able to take a sufficient number of data in this case, each observable will attain comparatively little precision relative to the fewer observables measured with higher precision in the case of compressed tomography for the same fixed amount of time T . One of our goals is to investigate if there are any trade-offs in this simple model.

For the relative cost c between switching measurement settings and taking one sample, we use $c = 20$, a number inspired by the current ion trap experiments done by the Blatt group in Innsbruck. However, we note that the conclusions do not seem to be very sensitive to this choice, as long as we do not do something outrageous like $c > t$, which would preclude measuring more than even one observable.

5.3. Other simulation parameters

We consider the following ensembles of quantum states. We initially choose $n = 5$ qubit states from the invariant Haar measure on $\mathbb{C}^{2^n}$. Then we add noise to the state by applying independent and identical depolarizing channels to each of the n qubits. Recall that the depolarizing channel with strength γ acts on a single qubit as

$$\mathcal{D}_\gamma(\rho) = \gamma \frac{1}{2} + (1 - \gamma)\rho; \quad (36)$$

that is, with probability γ the qubit is replaced by a maximally mixed state and otherwise it is left alone. Our simulations assume very weak decoherence, and we use $\gamma = 0.01$.

We use two measures to quantify the quality of a state reconstruction $\hat{\rho}$ relative to the underlying true state ρ , namely the (squared) fidelity $\|\sqrt{\hat{\rho}}\sqrt{\rho}\|_{\text{tr}}^2$ and the trace distance $\frac{1}{2}\|\rho - \hat{\rho}\|_{\text{tr}}$.

We used the interior-point solver SeDuMi [26] to do the optimizations in equations (34) and (35). Although these algorithms are not suitable for large numbers of qubits, they produce very accurate solutions, which allows for a more reliable comparison between the different estimators. The data and the code used to generate the data for these plots can be found with the source code for the arXiv version of this paper. For larger numbers of qubits, one may instead

use specialized algorithms such as singular value thresholding (SVT) or fixed-point continuation (FPCA) to solve these convex programs [27–29]; however, the quality of the solutions depends somewhat on the algorithm, and it is an open problem to explore this in more detail.

For the MLE, we used the iterative algorithm of [73], which has previously been used in e.g. [74]. This method converged on every example we tried, so we did not have to use the more sophisticated ‘diluted’ method of [75] that guarantees convergence.

For the purposes of our simulations, we sampled our Pauli operators *without* replacement.

5.4. Results and analysis

The data are depicted in figure 1. The three subfigures (a)–(c) use three different values for the total sampling time T in increasing order. We have plotted both the fidelity and the trace distance (inset) between the recovered solution and the true state.

Several features are immediately apparent. First, we see that both of the compressed sensing estimators consistently outperform MLE along a wide range of values of m , the number of Paulis sampled, even when we choose the optimal value of m just for the MLE. Once m is sufficiently large (but still $m \ll d^2$) *all* of the estimators converge to a reasonably high-fidelity estimate. Thus, it is not just the compressed sensing estimators which are capable of reconstructing nearly low-rank states with few measurement settings, but rather this seems to be a generic feature of many estimators. However, the compressed sensing estimators seem to be particularly well suited for the task.

When the total amount of time T is *very* small (figure 1(a)), then there is a large advantage of using compressed sensing. In this regime, if we insist on making an informationally complete measurement, there is barely enough time to make one measurement per setting, so the measurement results are dominated by statistical fluctuations, i.e. the signal-to-noise ratio is poor. Compressed sensing techniques—measuring an incomplete set of observables, and regularizing the solution to favor low-rank states—lead to much better results. However, note that in this situation, it is very difficult for *any* method, including compressed sensing, to achieve extremely high-fidelity state reconstructions (with fidelity greater than 95%, say).

As the total amount of time available increases, all of the estimators of course converge to higher fidelity estimates. Interestingly, for the compressed sensing estimators, there seems to be *no trade-off whatsoever* between the number of measurement settings and the fidelity, once T and m are sufficiently large. That is, the curve is completely flat as a function of m above a certain cutoff. This is most notable for the matrix Lasso, which consistently performs at least as well as the matrix Dantzig selector, and often better. These observations are consistent with the bounds proven earlier.

The flat curve as a function of m is especially interesting, because it suggests that there is no real drawback to using small values of m . Smaller values of m are attractive from the computational point of view because the runtime of each reconstruction algorithm scales with m . This makes a strong case for adopting these estimators in the future, and at the very least more numerical studies are needed to investigate how well these estimators perform more generally.

We draw the following conclusion from these simulations. The best performance for a fixed value of T is given by the matrix Lasso estimator of equations (4) and (35) in the regime where m is nearly as small as possible. Here the fidelity is larger than the other estimators (if only by a small or negligible amount when T is large), but using a small value for m means smaller memory and processing time requirements when doing the state reconstruction. MLE

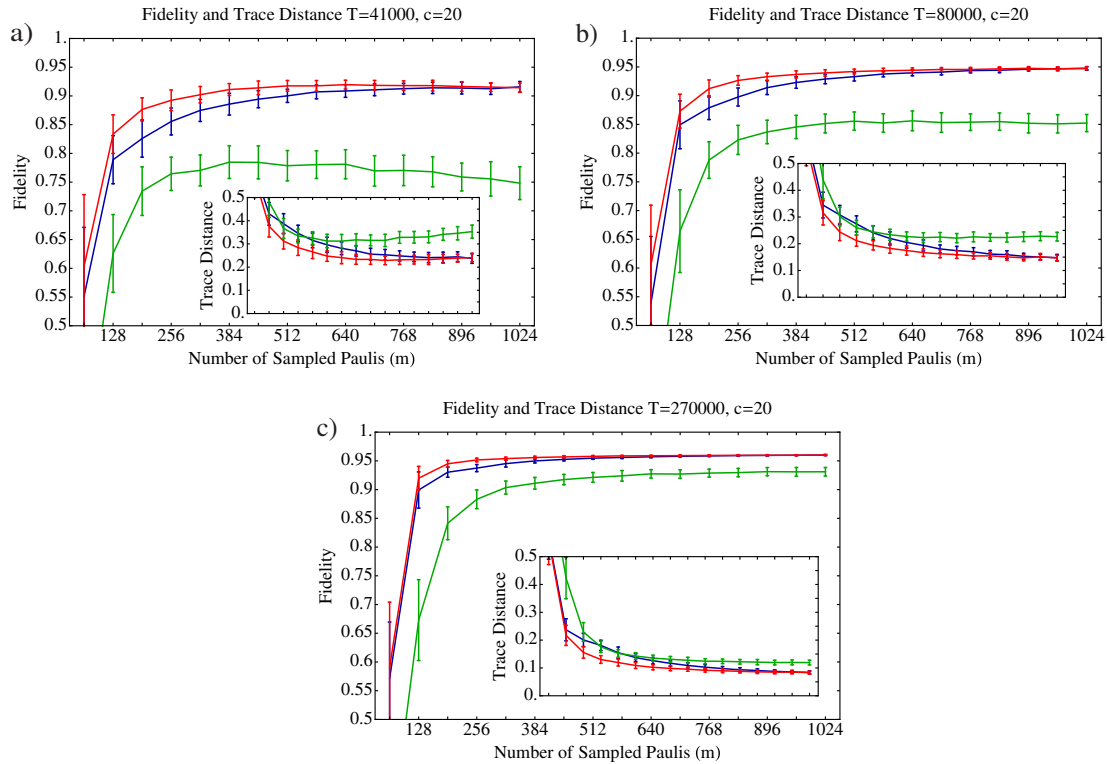


Figure 1. Fidelity and trace distance (inset) as a function of m , the number of sampled Paulis. Plots (a), (b) and (c) are for a total measurement time of $T = 41\text{k}$, 80k and 270k , respectively, in units where the cost to measure a single copy is one unit of time, and the cost to switch measurement settings is $c = 20$ units. The three estimators plotted are the matrix Dantzig selector (equation (34), blue), the matrix Lasso (equation (35), red) and a standard MLE (green). Each bar shows the mean and standard deviation from 120 Haar-random five-qubit states with 1% local depolarizing noise. Our estimators consistently outperform MLE, even after optimizing the MLE over m . (We do not know what precisely causes the degradation of the performance of MLE in figure 1a, but it is plausible to speculate that the cost term c is responsible for this effect). See the main text for more details.

consistently underperforms the compressed sensing estimators, but still seems to ‘see’ the low-rank nature of the underlying state and converges to a reasonable estimate even when m is small.

6. Process tomography

Compressed sensing techniques can also be applied to quantum *process* tomography. Here, our method has an advantage when the unknown quantum process has a small Kraus rank, meaning that it can be expressed using only a few Kraus operators. This occurs, for example, when the unknown process consists of unitary evolution (acting globally on the entire system) combined

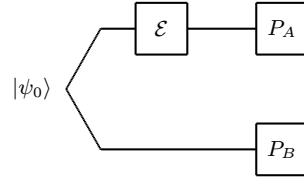


Figure 2. Compressed quantum process tomography using an ancilla. The quantum circuit represents a single measurement setting, where one measures the observable P_A on the system and P_B on the ancilla.

with local noise (acting on each qubit individually, or more generally, acting on small subsets of the qubits).

Consider a system of n qubits, with Hilbert space dimension $d = 2^n$. Let $\mathcal{E}$ be a completely positive trace-preserving (CP) map from $\mathbb{C}^{d \times d}$ to $\mathbb{C}^{d \times d}$, and suppose that $\mathcal{E}$ has Kraus rank r . Using compressed sensing, we will show how to characterize $\mathcal{E}$ using $m = O(rd^2 \log d)$ settings. (For comparison, standard process tomography requires d^4 settings, since $\mathcal{E}$ contains d^4 independent parameters in general.) Furthermore, our compressed sensing method requires only the ability to prepare product eigenstates of Pauli operators and make Pauli measurements, and it does not require any ancilla qubits.

We remark that, except for the ancilla-assisted method, just the notion of ‘measurement settings’ for process tomography does not capture all of the complexity because of the need to have an *input* to the channel. Here we define one ‘input setting’ to be a basis of states from which the channel input should be sampled uniformly. Then the total number of settings m is the sum of the number of measurement settings (Paulis, in our case) and input settings. This definition justifies the claim that the number of settings for our compressed process tomography scheme is $m = O(rd^2 \log d)$.

The analysis here focuses entirely on the number of settings m . We forgo a detailed analysis of t , the sample complexity, and instead leave this open for future work.

Note that there is a related set of techniques for estimating an unknown process that is *element-wise sparse* with respect to some known, experimentally accessible basis [19]. These techniques are not directly comparable to ours, since they assume a greater amount of prior knowledge about the process, and they use measurements that depend on this prior knowledge. We will discuss this in more detail at the conclusion of this section.

6.1. Our method

First, recall the Jamiołkowski isomorphism [76]: the process $\mathcal{E}$ is completely and uniquely characterized by the state

$$\rho_{\mathcal{E}} = (\mathcal{E} \otimes \mathcal{I})(|\psi_0\rangle\langle\psi_0|),$$

where $|\psi_0\rangle = \frac{1}{\sqrt{d}} \sum_{j=1}^d |j\rangle \otimes |j\rangle$. Note that when $\mathcal{E}$ has Kraus rank r , the state $\rho_{\mathcal{E}}$ has rank r . This immediately gives us a way to do compressed quantum process tomography: first prepare the Jamiołkowski state $\rho_{\mathcal{E}}$ (by adjoining an ancilla, preparing the maximally entangled state $|\psi_0\rangle$, and applying $\mathcal{E}$); then do compressed quantum state tomography on $\rho_{\mathcal{E}}$; see figure 2.

We now show a more direct implementation of compressed quantum process tomography that is equivalent to the above procedure, but does not require an ancilla. Observe that in the

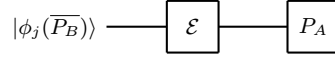


Figure 3. Compressed quantum process tomography, implemented directly without an ancilla. Here one prepares a random eigenstate of $\overline{P_B}$, applies the process $\mathcal{E}$ and measures the observable P_A on the output.

above procedure, we need to estimate expectation values of the form

$$\text{Tr}((P_A \otimes P_B)\rho_{\mathcal{E}}) = \text{Tr}((P_A \otimes P_B)(\mathcal{E} \otimes \mathcal{I})(|\psi_0\rangle\langle\psi_0|)),$$

where P_A and P_B are Pauli matrices. By using the Kraus decomposition, it is straightforward to derive the equivalent expression

$$\text{Tr}((P_A \otimes P_B)\rho_{\mathcal{E}}) = \frac{1}{d} \text{Tr}(P_A \mathcal{E}(\overline{P_B})), \quad (37)$$

where the bar denotes complex conjugation in the standard basis.

We now show how to estimate the expectation value (37). Let λ_j and $|\phi_j\rangle$ denote the eigenvalues and eigenvectors of $\overline{P_B}$. Then we have

$$\text{Tr}((P_A \otimes P_B)\rho_{\mathcal{E}}) = \frac{1}{d} \sum_{j=1}^d \lambda_j \text{Tr}(P_A \mathcal{E}(|\phi_j\rangle\langle\phi_j|)).$$

To estimate this quantity, we repeat the following experiment many times, and average the results: choose $j \in [d]$ uniformly at random, prepare the state $|\phi_j\rangle$, apply the process $\mathcal{E}$, measure the observable P_A , and multiply the measurement result by λ_j . (See figure 3.) In this way, we learn the expectation values of the Jamiołkowski state $\rho_{\mathcal{E}}$ without using an ancilla. We then use compressed quantum state tomography to learn $\rho_{\mathcal{E}}$, and from this we recover $\mathcal{E}$. Note that since the approach described here can also be applied to the case that $\mathcal{E}$ is not trace-preserving, it applies to detector tomography as well.

6.2. Related work

Our method is somewhat different from the method described in [19]. Essentially, the difference is that our method works for any quantum process with a small Kraus rank, whereas the method of Shabani *et al* works for a quantum process that is element-wise sparse in a known basis (provided this basis is experimentally accessible in a certain sense). The main advantage of the Shabani *et al* method is that it can be much faster: for a quantum process $\mathcal{E}$ that is s -sparse (i.e. has s non-zero matrix elements), it requires only $O(s \log d)$ settings. The main disadvantage is that it requires more prior knowledge about $\mathcal{E}$ and is more difficult to apply. While it has been demonstrated in a number of scenarios, there does not seem to be a general recipe for designing measurements that are both experimentally feasible and effective in the sparse basis of $\mathcal{E}$.

To clarify these issues, we now briefly review the Shabani *et al* method. We assume that we know a basis $\Gamma = \{\Gamma_{\alpha} | \alpha \in [d^2]\}$ in which the process $\mathcal{E}$ is s -sparse. For example, when $\mathcal{E}$ is close to some unitary evolution U , one can construct Γ using the singular-value decomposition (SVD) basis of U . This guarantees that, if $\mathcal{E}$ contains no noise, it will be perfectly sparse in the basis Γ . However, in practice, $\mathcal{E}$ will contain noise, which need not be sparse in the basis Γ ; any non-sparse components will not in general be estimated accurately. The success of the

Shabani *et al* method therefore rests on the assumption that the *noise* is also sparse in the basis Γ . Although this assumption has been verified in a few specific scenarios, it seems less clear why it should hold in general. By comparison, our method simply assumes that the noise is described by a process matrix that is low rank; this can be rigorously justified for any noise process that involves only local interactions or few-body processes.

The other complication with the Shabani *et al* method concerns the design of the state preparations and measurements. On the one hand, these must satisfy the RIP condition for s -sparse matrices over the basis Γ ; on the other hand, they must be easy to implement experimentally. This has been demonstrated in some cases, by using states and measurements that are ‘random’ enough to show concentration of measure behaviors, but also have tensor product structure. However, these constructions are not guaranteed to work in general for an arbitrary basis Γ .

We leave open the problem of doing a comparative study between these and other methods [77, 78], akin to [79].

7. Conclusion

In this work, we have introduced two new estimators for compressed tomography: the matrix Dantzig selector and the matrix Lasso (equations (3) and (4)). We have proved that the sample complexity for obtaining an estimate that is accurate to within ϵ in the trace distance scales as $O(\frac{r^2 d^2}{\epsilon^2} \log d)$ for rank- r states and that for higher rank states, the additional error is proportional to the truncation error. This error scaling is optimal up to constant multiplicative factors, and requires measuring only $O(rd \text{ poly } \log d)$ Pauli expectation values, a fact we proved using the RIP [30]. We also proved that our sample complexity upper bound is within $\text{poly } \log d$ of the sample complexity of the optimal minimax estimator, where the risk function is given by a trace norm confidence interval. We showed how a modification of DFE can be used to unconditionally certify the estimate using a number of measurements which is asymptotically negligible compared to obtaining the original estimate. We numerically simulated our estimators and found that they outperform MLE, giving higher fidelity estimates from the same amount of data. Finally, we generalized our method to quantum process tomography using only Pauli measurements and preparation of product eigenstates of Pauli operators.

There are many interesting directions for future work. One major open problem is to switch the focus from two-outcome Pauli measurements to alternative measurements which are still experimentally feasible. For example, measurements in a local basis have 2^n outcomes and are not directly analyzable using our techniques. It would be very interesting to give an analysis of our estimators from the perspective of such local basis measurements. One difficulty, however, is that something like the RIP is not likely to hold in this case, so we will need additional techniques.

Another direction is to find better methods for deriving error bars or confidence regions for compressed tomography. In principle, one can use the error bounds in this paper to estimate confidence regions (based on mild assumptions about the experimental setup), and then use DFE to check (unconditionally, i.e. without making any assumptions) whether the true state lies within that confidence region. However, while the error bounds in this paper have roughly the right asymptotic scaling (with the rank r and dimension d), in practice one may want to know finer properties of the confidence region, such as its shape [63, 64].

A related direction is to find methods for rank-based model selection (or, more simply, estimating the rank of the unknown state) [80]. In principle, one can always ‘guess and check’ the rank, using DFE; but a more sophisticated approach may perform better in practice.

Some other open problems are to tighten the gap that remains between the sample complexity upper and lower bounds, and to try to prove optimality with respect to alternative criteria, in addition to minimax risk. For example, it would be interesting to find a useful notion of average case optimality.

On the numerical side, some of the open questions are the following. First, it would be very interesting to carry out a more detailed numerical study of the performance of our estimators. While they have clearly outperformed MLE in the simulations reported here, there is no question that this is a narrow range of parameters on which we have tested these estimators. It would be interesting to do additional comparative studies between these and other estimators to see how robust these performance enhancements are. It would also be very interesting to study fast first-order solvers such as [27–29] which work particularly well for approximately low-rank matrices and which could compute estimates on a large number of qubits (ten or more).

The success of the estimators we studied depends on being able to find good values for the parameters λ and μ . While we have used simple heuristics to pick particular values, a detailed study of the optimal values for these parameters could only improve the quality of our estimators. Moreover, MLE seems to enjoy the same ‘plateau’ phenomenon, where the quality of the estimate is insensitive to m above a certain cutoff. This leads us to speculate that this is a generic phenomenon among many estimators, and that perhaps there are even better choices for estimators than the ones we benchmark here.

Acknowledgments

We thank R Blume-Kohout, M Kleinmann and T Monz for discussions, B Brown for some preliminary numerical investigations and S Glancy for helpful comments on a draft of the paper. We additionally thank B Englert for hosting the workshop on quantum tomography in Singapore where some of this work was completed. STF was supported by the IARPA MQCO program. DG was supported by the Excellence Initiative of the German Federal and State Governments (grant ZUK 43). Contributions to this work by NIST, an agency of the US government, are not subject to copyright laws. Although this paper refers to specific computer software, NIST does not endorse commercial products. JE was supported by EURYL, the EU (Q-Essence) and BMBF (QuOReP).

References

- [1] Nature Insight Supplement 2008 Quantum coherence *Nature* **453** 1003
- [2] Smithey D T, Beck M, Raymer M G and Faridani A 1993 Measurement of the Wigner distribution and the density matrix of a light mode using optical homodyne tomography: application to squeezed states and the vacuum *Phys. Rev. Lett.* **70** 1244
- [3] Altepeter J B, Branning D, Jeffrey E, Wei T C, Kwiat P G, Thew R T, O’Brien J L, Nielsen M A and White A G 2003 Ancilla-assisted quantum process tomography *Phys. Rev. Lett.* **90** 193601
- [4] O’Brien J L, Pryde G J, White A G, Ralph T C and Branning D 2003 Demonstration of an all-optical quantum controlled-NOT gate *Nature* **426** 264

- [5] O'Brien J L, Pryde G J, Gilchrist A, James D F V, Langford N K, Ralph T C and White A G 2004 Quantum process tomography of a controlled-NOT gate *Phys. Rev. Lett.* **93** 080502
- [6] Roos C F, Lancaster G P T, Riebe M, Häffner H, Hänsel W, Gulde S, Becher C, Eschner J, Schmidt-Kaler F and Blatt R 2004 Bell states of atoms with ultralong lifetimes and their tomographic state analysis *Phys. Rev. Lett.* **92** 220402
- [7] Resch K J, Walther P and Zeilinger A 2005 Full characterization of a three-photon Greenberger–Horne–Zeilinger state using quantum state tomography *Phys. Rev. Lett.* **94** 070402
- [8] Häffner H *et al* 2005 Scalable multiparticle entanglement of trapped ions *Nature* **438** 643
- [9] Myrskog S H, Fox J K, Mitchell M W and Steinberg A M 2005 Quantum process tomography on vibrational states of atoms in an optical lattice *Phys. Rev. A* **72** 013615
- [10] Smith G A, Silberfarb A, Deutsch I H and Jessen P S 2006 Efficient quantum-state estimation by continuous weak measurement and dynamical control *Phys. Rev. Lett.* **97** 180403
- [11] Riebe M, Kim K, Schindler P, Monz T, Schmidt P O, Körber T K, Hänsel W, Häffner H, Roos C F and Blatt R 2006 Process tomography of ion trap quantum gates *Phys. Rev. Lett.* **97** 220407
- [12] Filipp S *et al* 2009 Two-qubit state tomography using a joint dispersive readout *Phys. Rev. Lett.* **102** 200402
- [13] Medendorp Z E D, Torres-Ruiz F A, Shalm L K, Tabia G N M, Fuchs C A and Steinberg A M 2011 Experimental characterization of qutrits using symmetric informationally complete positive operator-valued measurements *Phys. Rev. A* **83** 051801
- [14] Barreiro J T, Muller M, Schindler P, Nigg D, Monz T, Chwalla M, Hennrich M, Roos C F, Zoller P and Blatt R 2011 An open-system quantum simulator with trapped ions *Nature* **470** 486
- [15] Fedorov A, Steffen L, Baur M, da Silva M P and Wallraff A 2012 Implementation of a Toffoli gate with superconducting circuits *Nature* **481** 170
- [16] Liu W-T, Zhang T, Liu J-Y, Chen P-X and Yuan J-M 2012 Experimental quantum state tomography via compressed sampling *Phys. Rev. Lett.* **108** 170403
- [17] Gross D, Liu Y-K, Flammia S T, Becker S and Eisert J 2010 Quantum state tomography via compressed sensing *Phys. Rev. Lett.* **105** 150401
- [18] Gross D 2011 Recovering low-rank matrices from few coefficients in any basis *IEEE Trans. Inf. Theory* **57** 1548
- [19] Shabani A, Kosut R L, Mohseni M, Rabitz H, Broome M A, Almeida M P, Fedrizzi A and White A G 2011 Efficient measurement of quantum dynamics via compressive sensing *Phys. Rev. Lett.* **106** 100401
- [20] Amiet J P and Weigert S 1999 Reconstructing the density matrix of a spin s through three Stern–Gerlach measurements *J. Phys. A: Math. Gen.* **32** 2777
- [21] Flammia S T, Silberfarb A and Caves C M 2005 Minimal informationally complete measurements for pure states *Found. Phys.* **35** 1985
- [22] Merkel S T, Riofrío C A, Flammia S T and Deutsch I H 2010 Random unitary maps for quantum state reconstruction *Phys. Rev. A* **81** 032126
- [23] Heinosaari T, Mazzarella L and Wolf M M 2011 Quantum tomography under prior information arXiv:1109.5478 [quant-ph]
- [24] Kaznady M S and James D F V 2009 Numerical strategies for quantum tomography: alternatives to full optimization *Phys. Rev. A* **79** 022109
- [25] Natarajan B K 1995 Sparse approximate solutions to linear systems *SIAM J. Comput.* **24** 227
- [26] Sturm J 1999 Using SeDuMi 1.02, a MATLAB toolbox for optimization over symmetric cones *Optim. Methods Softw.* **11–2** 625
- [27] Cai J-F, Candès E J and Shen Z 2010 A singular value thresholding algorithm for matrix completion *SIAM J. Optim.* **20** 1956
- [28] Becker S R, Candès E J and Grant M C 2011 Templates for convex cone problems with applications to sparse signal recovery *Math. Program. Comput.* **3** 165–218
- [29] Ma S, Goldfarb D and Chen L 2011 Fixed point and Bregman iterative methods for matrix rank minimization *Math. Program.* **128** 321

- [30] Liu Y-K 2011 Universal low-rank matrix recovery from Pauli measurements *Advances in Neural Information Processing Systems (NIPS)* **24** 1638
- [31] Candès E J and Plan Y 2011 Tight oracle inequalities for low-rank matrix recovery from a minimal number of noisy random measurements *IEEE Trans. Inf. Theory* **57** 2342
- [32] Recht B, Fazel M and Parrilo P A 2010 Guaranteed minimum-rank solutions of linear matrix equations via nuclear norm minimization *SIAM Rev.* **52** 471–501
- [33] Flammia S T and Liu Y-K 2011 Direct fidelity estimation from few Pauli measurements *Phys. Rev. Lett.* **106** 230501
- [34] da Silva M P, Landon-Cardinal O and Poulin D 2011 Practical characterization of quantum devices without tomography *Phys. Rev. Lett.* **107** 210404
- [35] Hradil Z 1997 Quantum-state estimation *Phys. Rev. A* **55** R1561
- [36] Banaszek K, D'Ariano G M, Paris M G A and Sacchi M F 1999 Maximum-likelihood estimation of the density matrix *Phys. Rev. A* **61** 010304
- [37] James D F V, Kwiat P G, Munro W J and White A G 2001 Measurement of qubits *Phys. Rev. A* **64** 052312
- [38] Sugiyama T, Turner P S and Murao M 2011 Error probability analysis in quantum tomography: a tool for evaluating experiments *Phys. Rev. A* **83** 012105
- [39] Shen X 2001 On Bahadur efficiency and maximum likelihood estimation in general parameter spaces *Statistica Sin.* **11** 479
- [40] Chakrabarti R and Ghosh A 2012 Optimal state estimation of controllable quantum dynamical systems *Phys. Rev. A* **85** 032305
- [41] Vogel K and Risken H 1989 Determination of quasiprobability distributions in terms of probability distributions for the rotated quadrature phase *Phys. Rev. A* **40** 2847
- [42] Jones K R W 1991 Principles of quantum inference *Ann. Phys.* **207** 140
- [43] Bužek V, Derka R, Adam G and Knight P L 1998 Reconstruction of quantum states of spin systems: from quantum bayesian inference to quantum tomography *Ann. Phys.* **266** 454
- [44] Gill R D and Massar S 2000 State estimation for large ensembles *Phys. Rev. A* **61** 042312
- [45] Schack R, Brun T A and Caves C M 2001 Quantum Bayes rule *Phys. Rev. A* **64** 014305
- [46] Ježek M, Fiurášek J and Hradil Z 2003 Quantum inference of states and processes *Phys. Rev. A* **68** 012305
- [47] Neri F 2005 Quantum Bayesian methods and subsequent measurements *Phys. Rev. A* **72** 062306
- [48] Tanaka F and Komaki F 2005 Bayesian predictive density operators for exchangeable quantum-statistical models *Phys. Rev. A* **71** 052323
- [49] Bagan E, Ballester M A, Gill R D, Monras A and Muñoz-Tapia R 2006 Optimal full estimation of qubit mixed states *Phys. Rev. A* **73** 032301
- [50] Audenaert K M R and Scheel S 2009 Quantum tomographic reconstruction with error bars: a Kalman filter approach *New J. Phys.* **11** 023028
- [51] Nunn J, Smith B J, Puentes G, Walmsley I A and Lundeen J S 2010 Optimal experiment design for quantum state tomography: fair, precise and minimal tomography *Phys. Rev. A* **81** 042109
- [52] Blume-Kohout R 2010 Optimal, reliable estimation of quantum states *New J. Phys.* **12** 043034
- [53] Blume-Kohout R 2010 Hedged maximum likelihood quantum state estimation *Phys. Rev. Lett.* **105** 200504
- [54] Khoon Ng H and Englert B-G 2012 A simple minimax estimator for quantum states *Int. J. Quantum Inform.* **10** 1250038
- [55] Bužek V 2004 Quantum tomography from incomplete data via MaxEnt principle *Lect. Notes Phys.* **649** 189
- [56] Teo Y S, Zhu H, Englert B-G, Řeháček J and Hradil Z 2011 Quantum-state reconstruction by maximizing likelihood and entropy *Phys. Rev. Lett.* **107** 020404
- [57] Teo Y S, Englert B-G, Řeháček J and Hradil Z 2011 Adaptive schemes for incomplete quantum process tomography *Phys. Rev. A* **84** 062125
- [58] Teo Y S, Stoklasa B, Englert B-G, Řeháček J and Hradil Z 2012 Incomplete quantum state estimation: a comprehensive study *Phys. Rev. A* **85** 042317

- [59] Yin J O S and van Enk S J 2011 Information criteria for efficient quantum state estimation *Phys. Rev. A* **83** 062110
- [60] Tóth G, Wieczorek W, Gross D, Krischek R, Schwemmer C and Weinfurter H 2010 Permutationally invariant quantum tomography *Phys. Rev. Lett.* **105** 250403
- [61] Cramer M, Plenio M B, Flammia S T, Somma R, Gross D, Bartlett S D, Landon-Cardinal O, Poulin D and Liu Y-K 2010 Efficient quantum state tomography *Nature Commun.* **1** 149
- [62] Landon-Cardinal O and Poulin D 2012 Practical learning method for multi-scale entangled states *New J. Phys.* **14** 085004
- [63] Christandl M and Renner R 2011 Reliable quantum state tomography arXiv:1108.5329 [quant-ph]
- [64] Blume-Kohout R 2012 Robust error bars for quantum tomography arXiv:1202.5270 [quant-ph]
- [65] Gross D and Nesme V 2010 Note on sampling without replacing from a finite collection of matrices arXiv:1001.2738
- [66] Fazel M, Candès E, Recht B and Parrilo P 2008 Compressed sensing and robust recovery of low rank matrices *Proc. IEEE 42nd Asilomar Conf. Signals, Systems and Computers* pp 1043–7
- [67] Candès E J and Recht B 2009 Exact matrix completion via convex optimization *Found. Comput. Math.* **9** 717
- [68] Tropp J A 2012 User-friendly tail bounds for sums of random matrices *Found. Comput. Math.* **12** 389–434
- [69] Cover T M and Thomas J A 1991 *Elements of Information Theory* (New York: Wiley)
- [70] Milman V D and Schechtman G 1986 *Asymptotic Theory of Finite Dimensional Normed Spaces* (Berlin: Springer)
- [71] Bhatia R 1996 *Matrix Analysis* (Berlin: Springer)
- [72] Boyd S and Vandenberghe L 2004 *Convex Optimization* (Cambridge: Cambridge University Press)
- [73] Řeháček J, Hradil Z and Ježek M 2001 Iterative algorithm for reconstruction of entangled states *Phys. Rev. A* **63** 040303
- [74] Molina-Terriza G, Vaziri A, Řeháček J, Hradil Z and Zeilinger A 2004 Triggered qutrits for quantum communication protocols *Phys. Rev. Lett.* **92** 167903
- [75] Řeháček J, Hradil Z, Knill E and Lvovsky A I 2007 Diluted maximum-likelihood algorithm for quantum tomography *Phys. Rev. A* **75** 042108
- [76] Jamiolkowski A 1972 Linear transformations which preserve trace and positive semidefiniteness of operators *Rep. Math. Phys.* **3** 275
- [77] Mohseni M and Lidar D A 2006 Direct characterization of quantum dynamics *Phys. Rev. Lett.* **97** 170501
- [78] Mohseni M and Lidar D A 2007 Direct characterization of quantum dynamics: general theory *Phys. Rev. A* **75** 062331
- [79] Mohseni M, Rezakhani A T and Lidar D A 2008 Quantum-process tomography: resource analysis of different strategies *Phys. Rev. A* **77** 032322
- [80] Guta M, Kypraios T and Dryden I 2012 Rank-based model selection for multiple ions quantum tomography arXiv:1206.4032 [quant-ph]