

Mixed state tomography reduces to pure state tomography

Angelos Pelecanos* Jack Spilecki* Ewin Tang* John Wright*

Abstract

A longstanding belief in quantum tomography is that estimating a mixed state is far harder than estimating a pure state. This is borne out in the mathematics, where mixed state algorithms have always required more sophisticated techniques to design and analyze than pure state algorithms. We present a new approach to tomography demonstrating that, contrary to this belief, state-of-the-art mixed state tomography follows easily and naturally from pure state algorithms.

We analyze the following strategy: given n copies of an unknown state ρ , convert them into copies of a purification $|\rho\rangle$; run a pure state tomography algorithm to produce an estimate of $|\rho\rangle$; and output the resulting estimate of ρ . The purification subroutine was recently discovered via the “acorn trick” of Tang, Wright, and Zhandry [TWZ25]. With this strategy, we obtain the first tomography algorithm which is sample-optimal in all parameters. For a rank- r d -dimensional state, it uses $n = O((rd + \log(1/\delta))/\varepsilon)$ samples to output an estimate which is ε -close in fidelity with probability at least $1 - \delta$. This algorithm also uses $\text{poly}(n)$ gates, making it the first gate-efficient tomography algorithm which is sample-optimal even in terms of the dimension d alone. Moreover, with this method we recover essentially all results on mixed state tomography, including its applications to tomography with limited entanglement, classical shadows, and quantum metrology. Our proofs are simple, closing the gap in conceptual difficulty between mixed and pure tomography. Our results also clarify the role of entangled measurement in mixed state tomography: the only step of the algorithm which requires entanglement across copies is the purification step, suggesting that, for tomography, the reason entanglement is useful is for *consistent purification*.

*UC Berkeley. {apelecan,jspilecki,ewin,jswright}@berkeley.edu

Contents

1	Introduction	3
1.1	Mixed state tomography: simpler, faster, and with high probability	4
1.2	Simpler unbiased estimators for mixed state tomography	6
1.2.1	Hayashi’s algorithm	6
1.2.2	The Grier–Pashayan–Schaeffer algorithm	6
1.2.3	Quasi-purification	8
1.2.4	Our new unbiased estimator	10
1.3	Discussion	12
1.4	Organization	14
2	Preliminaries	14
2.1	The symmetric group	14
2.2	The symmetric subspace	17
2.3	Representation theory	17
3	The random purification channel	19
3.1	The symmetric subspace across two registers	19
3.2	A formula for a many-copy random purification	21
3.3	The random purification channel	22
4	Pure state tomography with unentangled measurements	24
5	Pure state tomography with entangled measurements	29
5.1	Hayashi’s algorithm	29
5.2	Grier, Pashayan, and Schaeffer’s algorithm	31
6	A new unbiased estimator for mixed state tomography	36
6.1	Warmup: direct reduction to pure state tomography	36
6.2	Improving the reduction by only quasi-purifying	37
A	Viewing our algorithms as pretty good measurements	43
A.1	PGM preliminaries	43
A.2	The PGM over the Hilbert–Schmidt measure	43
A.3	Viewing $\text{Mix}(\mathcal{A}_{\text{GPS}})$ as a PGM	44

1 Introduction

Given n copies of a mixed state $\rho \in \mathbb{C}^{d \times d}$, *mixed state tomography* refers to the task of producing an estimate of ρ . *Pure state tomography* refers to the special case when ρ is promised to be a pure state. It has long been believed that pure state tomography is significantly easier than mixed state tomography. As Holevo [Hol11, Section 4.11] puts it:

*The full model clearly displays another feature of quantum estimation problem:
the complexity sharply increases with passage from pure to mixed states.*

This sharp increase in complexity appears not just in concrete resource requirements, but also in abstract conceptual difficulty. Pure state tomography can be solved using $n = \Theta(d)$ copies [Hay98, SSW25] with measurements which are unentangled across the n copies of ρ , computationally efficient, and conceptually simple [KRT14, GKKT20]. On the other hand, optimal mixed state tomography requires $n = \Theta(d^2)$ copies [HHJ⁺16, OW16], and the measurements which achieve this *must* be entangled [CHL⁺23, CLL24b], are not known to be computationally efficient, and are conceptually difficult, making use of heavy representation theory [HHJ⁺16, OW16, PSW25]. Put together, this state of affairs suggests that pure and mixed state tomography could not be more hopelessly dissimilar.

We overturn this intuition by showing that mixed state tomography actually *reduces* to pure state tomography. To do so, we make use of the following recent result of Tang, Wright, and Zhandry [TWZ25].

Theorem 1.1 (Efficiently generating random purifications). *Let $n \geq 1$, $d \geq 1$, and $r \leq d$ be integers. There is a quantum channel $\Phi_{\text{Purify}}^{d,r}(\cdot)$ which acts as follows. Given n copies of a rank- r mixed state $\rho \in \mathbb{C}^{d \times d}$,*

$$\Phi_{\text{Purify}}^{d,r}(\rho^{\otimes n}) = \mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n},$$

where the expectation is over a uniformly random purification¹ $|\rho\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r$ of ρ . In addition, $\Phi_{\text{Purify}}^{d,r}(\cdot)$ can be implemented to δ error in diamond distance in time $\text{poly}(n, \log(d), \log(1/\delta))$.

This result is an example of their “acorn trick”, in which copies of a resource (the mixed state ρ) are randomly “lifted” to copies of a stronger resource (the purification $|\rho\rangle$) in a manner which is consistent across all n copies. Similar results have previously appeared in the quantum property testing literature in the works of [SW22, Theorem 35] and especially [CWZ24]. These works showed that having access to purifications does not help solve property testing tasks, and used this to convert lower bounds for mixed state property testing problems into lower bounds for purified property testing problems. Here, we do the opposite: we will use Theorem 1.1 to convert pure state learning *upper* bounds into mixed state learning upper bounds.

In particular, we advocate for constructing mixed state tomography algorithms in the following manner.

Given n copies of ρ :

1. Apply $\Phi_{\text{Purify}}^{d,r}$ to produce n copies of a random purification $|\rho\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r$.
2. Run an off-the-shelf pure state tomography algorithm $\mathcal{A}$ to learn an estimate $\hat{\sigma}$ of $|\rho\rangle\langle\rho|$.
3. Convert $\hat{\sigma}$ to an estimate $\hat{\rho} = \text{tr}_2(\hat{\sigma})$ of ρ . Output $\hat{\rho}$.

Figure 1: A generic reduction from mixed state tomography to pure state tomography. We refer to the resulting mixed state tomography algorithm as $\text{Mix}(\mathcal{A})$.

Both the mixed state ρ and its random purification $|\rho\rangle$ have $\Theta(rd)$ real parameters, so this reduction should incur little loss in terms of sample complexity. In exchange, it makes mixed state tomography significantly more intuitive to understand. For example, while most entangled mixed state tomography algorithms involve complicated representation theory, here the representation theory is fairly elementary and confined to the

¹When we refer to a uniformly random purification, we are referring to the distribution over purifications which is unitarily invariant in the purification register. For example, one can generate a sample from this distribution by taking a fixed purification $|\rho\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r$ and applying a Haar-random unitary U to the second register.

description and analysis of the random purification channel. And if one takes the random purification channel as a black box, then one only needs to analyze $|\rho\rangle^{\otimes n}$, which can be done entirely with symmetric subspace computations.

Through this reduction, we give the first mixed state tomography algorithm which is sample-optimal in all parameters and gate-efficient. Moreover, we reproduce essentially all mixed state tomography bounds which exist in the literature with conceptually simpler algorithms and proofs. Occasionally, this will require us to make a minor modification to this generic reduction; we will discuss what this modification is and why it seems to be necessary when it arises below.

There are two main types of pure state tomography algorithms, corresponding to the unentangled and entangled measurement settings. This leads to two types of mixed state tomography algorithms, both of which we investigate.

1.1 Mixed state tomography: simpler, faster, and with high probability

We begin by instantiating the reduction with an unentangled pure state tomography algorithm. The “standard” such algorithm is given in [Figure 2](#).

Given n copies of $\sigma \in \mathbb{C}^{d \times d}$:

1. For each $1 \leq i \leq n$:
 - (a) Measure the i -th copy of σ with the uniform POVM $\{d \cdot |u\rangle\langle u| \cdot du\}$. Let $|\mathbf{v}_i\rangle$ be the outcome.
 - (b) Set $\hat{\sigma}_i = (d+1) \cdot |\mathbf{v}_i\rangle\langle \mathbf{v}_i| - I_d$.
2. Output $\hat{\sigma}_{\text{avg}} = \frac{1}{n} \cdot (\hat{\sigma}_1 + \dots + \hat{\sigma}_n)$.

Figure 2: The standard unentangled measurement tomography algorithm.

When σ is promised to be a pure state $|\psi\rangle\langle\psi|$, one might hope for the estimator produced by the algorithm to also be a pure state. In this case, it is natural to post-process the output of this algorithm by computing the top eigenvector $|\mathbf{v}\rangle$ of $\hat{\sigma}_{\text{avg}}$ and using this as the estimator for $|\psi\rangle$ instead. This pure state tomography algorithm was introduced in the work of Guta, Kahn, Kueng, and Tropp, who showed that $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$ when using $n = O((d + \log(1/\delta))/\varepsilon)$ copies [[GKKT20](#), Theorem 5]. One can even make this algorithm computationally efficient via a straightforward application of unitary t -designs, as we show in the next theorem. We prove this result in [Section 4](#) below.

Theorem 1.2 (Efficient pure state tomography). *There is an algorithm $\mathcal{A}_{\text{GKKT}}$ which, given*

$$n = O\left(\frac{d + \log(1/\delta)}{\varepsilon}\right)$$

copies of a pure state $|\psi\rangle \in \mathbb{C}^d$, outputs a pure state $|\mathbf{v}\rangle \in \mathbb{C}^d$ such that $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$. Furthermore, this algorithm can be implemented in $\text{poly}(n)$ time² and performs independent measurements across the copies of $|\psi\rangle$.

Combining our reduction from mixed state tomography to pure state tomography with this pure state tomography algorithm results in the strongest mixed state tomography algorithm currently known. We attain the correct dependence on failure probability δ , making this algorithm sample-optimal in all parameters. Further, we attain this optimal dependence with a gate-efficient algorithm.

Theorem 1.3 (Efficient mixed state tomography). *Let $\hat{\rho}$ be the output of the algorithm $\text{Mix}(\mathcal{A}_{\text{GKKT}})$ when run on*

$$n = O\left(\frac{rd + \log(1/\delta)}{\varepsilon}\right)$$

²When we state running time of algorithms, we imagine that the Hilbert space $\mathbb{C}^d$ is being represented on a system of $\lceil \log_2(d) \rceil$ qubits on a quantum computer. So, running time refers to the number of one- and two-qubit gates used by the algorithm, including classical post-processing.

copies of a rank- r mixed state $\rho \in \mathbb{C}^{d \times d}$. Then $F(\rho, \hat{\rho}) \geq 1 - \varepsilon$ with probability at least $1 - \delta$. Furthermore, $\text{Mix}(\mathcal{A}_{\text{GKKT}})$ can be implemented in $\text{poly}(n)$ time.

Proof. The first step of $\text{Mix}(\mathcal{A}_{\text{GKKT}})$ is to apply $\Phi_{\text{Purify}}^{d,r}$ to $\rho^{\otimes n}$ to produce n copies of a random purification $|\rho\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r \cong \mathbb{C}^{dr}$. Next, we apply $\mathcal{A}_{\text{GKKT}}$ to learn an estimate $|\mathbf{v}\rangle$ of $|\rho\rangle$. By [Theorem 1.2](#), with probability at least $1 - \delta/2$, this estimate will satisfy $|\langle \mathbf{v} | \rho \rangle|^2 \geq 1 - \varepsilon$. Now, set $\hat{\rho} = \text{tr}_2(|\mathbf{v}\rangle\langle \mathbf{v}|)$. By Uhlmann's theorem,

$$F(\rho, \hat{\rho}) = \max_{|\psi_\rho\rangle, |\psi_{\hat{\rho}}\rangle} |\langle \psi_{\hat{\rho}} | \psi_\rho \rangle|^2 \geq |\langle \mathbf{v} | \rho \rangle|^2,$$

where the maximization is over all purifications $|\psi_\rho\rangle$ and $|\psi_{\hat{\rho}}\rangle$ of ρ and $\hat{\rho}$, respectively. Therefore, $F(\rho, \hat{\rho}) \geq 1 - \varepsilon$ with probability at least $1 - \delta/2$.

To make this efficient, note that the pure state tomography algorithm $\mathcal{A}_{\text{GKKT}}$ can be implemented in time $\text{poly}(n) = \text{poly}(d, 1/\varepsilon, \log(1/\delta))$ by [Theorem 1.2](#). Similarly, we can implement the purifying channel $\Phi_{\text{Purify}}^{d,r}$ to error $\delta/2$ in diamond distance in time $\text{poly}(n, \log(d), \log(1/\delta)) = \text{poly}(d, 1/\varepsilon, \log(1/\delta))$ by [Theorem 1.1](#). This will introduce an additional $\delta/2$ probability of failure, meaning that the algorithm succeeds with probability at least $1 - \delta$. This completes the proof. $\square$

This is the first tomography algorithm which is sample-optimal in all four parameters d, r, ε , and δ ; optimality of the $O(rd/\varepsilon)$ term follows from the lower bound of [\[Yue23\]](#) (see also the lower bound of [\[SSW25\]](#)), and optimality of the $O(\log(1/\delta)/\varepsilon)$ term follows from estimating the bias of a coin with high probability. Prior to this work, the best known bounds were

$$n = O\left(\frac{rd}{\varepsilon} \cdot \log\left(\frac{rd + \log(1/\delta)}{\varepsilon}\right) + \frac{\log(1/\delta)}{\varepsilon}\right) \quad \text{and} \quad n = O\left(\frac{rd}{\varepsilon} \cdot \log(1/\delta)\right),$$

which follow from [\[HHJ⁺16, Equation \(14\)\]](#) and from combining [\[PSW25, Theorem 1.6\]](#) with standard amplification results (for example, [\[HKOT23, Proposition 2.4\]](#)), respectively. Our bound gives a strict improvement if δ is smaller than constant and larger than $(rd)^{-rd}$.

This also marks a major improvement in the sample-complexity of time-efficient tomography; we are not aware of an explicit theorem statement in the existing literature, but the works which could be made time-efficient with standard techniques [\[GKKT20, KRT14\]](#) achieve a far-from-optimal scaling of $n = O(r^2d)$ when ε and δ are constant.

Even including proofs of intermediate results, the proof of [Theorem 1.3](#) is simple: ignoring time-efficiency, all it needs are basic representation theory (used in the proof of [Theorem 1.1](#)) and a scalar concentration inequality (used in the proof of [Theorem 1.2](#)). By comparison, all prior work requires more advanced representation theory, especially of the unitary group. In addition, most prior works achieve suboptimal sample complexities for learning in fidelity even for constant settings of the error probability δ [\[HHJ⁺16, OW16, OW17\]](#), and so their techniques seem unable to establish optimal sample complexity bounds in terms of all four parameters. The one exception is the recent work of [\[PSW25\]](#), which shows a tight sample complexity bound of $n = O(rd/\varepsilon)$ for learning with constant error probability δ . They do so by analyzing the second moment of their estimator, and it seems plausible that one could extend their result to arbitrary error probabilities δ by analyzing higher moments of their estimator. However, their second moment proof is already quite involved, spanning dozens of pages of complicated representation-theoretic calculations. A proof for higher moments, without further ideas, seems like it would be at best a tremendous chore and at worst completely intractable.³

An interesting feature of the algorithm in [Theorem 1.3](#) is that it only uses entanglement across the copies of ρ in order to produce consistent purifications of ρ . After doing this, as our algorithm shows, it suffices to only use unentangled measurements on the purified copies. This suggests that entangled measurements are helpful in mixed state tomography *because* they allow us to produce consistent purifications across all the copies of ρ . This is consistent with our understanding of tomography in these settings: for mixed states, it is

³One could also consider the easier task of learning ρ to error ε in trace distance. There, the situation is similar. [Theorem 1.3](#) implies an algorithm with an optimal sample complexity of $n = O((rd + \log(1/\delta))/\varepsilon^2)$. This bound was not known in the literature; the previous best bounds come from the fidelity algorithms discussed here, translated to trace distance. Note that, for this easier regime, the upper bound of $O(rd \log(1/\delta)/\varepsilon^2)$ was proved earlier [\[OW15\]](#) and the lower bound of $\Omega(rd/\varepsilon^2)$ was proved later [\[SSW25\]](#). Even for trace distance, it does not seem like any prior works had a clear pathway towards achieving a fully sample-optimal algorithm.

known that sample-optimal mixed state tomography requires entangled measurements [CHL⁺23, CLL24b], but for pure states, as [Theorem 1.2](#) shows, sample-optimal tomography can be achieved with algorithms which perform independent measurements across the input copies. We will explore this theme in further detail in [Section 1.3](#) below.

1.2 Simpler unbiased estimators for mixed state tomography

The estimator described above is sample-optimal for the task of generic full-state tomography. However, there are other, more fine-grained tomographic tasks which have gained importance in the literature, and it turns out that this estimator performs sub-optimally at these tasks. We will discuss exactly why this is the case in [Section 1.3](#), but at a high level, these tasks require having an estimator which has low variance about its mean, and the above estimator has prohibitively high variance. To tackle these other applications, we will consider applying our reduction to another pure state tomography algorithm, one due to Hayashi.

1.2.1 Hayashi’s algorithm

Next, we instantiate the reduction with the “standard” *entangled* pure state tomography algorithm. When performing tomography on n copies of a pure state $|\psi\rangle \in \mathbb{C}^d$, the input $|\psi\rangle^{\otimes n}$ is an element of the symmetric subspace $\vee^n \mathbb{C}^d$. This means that a pure state tomography algorithm’s measurement operators need only be specified on the symmetric subspace. Motivated by this, Hayashi [Hay98] introduced the following natural pure state tomography algorithm $\mathcal{A}_{\text{Hayashi}}$: simply perform the POVM

$$\{d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du\} \quad (1)$$

and output the pure state $|\mathbf{v}\rangle$ that it returns. (Here, $d[n] = \dim(\vee^n \mathbb{C}^d)$.) This is indeed a valid POVM on the symmetric subspace, as

$$\int_{|u\rangle} d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du = d[n] \cdot \mathbf{E}_{|\mathbf{u}\rangle \sim \text{Haar}} |\mathbf{u}\rangle\langle \mathbf{u}|^{\otimes n} = \Pi_{\text{sym}},$$

where Π_{sym} is the projector onto $\vee^n \mathbb{C}^d$, as we discuss in the Preliminaries below (cf. [Equation \(9\)](#)). Hayashi showed that this algorithm is in fact the *optimal* pure state tomography algorithm in a certain precise technical sense, and so it is perhaps the most natural algorithm to apply our reduction to.

It is well-known that the output of $\mathcal{A}_{\text{Hayashi}}$ satisfies $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability 99% when $n = O(d/\varepsilon)$. We show that it also achieves optimal dependence on the failure probability δ .

Proposition 1.4 (Hayashi’s algorithm with high probability). *Given n copies of a pure state $|\psi\rangle \in \mathbb{C}^d$, suppose $\mathcal{A}_{\text{Hayashi}}$ outputs the state $|\mathbf{v}\rangle$. Then $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$ when*

$$n = O\left(\frac{d + \log(1/\delta)}{\varepsilon}\right).$$

To our knowledge, this bound has not been previously observed in the literature for Hayashi’s algorithm. It actually follows immediately from combining Hayashi’s optimality result for his algorithm with the fact that the Guta–Kahn–Kueng–Tropp pure state tomography algorithm also achieves this bound [GKKT20, Theorem 5]. We give an alternative proof of this fact which analyzes the output of Hayashi’s algorithm directly. Combined with our reduction, this gives a second mixed state algorithm $\text{Mix}(\mathcal{A}_{\text{Hayashi}})$ which achieves the sample complexity bound in [Theorem 1.3](#). However, it is not computationally efficient, as Hayashi’s algorithm is not known to be computationally efficient.

1.2.2 The Grier–Pashayan–Schaeffer algorithm

Grier, Pashayan, and Schaeffer [GPS24] considered the following modification to Hayashi’s algorithm.

Given n copies of $\sigma \in \mathbb{C}^{d \times d}$:

1. Measure the copies with the POVM $\{d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du\}$. Let $|\mathbf{v}\rangle$ be the outcome.
2. Output the estimator

$$\hat{\sigma}_{\mathbf{v}} := \frac{d+n}{n} \cdot |\mathbf{v}\rangle\langle \mathbf{v}| - \frac{1}{n} \cdot I_d.$$

Figure 3: The Grier–Pashayan–Schaeffer tomography algorithm $\mathcal{A}_{\text{GPS}}$.

They observed that this modification results in an *unbiased estimator* for pure state tomography, meaning that if this algorithm is performed on n copies of a pure state $\sigma \in \mathbb{C}^{d \times d}$, then its output satisfies $\mathbf{E}[\hat{\sigma}_{\mathbf{v}}] = \sigma$. The quality of an unbiased estimator is governed by how much it deviates from its mean, which we can quantify using its variance $\mathbf{E}[(\hat{\sigma}_{\mathbf{v}} - \sigma)^{\otimes 2}] = \mathbf{E}[\hat{\sigma}_{\mathbf{v}}^{\otimes 2}] - \sigma^{\otimes 2}$. This entails calculating its second moment $\mathbf{E}[\hat{\sigma}_{\mathbf{v}}^{\otimes 2}]$, which we do as follows.

Theorem 1.5 (Moments of Grier–Pashayan–Schaeffer). *Let $\hat{\sigma}_{\mathbf{v}}$ be the output of $\mathcal{A}_{\text{GPS}}$ when run on n copies of a pure state $\sigma \in \mathbb{C}^{d \times d}$. Then $\hat{\sigma}_{\mathbf{v}}$ is an unbiased estimator for σ , i.e. $\mathbf{E} \hat{\sigma}_{\mathbf{v}} = \sigma$. In addition,*

$$\mathbf{E}[\hat{\sigma}_{\mathbf{v}} \otimes \hat{\sigma}_{\mathbf{v}}] = \frac{n-1}{n} \cdot \sigma \otimes \sigma + \frac{1}{n} \cdot (\sigma \otimes I_d + I_d \otimes \sigma) \cdot \text{SWAP} + \frac{1}{n^2} \cdot \text{SWAP} - \text{Lower}_{\sigma},$$

where $\text{Lower}_{\sigma} \in \text{SoS}(d)$.

This theorem uses the following definition for the Lower_{σ} term.

Definition 1.6 (Sum of Hermitian squares). Given an integer d , we define

$$\text{SoS}(d) := \text{cone}(\{X \otimes X \mid X \in \mathbb{C}^{d \times d} \text{ is Hermitian}\}),$$

where $\text{cone}(\cdot)$ is the conical hull of its input, i.e. the set of all nonnegative linear combinations of matrices in its input set.

The Lower_{σ} term is so-named because it is lower order in the parameter d and tends towards 0 as $d \rightarrow \infty$. What is important to us is that for the applications we care about, its contribution always turns out to be negative and hence can be discarded, as we will discuss in more detail below.

This suggests a natural unbiased estimator for mixed state tomography: simply apply our reduction to the Grier–Pashayan–Schaeffer algorithm. The result is the following algorithm.

Given n copies of ρ :

1. First apply $\Phi_{\text{Purify}}^{d,r}$ to produce n copies of a random purification $|\rho\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r$.
2. Apply the Grier–Pashayan–Schaeffer algorithm to learn an estimate $\hat{\sigma}_{\mathbf{v}}$ of $|\rho\rangle\langle \rho|$.
3. Set $\hat{\rho}_{\mathbf{v}} = \text{tr}_2(\hat{\sigma}_{\mathbf{v}})$ of ρ . Output $\hat{\rho}_{\mathbf{v}}$.

Figure 4: The mixed state tomography algorithm $\text{Mix}(\mathcal{A}_{\text{GPS}})$.

By construction, this produces an unbiased estimator, as for any purification $|\rho\rangle$ of ρ ,

$$\mathbf{E}[\hat{\rho}_{\mathbf{v}}] = \mathbf{E}[\text{tr}_2(\hat{\sigma}_{\mathbf{v}})] = \text{tr}_2(\mathbf{E}[\hat{\sigma}_{\mathbf{v}}]) = \text{tr}_2(|\rho\rangle\langle \rho|) = \rho.$$

Using [Theorem 1.5](#), we will show the following expression for the second moment of this estimator.

Theorem 1.7 (Moments of the Grier–Pashayan–Schaeffer mixed state tomography algorithm). *Let $\hat{\rho}_{\mathbf{v}}$ be the output of $\text{Mix}(\mathcal{A}_{\text{GPS}})$ when run on n copies of a rank- r state $\rho \in \mathbb{C}^{d \times d}$. Then $\hat{\rho}_{\mathbf{v}}$ is an unbiased estimator for ρ with second moment*

$$\mathbf{E}[\hat{\rho}_{\mathbf{v}} \otimes \hat{\rho}_{\mathbf{v}}] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{r}{n^2} \cdot \text{SWAP} - \text{Lower}_{\rho},$$

where $\text{Lower}_{\rho} \in \text{SoS}(d)$.

Historically, designing good unbiased estimators for mixed state tomography has been a challenging task. Until recently, our only known unbiased estimators, such as the estimator from [Figure 2](#), were not sample optimal, and our only known sample-optimal estimators were not unbiased [\[OW16, HHJ⁺16\]](#). This changed with the work of Pelecanos, Spilecki, and Wright [\[PSW25\]](#), who introduced the *debiased Keyl's algorithm*, the first estimator for mixed state tomography which is both sample-optimal and unbiased. Using it, they proved a number of new and optimal sample complexity upper bounds on a number of interesting tomographic tasks, which we will describe in [Section 1.2.4](#) below.

Let us compare the performance of their estimator to the estimator produced by $\text{Mix}(\mathcal{A}_{\text{GPS}})$. Writing $\hat{\rho}$ for the output of their debiased Keyl's algorithm on n copies of a rank- r mixed state $\rho \in \mathbb{C}^{d \times d}$, they showed that

$$\mathbf{E}[\hat{\rho} \otimes \hat{\rho}] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \cdot \text{SWAP} - \text{Lower}_\rho, \quad (2)$$

where $\text{Lower}_\rho \in \text{SoS}(d)$ [\[PSW25, Theorem 1.4\]](#). This matches the bound in [Theorem 1.7](#) on all terms except the factor of r on the third term is replaced with the term $\mathbf{E}[\ell(\boldsymbol{\lambda})]$. We will explain what exactly this notation means below; for now, it suffices to know that it can be upper-bounded by $\mathbf{E}[\ell(\boldsymbol{\lambda})] \leq \min\{r, 2\sqrt{n}\}$, which is smaller than r whenever $n < r^2/4$. This means that the debiased Keyl's algorithm will actually outperform $\text{Mix}(\mathcal{A}_{\text{GPS}})$ when the number of copies $n = o(r^2)$. As we describe below, some of our applications operate in the regime of $n = o(r^2)$ copies, in which case this difference is significant, and some of our applications operate in the regime of $n = \omega(r^2)$ copies, in which case this difference is insignificant and the two algorithms behave similarly.

Our main goal, however, will be to improve our algorithm so that its second moment matches that of the debiased Keyl's algorithm. Doing so will require us to design a modified version of the purification channel, which we describe in the next section.

Remark 1.8 (On the lower term). Let us point out one subtle difference between the way we have stated the second moment formula for the debiased Keyl's algorithm and the way it appears in [\[PSW25, Theorem 1.4\]](#). In [\[PSW25\]](#), the Lower_ρ term is characterized as being a positive linear combination of matrices of the form $(P \otimes P) \cdot \text{SWAP}$, where P is Hermitian and positive semidefinite. As it turns out, any matrix of this form can be shown to be an element of $\text{SoS}(d)$, and so their $\text{Lower}_\rho \in \text{SoS}(d)$ as well. This means that [Equation \(2\)](#) is actually a slightly weaker characterization of the second moment than the one given in [\[PSW25\]](#), but as we show below, this is still sufficient to derive all of their applications.

1.2.3 Quasi-purification

The single copy case. To understand why $\text{Mix}(\mathcal{A}_{\text{GPS}})$ performs sub-optimally in the regime of small n , it will help to first gain some understanding for how the purification channel operates in the simplest case of $n = 1$ copy. Let $\rho \in \mathbb{C}^{d \times d}$ be the rank- r state to be purified, and write $\rho = \sum_{i=1}^r \alpha_i \cdot |v_i\rangle\langle v_i|$ for its eigendecomposition. Consider the purification of ρ given by

$$|\rho_0\rangle_{\text{AB}} := \sum_{i=1}^r \sqrt{\alpha_i} \cdot |v_i\rangle_{\text{A}} \otimes |i\rangle_{\text{B}},$$

where B is an r -dimensional register. One can generate a random purification of ρ by sampling a Haar random unitary $U \in U(r)$ and outputting $|\rho\rangle := U_{\text{B}} \cdot |\rho_0\rangle_{\text{AB}}$. The mixture over these random purifications is given by

$$\mathbf{E}|\rho\rangle\langle\rho| = \mathbf{E}[U_{\text{B}} \cdot |\rho_0\rangle\langle\rho_0|_{\text{AB}} \cdot U_{\text{B}}^\dagger] = \text{tr}_{\text{B}}(|\rho_0\rangle\langle\rho_0|) \otimes (I_r/r)_{\text{B}} = \rho_{\text{A}} \otimes (I_r/r)_{\text{B}}. \quad (3)$$

Implementing the random purification channel $\Phi_{\text{Purify}}^{d,r}(\rho)$ is therefore easy: simply take ρ and append to it a maximally mixed state in a second register.

After performing this purification, $\text{Mix}(\mathcal{A}_{\text{GPS}})$ will perform pure state tomography on both registers of the purified state in [Equation \(3\)](#). By [Theorem 1.7](#), it will output an estimator $\hat{\rho}_{\mathbf{v}}$ satisfying

$$\mathbf{E}[\hat{\rho}_{\mathbf{v}} \otimes \hat{\rho}_{\mathbf{v}}] = (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + r \cdot \text{SWAP} - \text{Lower}_\rho. \quad (4)$$

However, the purification register in [Equation \(3\)](#) contains no information whatsoever, and so it seems wasteful to include it when performing the pure state tomography step. Instead, what we can do is omit the

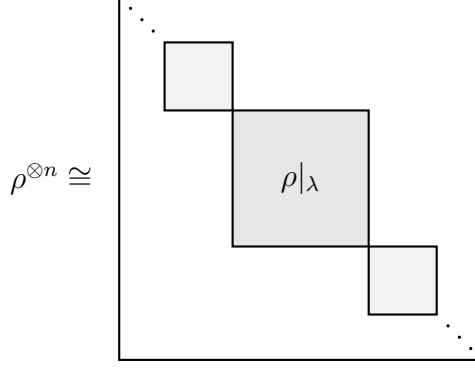


Figure 5: Schur–Weyl duality applied to $\rho^{\otimes n}$. Here, $\rho|_{\lambda}$ is the normalized restriction of ρ to the λ -block.

purification register altogether and simply perform pure state tomography on ρ itself. This at least “type checks” because although ρ is itself not necessarily a pure state, it is still a mixture over the pure states corresponding to its eigenvectors. So perform the $\mathcal{A}_{\text{GPS}}$ algorithm on ρ directly and let $\hat{\sigma}_{\mathbf{v}}$ be its output. Then by [Theorem 1.5](#),

$$\mathbf{E}[\hat{\sigma}_{\mathbf{v}} \otimes \hat{\sigma}_{\mathbf{v}}] = (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \text{SWAP} - \text{Lower}_{\rho}.$$

This improves on [Equation \(4\)](#) by a factor of r on the SWAP term.

The general case. Generalizing this to larger values of n is conceptually more interesting than simply discarding the purification register and requires the use of some representation theory. In particular, let us recall *Schur–Weyl duality*, which states that there is a unitary change of basis U_{Schur} known as the *Schur transform* under which the n copy state $\rho^{\otimes n}$ becomes block diagonal, with a block for every partition $\lambda \vdash n$ of height $\ell(\lambda) \leq d$. We illustrate this in [Figure 5](#). It is common for entangled tomography algorithms to begin with a step known as *weak Schur sampling*, in which one measures $\rho^{\otimes n}$ with the projective measurement $\{\Pi_{\lambda}\}$, where Π_{λ} is the projector onto the subspace corresponding to λ in the Schur basis. This produces a random partition $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_d)$ as a measurement outcome and collapses $\rho^{\otimes n}$ to the state $\rho|_{\boldsymbol{\lambda}}$. Typically, one then performs a further measurement within the $\boldsymbol{\lambda}$ -subspace on the state $\rho|_{\boldsymbol{\lambda}}$ in order to learn ρ .

The length $\ell(\boldsymbol{\lambda})$ of the measurement outcome $\boldsymbol{\lambda}$, defined to be the number of nonzero coordinates λ_i of the partition, can be viewed as a loose estimate for the rank r of ρ . It is in fact an *underestimate*, as it always satisfies $\ell(\boldsymbol{\lambda}) \leq r$, and for small values of n it can underestimate r by a significant amount. For example, when $n = 1$, the length $\ell(\boldsymbol{\lambda}) = 1$ always. More generally, for larger n , its expectation satisfies $\mathbf{E}[\ell(\boldsymbol{\lambda})] \leq 2\sqrt{n}$, which is significantly smaller than r so long as $n = o(r^2)$.

What is so nice about this is that for the sake of purification, it turns out that we can treat $\rho|_{\ell(\boldsymbol{\lambda})}$ as if it came from a rank- $\ell(\boldsymbol{\lambda})$ state rather than a rank- r state. In particular, we will see that the rank- k purification channel applied to this state, which outputs the state

$$\Phi_{\text{Purify}}^{d,k}(\rho|_{\boldsymbol{\lambda}}),$$

is at the very least well-defined so long as $k \geq \ell(\boldsymbol{\lambda})$. In addition, its output, while no longer necessarily a mixture over states of the form $|\rho\rangle^{\otimes n}$, is still guaranteed to be an element of the symmetric subspace $\vee^n(\mathbb{C}^d \otimes \mathbb{C}^{\ell(\boldsymbol{\lambda})})$, and so it at least “type checks” to run pure state tomography on it.

Our basic purification algorithm can be viewed as the algorithm which always performs rank $k = r$ purification on this state, no matter what the length $\ell(\boldsymbol{\lambda})$ of the partition is. We will now consider a more fine-grained purification algorithm which performs rank $k = \ell(\boldsymbol{\lambda})$ purification, the smallest value of k possible. We refer to this operation as *quasi-purification*, as although it is applying the purification channel, it can no longer be viewed as outputting purified copies of ρ . In addition, it is not even a channel anymore, as its output space $\vee^n(\mathbb{C}^d \otimes \mathbb{C}^{\ell(\boldsymbol{\lambda})})$ is no longer fixed, but depends on the measured $\boldsymbol{\lambda}$. However, we can still integrate into our generic reduction, which yields the following “upgraded” reduction from mixed to pure state tomography.

Given n copies of ρ :

1. Run weak Schur sampling to produce a Young diagram $\lambda \vdash n$; the n states collapse to $\rho|_\lambda$.
2. Apply $\Phi_{\text{Purify}}^{d, \ell(\lambda)}$ to the resulting state to produce $\Phi_{\text{Purify}}^{d, \ell(\lambda)}(\rho|_\lambda) \in (\mathbb{C}^d \otimes \mathbb{C}^{\ell(\lambda)})^{\otimes n}$.
3. Run an off-the-shelf pure state tomography algorithm $\mathcal{A}$ on $\Phi_{\text{Purify}}^{d, \ell(\lambda)}(\rho|_\lambda)$; call the output $\hat{\sigma}^\lambda$.
4. Convert $\hat{\sigma}^\lambda$ to an estimate $\hat{\rho}^\lambda = \text{tr}_2(\hat{\sigma}^\lambda)$ of ρ . Output $\hat{\rho}^\lambda$.

Figure 6: A tighter version of our main reduction. We refer to the resulting mixed state tomography algorithm as $\text{Mix}^+(\mathcal{A})$.

Note that this generalizes our $n = 1$ example from above: when $n = 1$, we have that $\ell(\lambda) = 1$ as well. In this case, we want to apply the purification channel $\Phi_{\text{Purify}}^{d, 1}(\cdot)$, but it is easy to see that this channel should not modify its input, as a rank-1 state is already pure.

1.2.4 Our new unbiased estimator

We now combine quasi-purification with the Grier–Pashayan–Schaeffer algorithm. The resulting algorithm is still unbiased and has the following second moment guarantee.

Theorem 1.9 (Moments of our unbiased estimator). *Let $\hat{\rho}_v^\lambda$ be the output of $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ when run on n copies of a rank- r state $\rho \in \mathbb{C}^{d \times d}$. Then $\hat{\rho}_v^\lambda$ is an unbiased estimator for ρ with second moment*

$$\mathbf{E}[\hat{\rho}_v^\lambda \otimes \hat{\rho}_v^\lambda] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \cdot \text{SWAP} - \text{Lower}_\rho,$$

where $\text{Lower}_\rho \in \text{SoS}(d)$.

This matches the second moment formula that Pelecanos, Spilecki, and Wright proved for the debiased Keyl’s algorithm, which we saw previously in Equation (2). They showed how to use this second moment formula to derive a number of applications of the debiased Keyl’s algorithm. We will briefly survey these applications below; for more background on these applications, see [PSW25, Section 1], and for proofs that they can be derived from the debiased Keyl’s algorithm, see [PSW25, Part II]. Because our estimator $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ has the same second moment formula as the debiased Keyl’s algorithm, these applications hold for it as well with essentially identical proofs. The one minor modification needed to adapt these proofs to our algorithm comes from the fact that the Lower_ρ term takes a slightly different form in our algorithm than it does in the debiased Keyl’s algorithm, as we saw in Remark 1.8. To address this, we will explain below what properties of the Lower_ρ term each of these applications need, and we will show the Lower_ρ term from our algorithm does indeed satisfy these properties.

Application 1: tomography with limited entanglement. In the k -entangled tomography problem, the goal is to estimate an unknown mixed state ρ while performing entangled measurements on at most k copies of ρ at a time. The natural algorithm for doing so is the following.

1. Divide the n copies of ρ into $n' := n/k$ batches of size k .
2. For each $1 \leq i \leq n'$, run a mixed state tomography algorithm on the i -th batch of copies and let $\hat{\rho}_i$ be its output.
3. Output the estimator $\hat{\rho} = \frac{1}{n'} \cdot (\hat{\rho}_1 + \dots + \hat{\rho}_{n'})$.

When the $\hat{\rho}_i$ ’s are produced by an unbiased estimator, then averaging them together to produce $\hat{\rho}$ results in an estimator which remains unbiased, but has significantly decreased variance. Pelecanos, Spilecki, and Wright showed that when the debiased Keyl’s algorithm is used, the $\hat{\rho}$ this algorithm produces is ε -close to ρ in trace distance with probability 99% when

$$n = O\left(\max\left(\frac{d^3}{\sqrt{k}\varepsilon^2}, \frac{d^2}{\varepsilon^2}\right)\right) \quad (5)$$

copies of ρ are used [PSW25, Theorem 1.8]. This improves on prior work of Chen, Li, and Liu [CLL24b] and matches their lower bound for this task of $n = \Omega(d^3/(\sqrt{k}\varepsilon^2))$ copies, which they showed for the case when $k \leq 1/\varepsilon^c$, for c a small constant.

The proof of [PSW25] uses the second moment formula for the debiased Keyl’s algorithm, and as a result their sample complexity bound also applies if we use our estimator $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$ instead. The one property of Lower_ρ that this proof needs is that $\text{tr}(\text{SWAP} \cdot \text{Lower}_\rho) \geq 0$ (cf. the proof of [PSW25, Lemma 4.2]). This holds in our case too, as our Lower_ρ is a nonnegative linear combination of terms of the form $X \otimes X$, where X is a Hermitian matrix, and

$$\text{tr}(\text{SWAP} \cdot X \otimes X) = \text{tr}(X^2) \geq 0,$$

which holds because X is Hermitian and therefore X^2 has all nonnegative eigenvalues. We note that the interesting regime of the sample complexity bound in Equation (5) is when $k = o(d^2)$; in this case, we truly do need the quasi-purification-based algorithm $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$ rather than $\mathbf{Mix}(\mathcal{A}_{\text{GPS}})$ to achieve optimal sample complexity.

Application 2: shadow tomography. In the shadow tomography problem, one is given m bounded observables $O_1, \dots, O_m \in \mathbb{C}^{d \times d}$ which satisfy $\|O_i\|_\infty \leq 1$, for all $1 \leq i \leq m$, and asked to estimate the observable values $\text{tr}(O_1 \cdot \rho), \dots, \text{tr}(O_m \cdot \rho)$ up to ε accuracy each. The natural strategy for doing so is the following “plug-in” approach.

1. Run a mixed state tomography algorithm on n' copies of ρ . Let $\hat{\rho}$ be the estimator it produces.
2. Output $\hat{o}_1 := \text{tr}(O_1 \cdot \hat{\rho}), \dots, \hat{o}_m := \text{tr}(O_m \cdot \hat{\rho})$.

When $\hat{\rho}$ is an unbiased estimator for ρ , the $\hat{o}_i$ ’s are unbiased estimators for the true observable values $\text{tr}(O_i \cdot \rho)$. If $\hat{\rho}$ has small variance about its mean, then these $\hat{o}_i$ ’s will also have small variance about their means; in particular, one wants to take n' large enough so that each $\hat{o}_i$ is within ε of its mean with 99% probability. To ensure that all $\hat{o}_i$ ’s are within ε if their mean at once, one can then perform the following algorithm.

1. Repeat the “plug-in” approach k times, producing the estimators $\hat{o}_1^1, \dots, \hat{o}_m^1$ through $\hat{o}_1^k, \dots, \hat{o}_m^k$.
2. For each $1 \leq i \leq m$, output the estimator $\hat{o}_i = \text{median}\{\hat{o}_i^1, \dots, \hat{o}_i^k\}$.

In total, the whole process takes $n = k \cdot n'$ copies of ρ . Pelecanos, Spilecki, and Wright showed that when the debiased Keyl’s algorithm is used for the plug-in estimator, then this algorithm succeeds with probability 99% when using

$$n = O\left(\log(m) \cdot \left(\min\left\{\frac{\sqrt{rF}}{\varepsilon}, \frac{F^{2/3}}{\varepsilon^{4/3}}\right\} + \frac{1}{\varepsilon^2}\right)\right) \quad (6)$$

copies of a rank r state ρ [PSW25, Theorem 1.9]. Here, each observable O_i is assumed to satisfy the bound $\text{tr}(O_i^2) \leq F$. Since the measurements this algorithm performs are independent of the observables O_i , it also solves the related “classical shadows” problem, in which one is provided the observables only after measuring, with this sample complexity, and in doing so it improves on the prior works of [HKP20, GLM24]. In addition, note that because the observables O_i satisfy $\|O_i\|_\infty \leq 1$, we have $F \leq d$. Thus, in the “high accuracy regime” of $\varepsilon = O(1/d)$, this shows that $n = O(\log(m)/\varepsilon^2)$ copies suffice, improving on a bound of [CLL24a].

The proof of [PSW25] uses the second moment formula for the debiased Keyl’s algorithm, and as a result their sample complexity bound also applies if we use our estimator $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$ instead. The one property of Lower_ρ that this proof needs is that $\text{tr}(O \otimes O \cdot \text{Lower}_\rho) \geq 0$ for any observable O (cf. the proof of [PSW25, Lemma 7.3]). This holds in our case too, as our Lower_ρ is a nonnegative linear combination of terms of the form $X \otimes X$, where X is a Hermitian matrix, and

$$\text{tr}(O \otimes O \cdot X \otimes X) = \text{tr}(OX)^2 \geq 0,$$

which holds because X is Hermitian and therefore $\text{tr}(OX)$ is a real number whose square is therefore nonnegative. Finally, we note that to achieve the full sample complexity bound of Equation (6), we do require using the quasi-purified algorithm $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$. However, in the high accuracy regime when $\varepsilon = O(1/d)$, if we assume bounds of $r, F \leq d$, then our algorithm requires using $n = \Omega(d^2)$ copies of ρ , in which case it suffices to use the simpler algorithm $\mathbf{Mix}(\mathcal{A}_{\text{GPS}})$.

Application 3: quantum metrology. In multiparameter quantum metrology, one is given n copies of a quantum state ρ_θ parameterized by a vector $\theta \in \mathbb{R}^m$, and the goal is to output an estimator $\hat{\theta}$ of θ . This estimator is *locally unbiased* at a point θ^* if it satisfies

$$(i) \quad \mathbf{E}[\hat{\theta} \mid \rho_{\theta^*}] = \theta^* \quad \text{and} \quad (ii) \quad \left. \frac{\partial}{\partial \theta_i} \mathbf{E}[\hat{\theta} \mid \rho_\theta] \right|_{\theta=\theta^*} = 0.$$

Given a locally unbiased estimator, we can evaluate its performance using the *mean squared error matrix* (MSEM) $V \in \mathbb{R}^{m \times m}$ defined as

$$V_{ij} := \mathbf{E}[(\hat{\theta}_i - \theta_i^*)(\hat{\theta}_j - \theta_j^*) \mid \rho_{\theta^*}], \quad \text{for all } i, j \in [m].$$

The *quantum Cramér–Rao bound* (QCRB) [Hel67] states that the performance of the MSEM can always be lower-bounded via $V \succeq \mathcal{F}^{-1}$, where $\mathcal{F}$ is a particular matrix known as the *Quantum Fisher Information* (QFI) matrix. A definition of this matrix can be found, for example, in [PSW25, Section 1.6].

Recently, Zhou and Chen gave a generic method for converting unbiased estimators for mixed state tomography into locally unbiased estimators for multiparameter quantum metrology [ZC25]. Plugging the debiased Keyl’s algorithm into this transformation, Pelecanos, Spilecki, and Wright gave locally unbiased estimator with the following guarantee. Writing V_n for the MSEM of their matrix when given n copies of ρ_θ , it satisfies $n \cdot V_n \rightarrow 2 \cdot \mathcal{F}^{-1}$ in the limit as $n \rightarrow \infty$ [PSW25, Theorem 1.10]. In other words, their estimator achieves twice the QCRB asymptotically. This is optimal, as there are examples of parameterized quantum states in which the factor of 2 is necessary [DDGG20, Section 3.1.1].

The proof of [PSW25] uses the second moment formula for the debiased Keyl’s algorithm, and as a result, we can also achieve twice the QCRB asymptotically by plugging our estimator $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$ into the Zhou and Chen transformation. Their proof requires two properties of $\text{Lower}_{\rho_\theta}$. First, they need that for any observables O_1 and O_2 , $\text{tr}(O_1 \otimes O_2 \cdot \text{Lower}_{\rho_\theta}) = \text{tr}(O_2 \otimes O_1 \cdot \text{Lower}_{\rho_\theta})$. This holds in our case too, as our $\text{Lower}_{\rho_\theta}$ is a nonnegative linear combination of terms of the form $X \otimes X$, where X is a Hermitian matrix, and

$$\text{tr}(O_1 \otimes O_2 \cdot X \otimes X) = \text{tr}(O_1 \cdot X) \cdot \text{tr}(O_2 \cdot X) = \text{tr}(O_2 \otimes O_1 \cdot X \otimes X).$$

In addition, they need that for any matrix Q , $\text{tr}(Q^\dagger \otimes Q \cdot \text{Lower}_{\rho_\theta}) \geq 0$. This too holds in our case, as

$$\text{tr}(Q^\dagger \otimes Q \cdot X \otimes X) = \text{tr}(Q^\dagger \cdot X) \cdot \text{tr}(Q \cdot X) = \overline{\text{tr}(Q \cdot X)} \cdot \text{tr}(Q \cdot X) = |\text{tr}(Q \cdot X)|^2 \geq 0,$$

where the third equality uses the fact that X is Hermitian. (See [PSW25, Proof of Theorem 8.1] for both of these required properties.) Note that for this result, we care about the asymptotic regime of $n \rightarrow \infty$, in which case it actually suffices to use the simpler algorithm $\mathbf{Mix}(\mathcal{A}_{\text{GPS}})$.

1.3 Discussion

Our results demonstrate that the purification channel provides a powerful tool for designing mixed state tomography algorithms. We believe that it will have more applications in the future, and we conclude with some open directions along these lines.

The purification channel. We still feel like we lack a deep understanding of the purification channel. For example, currently we only know how to implement it using Schur transforms; could this channel be implemented in a more elementary way, without using representation theory? Similarly, for some of our applications, we must apply the purification channel to a state which has larger rank than the purification channel “supports”. Is there a natural interpretation of the behavior of the channel in this case?

Other applications of the purification channel. Are there more applications of the mixed state to pure state reduction beyond the ones we considered in this work? One possibility is the shadow tomography problem: an immediate consequence of our reduction is that mixed state shadow tomography generically reduces to pure state shadow tomography. Could this help design better shadow tomography algorithms in the future? More generally, are there additional applications of the purification channel, even beyond quantum learning?

Efficient algorithms for Hayashi’s algorithm. Although we now have an efficient quantum algorithm for vanilla mixed state tomography due to [Theorem 1.3](#), we still do not have efficient algorithms for any of our other applications, as these all require performing Hayashi’s algorithm. This motivates the following question: can Hayashi’s algorithm be made efficient? We believe the answer is yes, and we leave this question to future work.

The debiased Keyl’s algorithm. In the course of this work, we have resolved a number of open problems stated in the debiased Keyl’s algorithm paper (admittedly much faster than we were expecting), namely open problems 1 (learning with high probability), 4 (efficient algorithms), and 6 (simpler proofs of the variance formula), except we resolved these questions for our algorithm $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ rather than for the debiased Keyl’s algorithm. Can any of our techniques help us give simpler proofs for the debiased Keyl’s algorithm? More concretely, does purification give us a simpler perspective to understand the measurements that the debiased Keyl’s algorithm performs? Also, are there any tomographic tasks which separate the performance of these two algorithms, or do they essentially behave identically for all tasks?

Unentangled measurements for pure state problems. For the task of vanilla mixed state tomography, our [Theorem 1.3](#) showed that after purifying ρ , it suffices to run a pure state tomography algorithm which uses unentangled measurements. On the other hand, our other applications involve entangled measurements across the purified copies. So, is this just a fluke of vanilla tomography? Or could it be that once ρ has been purified, one never requires entangled measurements? In other words, could it be that entanglement is only useful for producing consistent purifications, for all tomographic tasks?

To investigate these questions, let us first explain why we did not use unentangled measurements for our mixed state unbiased estimators. The natural first thing to try would be to plug in the standard unentangled measurement tomography algorithm into our mixed state to pure state reduction, giving the algorithm $\text{Mix}(\mathcal{A}_{\text{standard}})$. To understand how well this performs, let us first look at $\mathcal{A}_{\text{standard}}$. Given n copies of a pure state $\sigma \in \mathbb{C}^{d \times d}$, its output $\hat{\sigma}_{\text{avg}}$ has second moment

$$\mathbf{E}[\hat{\sigma}_{\text{avg}} \otimes \hat{\sigma}_{\text{avg}}] = \frac{n-1}{n} \cdot \sigma \otimes \sigma + \frac{1}{n} \cdot (\sigma \otimes I_d + I_d \otimes \sigma) \cdot \text{SWAP} + \frac{1}{n} \cdot \text{SWAP} - \text{Lower}_{\sigma}.$$

(In fact, this expression also holds when σ is a mixed state.) This is easy to show by direct calculation and we omit the proof. Note that the third term has a $1/n$ factor on the SWAP, whereas the Grier–Pashayan–Schaeffer algorithm has a significantly smaller factor of $1/n^2$ (see [Theorem 1.5](#)). This means that this algorithm is not useful for any of our unbiased estimator applications. The ultimate issue is that although $\mathcal{A}_{\text{standard}}$ has strong ℓ_{∞} guarantees, which are sufficient for achieving optimal pure state tomography bounds, it has relatively weak ℓ_2 guarantees, and these are what many of our other applications rely on.

This could be a fundamental limitation of all algorithms which use unentangled measurements, but it could also be a sign that we are simply using the wrong unentangled pure state tomography algorithm. Perhaps a better algorithm would be to first compute the top eigenvector $|\mathbf{v}\rangle$ of $\hat{\sigma}_{\text{avg}}$ and consider the density matrix $|\mathbf{v}\rangle\langle\mathbf{v}|$. This will certainly be a biased estimator of σ , but one can certainly correct for its bias, much as Grier, Pashayan, and Schaeffer corrected for the bias of Hayashi’s algorithm [[GPS24](#)]. How well does the resulting algorithm perform? We don’t have a clear answer to this question, but there is limited evidence suggesting that it may indeed be a better estimator than $\hat{\sigma}_{\text{avg}}$ on its own. In particular, Grier, Pashayan, and Schaeffer showed that when σ is a pure state, there is an unbiased estimator for σ which is closely related to $\hat{\sigma}_{\text{avg}}^2$, and this unbiased estimator outperforms the unsquared $\hat{\sigma}_{\text{avg}}$ at the classical shadows task, at least in some regimes of parameters. Note that squaring $\hat{\sigma}_{\text{avg}}$ has the effect of putting more weight on its larger eigenvectors and less weight on its smaller eigenvectors, which is a step in the direction of $|\mathbf{v}\rangle\langle\mathbf{v}|$.

We note that there are some quantum learning tasks where entangled measurements help, even when the input states are pure. These include testing if a bipartite pure state is product or entangled [[CCHL22](#), [Har23](#)] and testing if a multipartite pure state has a hidden cut [[BCS⁺25](#)]. As for tasks with a more tomographic flavor, the only separation we are aware of is learning stabilizer states, which can be solved with $O(n)$ copies using entangled measurements but requires $\Omega(n^2)$ copies using unentangled measurements [[ABDY22](#)], though we note that this lower bound proof only seems to apply to algorithms which make nonadaptive measurements. (We thank Sitan Chen for pointing out these applications to us.) However, it seems possible to us that the unentangled pure state tomography algorithm we suggested above could be competitive with Hayashi’s

algorithm for many tomographic tasks. This would of course not yield improved sample complexity bounds, as it seems clear that Hayashi’s algorithm is performing the “right” measurement for pure state tomography, but it would help us understand the philosophical question of when and why entangled measurements help.

1.4 Organization

The rest of this document closely follows the introduction in organization. We begin with relevant preliminaries in [Section 2](#) below. Then, in [Section 3](#), we describe the random purification channel ([Theorem 1.1](#)) and provide the relevant tools for analyzing quasi-purification.

Next, we investigate pure state tomography algorithms in the unentangled and entangled measurement settings in [Sections 4](#) and [5](#), respectively. In the unentangled section, we give a sample-optimal gate-efficient pure state tomography algorithm ([Theorem 1.2](#)), which through our reduction produces a sample-optimal gate-efficient mixed state tomography algorithm ([Theorem 1.3](#)). In the entangled tomography section, we study the Grier–Pashayan–Schaeffer unbiased estimator for pure states and compute its second moment ([Theorem 1.5](#)). Through our reduction, this produces a state-of-the-art unbiased estimator for mixed states ([Theorem 1.9](#)); the proof that this reduction works is the focus of [Section 6](#).

On the way, we prove the intermediate results about Hayashi’s algorithm and the Grier–Pashayan–Schaeffer algorithm discussed in the introduction. We include these, as they give simple proofs of slightly weaker statements: sample-optimal tomography without time-efficiency ([Proposition 1.4](#)) and unbiased estimators which do not use quasi-purification ([Theorem 1.7](#)).

Finally, in [Appendix A](#), we analyze the behavior of $\text{Mix}(\mathcal{A}_{\text{Hayashi}})$, $\text{Mix}(\mathcal{A}_{\text{GPS}})$, and $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$, and find that the measurements they perform are Pretty Good Measurements (PGMs) over the Hilbert–Schmidt measure. This gives a way to describe these algorithms without using the purification channel.

2 Preliminaries

We use **boldface** to denote random variables and **sans serif** to denote registers, i.e. subsystems of a larger system. For example, $\rho = \rho_{A_1 \dots A_n}$ denotes a state on n registers, where we drop the subscript if the registers are clear; we can then denote the partial trace on a subsystem as, for example, $\rho_{A_1} = \text{tr}_{A_2 \dots A_n}(\rho)$. If not otherwise stated, 1 denotes the first subsystem, 2 the second, and so on.

We use $\|X\|_1$ and $\|X\|_\infty$ to denote the trace norm and operator norm (i.e. Schatten 1-norm and Schatten ∞ -norm, respectively). For two quantum states ρ and σ , the fidelity between them is $F(\rho, \sigma) = \|\rho^{1/2} \sigma^{1/2}\|_1^2$, so that, in particular, $F(|u\rangle\langle u|, |v\rangle\langle v|) = |\langle u|v\rangle|^2$. We use $A \preceq B$ to denote PSD ordering, i.e. $B - A$ is positive semidefinite.

Recall the definition of $\text{SoS}(d)$ from [Definition 1.6](#). We observe that this cone is closed under partial trace.

Proposition 2.1 (Partial trace preserves the cone of Hermitian squares). *Given integers d and r , consider the Hilbert space $\mathbb{C}^d \otimes \mathbb{C}^r \cong \mathbb{C}^D$, where $D = d \cdot r$. Call d -dimensional registers A and r -dimensional registers B . Let $M_{A_1 B_1 A_2 B_2} \in \text{SoS}(D)$. Then $\text{tr}_{B_1 B_2}(M) \in \text{SoS}(d)$.*

Proof. Since $M_{A_1 B_1 A_2 B_2} \in \text{SoS}(D)$, it can be written as a positive linear combination of matrices of the form $X_{A_1 B_1} \otimes X_{A_2 B_2}$, where X is a Hermitian matrix acting on $\mathbb{C}^D$. Then $\text{tr}_{B_1 B_2}(M)$ can be written as a positive linear combination of matrices of the form

$$\text{tr}_{B_1 B_2}(X_{A_1 B_1} \otimes X_{A_2 B_2}) = \text{tr}_{B_1}(X_{A_1 B_1}) \otimes \text{tr}_{B_2}(X_{A_2 B_2}).$$

Since $\text{tr}_B(X_{AB})$ is Hermitian, $\text{tr}_{B_1 B_2}(M)$ is in $\text{SoS}(d)$. □

2.1 The symmetric group

We write S_n for the symmetric group on n elements, and e for the identity element of S_n . We sometimes use cycle notation to represent permutations; for example, (i, j) is the transposition swapping elements $i, j \in [n]$, and $(1, 2)(2, 3) = (1, 2, 3)$.

For $m < n$, the subgroup S_m naturally embeds into S_n by associating the permutation $\sigma \in S_m$ with the permutation $\sigma' \in S_n$ defined by taking $\sigma'(i) = \sigma(i)$ when $i \leq m$ and taking $\sigma'(i) = i$ otherwise. The

corresponding *group algebra* of S_n consists of all linear combinations of symmetric group elements $\sum_{\pi \in S_n} \alpha_\pi \cdot \pi$ with coefficients $\alpha_\pi \in \mathbb{C}$. An especially important subset of the symmetric group algebra are the *Jucys–Murphy elements*, given by

$$X_1 = 0, \quad \text{and} \quad X_i := (1, i) + \cdots + (i-1, i), \quad \text{for } 2 \leq i \leq n.$$

For us, their significance arises from the following formula, which relates them to the uniform sum over all permutations.

Proposition 2.2 (Product of Jucys–Murphy elements). *Let $n \geq 1$. Then*

$$(e + X_n) \cdots (e + X_1) = \sum_{\pi \in S_n} \pi.$$

Proof. We prove this by induction on n . The $n = 1$ base case is trivial. For the inductive step, let us assume that this statement is true for $n - 1$, i.e.

$$(e + X_{n-1}) \cdots (e + X_1) = \sum_{\sigma \in S_{n-1}} \sigma.$$

Then our goal is to show that

$$\begin{aligned} \sum_{\pi \in S_n} \pi &= (e + X_n) \cdots (e + X_1) = (e + X_n) \cdot \sum_{\sigma \in S_{n-1}} \sigma \\ &= \sum_{\sigma \in S_{n-1}} \sigma + \sum_{i=1}^{n-1} \sum_{\sigma \in S_{n-1}} (i, n) \cdot \sigma. \end{aligned} \tag{7}$$

To prove this, we will show that every $\pi \in S_n$ occurs as a summand in Equation (7). First, suppose that $\pi(n) = n$. Then $\pi \in S_{n-1}$ as well, and so it can be found in the first sum in Equation (7). Otherwise, suppose $\pi(n) = i$ for some $1 \leq i \leq n-1$. Then $(i, n) \cdot \pi$, fixes n , and so we can write $(i, n) \cdot \pi = \sigma$, for some $\sigma \in S_{n-1}$. But then $\pi = (i, n) \cdot \sigma$, which can be found in the second sum in Equation (7). Thus, every $\pi \in S_n$ occurs as a summand in Equation (7). As there are exactly $(n-1)! + (n-1) \cdot (n-1)! = n!$ terms in this equation, each must correspond to a unique permutation in S_n , completing the proof. $\square$

Given $\pi \in S_n$, $P_d(\pi)$ is defined to be the unitary operation which acts on $(\mathbb{C}^d)^{\otimes n}$ by permuting the n registers via the equation

$$P_d(\pi) \cdot |i_1, \dots, i_n\rangle = |i_{\pi^{-1}(1)}, \dots, i_{\pi^{-1}(n)}\rangle, \quad \text{for all } i_1, \dots, i_n \in [d].$$

When d is clear from context, we will often write this as $P(\pi)$ for simplicity. We may sometimes even drop the notation $P(\cdot)$ altogether and write this simply as π . When $n = 2$, we will often write $\text{SWAP} := P((1, 2))$.

We will commonly encounter having two sets of n registers: n registers of Hilbert space $\mathbb{C}^d$, named $A_1, \dots, A_n$ and n registers of Hilbert space $\mathbb{C}^r$ named $B_1, \dots, B_n$. Pairing these up gives n registers of Hilbert space $(\mathbb{C}^d \otimes \mathbb{C}^r) \cong \mathbb{C}^D$, for $D = d \cdot r$, named $A_1 B_1, \dots, A_n B_n$. Given a permutation $\pi \in S_n$, we will write:

- either $P_d(\pi)$ or $P_A(\pi)$ for the permutation matrix acting on $A_1, \dots, A_n$,
- either $P_r(\pi)$ or $P_B(\pi)$ for the permutation matrix acting on $B_1, \dots, B_n$, and
- either $P_D(\pi)$ or $P_{AB}(\pi)$ for the permutation matrix acting on $A_1 B_1, \dots, A_n B_n$.

When we are in the above situation with $n = 2$, we will apply similar notational choices to the SWAP matrix (so that we write SWAP_d or SWAP_A for the SWAP matrix acting on $A_1 A_2$, and so on). We will abuse notation and use e.g. SWAP_A to refer to both the operator on A alone, along with the operator $\text{SWAP}_A \otimes I_B$.

The operator $P_{AB}(\pi)$ decomposes across the A and B registers in the natural way.

Proposition 2.3. *Let $A_1, \dots, A_n$ be registers with Hilbert space $\mathbb{C}^d$ and $B_1, \dots, B_n$ be registers with Hilbert space $\mathbb{C}^r$. Then for all $\pi \in S_n$,*

$$P_{AB}(\pi) = P_A(\pi) \otimes P_B(\pi).$$

Proof. Let $\pi \in S_n$. Then for all $a_1, \dots, a_n \in [d]$ and $b_1, \dots, b_n \in [r]$,

$$\begin{aligned}
& P_{AB}(\pi) \cdot \left(|a_1\rangle_{A_1} |b_1\rangle_{B_1} \otimes \dots \otimes |a_n\rangle_{A_n} |b_n\rangle_{B_n} \right) \\
&= \left(|a_{\pi^{-1}(1)}\rangle_{A_1} |b_{\pi^{-1}(1)}\rangle_{B_1} \right) \otimes \dots \otimes \left(|a_{\pi^{-1}(n)}\rangle_{A_n} |b_{\pi^{-1}(n)}\rangle_{B_n} \right) \\
&= \left(|a_{\pi^{-1}(1)}\rangle_{A_1} \otimes \dots \otimes |a_{\pi^{-1}(n)}\rangle_{A_n} \right) \otimes \left(|b_{\pi^{-1}(1)}\rangle_{B_1} \otimes \dots \otimes |b_{\pi^{-1}(n)}\rangle_{B_n} \right) \\
&= \left(P_A(\pi) \cdot |a_1\rangle_{A_1} \otimes \dots \otimes |a_n\rangle_{A_n} \right) \otimes \left(P_B(\pi) \cdot |b_1\rangle_{B_1} \otimes \dots \otimes |b_n\rangle_{B_n} \right) \\
&= P_A(\pi) \otimes P_B(\pi) \cdot \left(|a_1\rangle_{A_1} |b_1\rangle_{B_1} \otimes \dots \otimes |a_n\rangle_{A_n} |b_n\rangle_{B_n} \right).
\end{aligned}$$

This completes the proof. $\square$

Our moment computations will feature partial traces of permuted operators very heavily. Simplifying these expressions is straightforward, either through manipulating them algebraically or inspecting their tensor network diagrams. We will now prove a couple of helper propositions to illustrate how this can be done. Later, we will perform similar computations, and their proofs will follow similarly.

Proposition 2.4. *Let $M = M_{AB}$ act on $\mathbb{C}^d \otimes \mathbb{C}^r$. Then*

$$\text{tr}_A(\text{SWAP}_{AC} \cdot (M_{AB} \otimes (I_d)_C)) = M_{CB} = \text{tr}_A((M_{AB} \otimes (I_d)_C) \cdot \text{SWAP}_{AC}).$$

Proof. We verify that every entry of $\text{tr}_A(\text{SWAP}_{AC} \cdot (M_{AB} \otimes (I_d)_C))$ equals the corresponding entry of M_{CB} : for arbitrary $b, b' \in [r]$ and $c, c' \in [d]$, we have

$$\begin{aligned}
\langle bc|_{BC} \text{tr}_A(\text{SWAP}_{AC} \cdot (M_{AB} \otimes (I_d)_C)) |b'c'\rangle_{BC} &= \sum_{a=1}^d \langle abc|_{ABC} \cdot \text{SWAP}_{AC} \cdot (M_{AB} \otimes (I_d)_C) \cdot |ab'c'\rangle_{ABC} \\
&= \sum_{a=1}^d \langle cba|_{ABC} \cdot (M_{AB} \otimes (I_d)_C) \cdot |ab'c'\rangle_{ABC} \\
&= \sum_{a=1}^d \langle cb|M|ab'\rangle \cdot \langle a|I_d|c'\rangle = \langle cb|M|c'b'\rangle = \langle bc|_{BC} M_{CB} |b'c'\rangle_{BC}.
\end{aligned}$$

The analogous statement for $\text{tr}_A((M_{AB} \otimes (I_d)_C) \cdot \text{SWAP}_{AC})$ follows by taking the conjugate transpose of the computation above. $\square$

We will occasionally make use of [Proposition 2.4](#) under the following guise, and so we record it here.

Proposition 2.5. *Let M act on $A_1 B_1$. Then*

$$\text{tr}_{B_1 B_2}(M_{A_1 B_1} \otimes I_{A_2 B_2} \cdot \text{SWAP}_{AB}) = (\text{tr}_{B_1}(M_{A_1 B_1}) \otimes I_{A_2}) \cdot \text{SWAP}_A.$$

Proof. Recall that SWAP_{AB} swaps both A_1 with A_2 and B_1 with B_2 . We calculate:

$$\begin{aligned}
\text{tr}_{B_1 B_2}(M_{A_1 B_1} \otimes I_{A_2 B_2} \cdot \text{SWAP}_{AB}) &= \text{tr}_{B_1 B_2}(M_{A_1 B_1} \otimes I_{A_2 B_2} \cdot \text{SWAP}_A \cdot \text{SWAP}_B) && \text{(by Proposition 2.3)} \\
&= \text{tr}_{B_1 B_2}(M_{A_1 B_1} \otimes I_{A_2 B_2} \cdot \text{SWAP}_B) \cdot \text{SWAP}_A \\
&= \text{tr}_{B_2} \left(\text{tr}_{B_1}(M_{A_1 B_1} \otimes I_{A_2 B_2} \cdot \text{SWAP}_B) \right) \cdot \text{SWAP}_A \\
&= \text{tr}_{B_2} \left(M_{A_1 B_2} \otimes I_{A_2} \right) \cdot \text{SWAP}_A && \text{(by Proposition 2.4)} \\
&= (\text{tr}_{B_2}(M_{A_1 B_2}) \otimes I_{A_2}) \cdot \text{SWAP}_A.
\end{aligned}$$

This completes the proof. $\square$

2.2 The symmetric subspace

Now, we will recall standard facts about the symmetric subspace. These can be found, for example, in the survey by Harrow [Har13]. The *symmetric subspace* on $(\mathbb{C}^d)^{\otimes n}$ is given by

$$\vee^n \mathbb{C}^d := \{|\psi\rangle \in (\mathbb{C}^d)^{\otimes n} \mid \forall \pi \in S_n, P(\pi) \cdot |\psi\rangle = |\psi\rangle\}.$$

It can be equivalently written as

$$\vee^n \mathbb{C}^d = \text{span}\{|\psi\rangle^{\otimes n} \mid |\psi\rangle \in \mathbb{C}^d\}.$$

We will denote by $d[n]$ the dimension of $\vee^n \mathbb{C}^d$, which is given by the formula

$$d[n] = \binom{n+d-1}{n}.$$

The ratio of successive dimensions satisfies the formula

$$\frac{d[n]}{d[n+1]} = \frac{n+1}{n+d}. \quad (8)$$

Finally, we write $\Pi_{\text{sym}}^{n,d}$ for the projector onto $\vee^n \mathbb{C}^d$. It can be computed by the formula

$$\Pi_{\text{sym}}^{n,d} = \mathbf{E}_{\pi \sim S_n} [P(\pi)] = d[n] \cdot \mathbf{E}_{|\mathbf{u}\rangle \sim \text{Haar}} |\mathbf{u}\rangle \langle \mathbf{u}|^{\otimes n}. \quad (9)$$

We will sometimes drop either n or d from the notation $\Pi_{\text{sym}}^{n,d}$ when they are clear from context. The projector onto the symmetric subspace obeys the following nice recurrence relation.

Proposition 2.6 (Symmetric subspace projector recurrence). *Let $n \geq 2$ and $1 \leq m \leq n$. Then*

$$\Pi_{\text{sym}}^{n,d} = \left(\frac{e + X_n}{n}\right) \cdots \left(\frac{e + X_{m+1}}{m+1}\right) \cdot \Pi_{\text{sym}}^{m,d} \otimes (I_d)^{\otimes (n-m)}.$$

Proof. By Equation (9),

$$\begin{aligned} \Pi_{\text{sym}}^{n,d} &= \mathbf{E}_{\pi \sim S_n} [\pi] = \frac{1}{n!} \cdot \sum_{\pi \in S_n} \pi = \frac{1}{n!} \cdot (e + X_n) \cdots (e + X_1) \\ &= \frac{m!}{n!} \cdot (e + X_n) \cdots (e + X_{m+1}) \cdot \frac{1}{m!} \sum_{\sigma \in S_m} \sigma \\ &= \frac{m!}{n!} \cdot (e + X_n) \cdots (e + X_{m+1}) \cdot \mathbf{E}_{\sigma \sim S_m} [\sigma] \\ &= \frac{m!}{n!} \cdot (e + X_n) \cdots (e + X_{m+1}) \cdot \Pi_{\text{sym}}^{m,d} \otimes (I_d)^{\otimes (n-m)}, \end{aligned}$$

where the third and the fourth equalities used Proposition 2.2. This completes the proof. $\square$

2.3 Representation theory

To analyze the random purification channel (Section 3) and its subsequent use in analyzing the unbiased estimator (Section 6.2 and Appendix A), we will need some standard facts from representation theory. These facts are not necessary for the rest of the document. For a more thorough treatment of these topics, see [Wri16].

Partitions and Young diagrams. A *partition* of n , denoted $\lambda \vdash n$, is a tuple of integers $\lambda = (\lambda_1, \dots, \lambda_k)$ such that $\lambda_1 \geq \dots \geq \lambda_k \geq 0$ and $\lambda_1 + \dots + \lambda_k = n$. The *length* of λ , denoted $\ell(\lambda)$, is equal to the number of nonzero components λ_i . Partitions are typically represented pictorially using *Young diagrams*, which consist of boxes arranged into rows of length λ_1 through λ_k . A *standard Young tableau (SYT)* S of shape λ is a Young diagram of shape λ in which each box has been filled in with a number from $[n]$, with the restriction

that the numbers in each row must be strictly increasing from left-to-right and the numbers in each column must be strictly increasing from top-to-bottom. A *semistandard Young tableau (SSYT)* T of shape λ and alphabet $[d]$ is a Young diagram of shape λ in which each box has been filled in with a number from $[d]$, with the restriction that the numbers in each row must be weakly increasing from left-to-right and the numbers in each column must be strictly increasing from top-to-bottom. We illustrate these concepts in [Figure 7](#).

1	2	5	7
3	6		
4			

1	1	1	3
2	2		
3			

Figure 7: Examples of tableaux of shape $\lambda = (4, 2, 1)$. Left: an SYT. Right: an SSYT, for $d \geq 3$.

Representation theory of the symmetric group. The irreducible representations of the symmetric group are indexed by partitions $\lambda \vdash n$ and are written $(\kappa_\lambda, \text{Sp}_\lambda)$, where Sp_λ is known as the *Specht module*. There is a convenient choice of basis for these irreducible representations known as *Young's orthogonal basis*, which gives rise to an explicit choice of the matrices $\kappa_\lambda(\pi)$ known as *Young's orthogonal representation*. For the irreducible representation corresponding to the Young diagram λ , Young's orthogonal basis has a basis vector $|S\rangle$ for each standard Young tableau S of shape λ . These vectors form an orthonormal basis of Sp_λ . Furthermore, when written in this basis, the matrices $\kappa_\lambda(\pi)$ have only real-valued entries. For shorthand, we write the dimension of the λ -irrep as $\dim(\lambda)$.

Representation theory of the general linear group. The *general linear group* $GL(d)$ is the group consisting of all invertible matrices M in $\mathbb{C}^{d \times d}$. The polynomial irreducible representations of the general linear group are indexed by partitions λ with $\ell(\lambda) \leq d$ and are written $(\nu_\lambda^d, V_\lambda^d)$, where V_λ^d is known as the *Weyl module*. Here, “polynomial” means that the matrix entries of each representation $\nu_\lambda^d(M)$ can be expressed as a polynomial in the matrix entries of M . There is an orthonormal basis of the Weyl module V_λ^d known as the *Gelfand-Tsetlin basis*, in which the basis vectors $|T\rangle$ are indexed by SSYT of shape λ and alphabet $[d]$. An important subgroup of $GL(d)$ is the unitary group $U(d)$. The irreducible representations $(\nu_\lambda^d, V_\lambda^d)$ of $GL(d)$ also form the irreducible representations of $U(d)$.

We will also sometimes want to plug in density matrices $\rho \in \mathbb{C}^{d \times d}$ into the irreducible representation ν_λ^d . When ρ has positive eigenvalues, $\nu_\lambda^d(\rho)$ is clearly well-defined as such a ρ is in $GL(d)$. But this is also well-defined even when ρ has some eigenvalues which are zero, by remembering that the matrix entries of $\nu_\lambda^d(\rho)$ are polynomials in the matrix entries of ρ and are therefore always well-defined. Equivalently, by continuity, we can view $\nu_\lambda^d(\rho)$ as the limit of $\nu_\lambda^d(\rho^+)$ for a sequence of matrices ρ^+ with positive eigenvalues which approach ρ in the limit.

Schur–Weyl duality. Given $\pi \in S_n$, let us recall the representation $P(\pi)$ of S_n which acts on $(\mathbb{C}^d)^{\otimes n}$ as follows:

$$P(\pi) \cdot |i_1, \dots, i_n\rangle = |i_{\pi^{-1}(1)}, \dots, i_{\pi^{-1}(n)}\rangle, \quad \text{for all } i_1, \dots, i_n \in [d].$$

Given $M \in GL(d)$, we can also define a representation $Q^d(M)$ of $GL(d)$ which acts on $(\mathbb{C}^d)^{\otimes n}$, as follows:

$$Q^d(M) \cdot |i_1, \dots, i_n\rangle = (M \cdot |i_1\rangle) \otimes \dots \otimes (M \cdot |i_n\rangle).$$

For all $\pi \in S_n$ and $M \in GL(d)$, we have that $P(\pi) \cdot Q^d(M) = Q^d(M) \cdot P(\pi)$, and so these representations commute with each other. As a result, there is a change of basis in which these representations are simultaneously block-diagonalized into their irreducible representations. The precise form of this simultaneous block-diagonalization is provided by *Schur–Weyl duality*, which states that there is a unitary U_{Schur}^d known as the *Schur transform* such that, for all $\pi \in S_n$ and $M \in GL(d)$,

$$U_{\text{Schur}}^d \cdot P(\pi) Q^d(M) \cdot (U_{\text{Schur}}^d)^\dagger = \sum_{\lambda \vdash n, \ell(\lambda) \leq d} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\pi) \otimes \nu_\lambda^d(M).$$

An efficient algorithm for computing the Schur transform was provided by Bacon, Chuang, and Harrow [BCH05]. By the footnote on page 160 of Harrow's Ph.D. thesis [Har05], it can be computed to ε accuracy in diamond distance in time $\text{poly}(n, \log(d), \log(1/\varepsilon))$. The fact that this efficient algorithm produces the Young orthogonal basis on the symmetric group register was recently shown by Pelecanos, Spilecki, and Wright [PSW25, Appendix A].

While in the Schur basis, it is natural to measure which irrep λ one is in. Doing so is known as *weak Schur sampling* and involves performing the projective measurement $\{\Pi_\lambda\}_{\lambda \vdash n, \ell(\lambda) \leq d}$ in which

$$\Pi_\lambda := |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes I_{\dim(V_\lambda^d)}.$$

3 The random purification channel

In this section, we formally define the random purification channel $\Phi_{\text{Purify}}^{d,r}$. Then, we reprove [Theorem 1.1](#), closely following the treatment given in [TWZ25, Section 2.3]. Finally, we prove some additional facts needed to understand the Mix^+ reduction, which we use later in [Section 6.2](#).

The purification channel $\Phi_{\text{Purify}}^{d,r}$ converts the state $\rho^{\otimes n}$ into the mixture $\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}$. To do so, it is crucial to understand what these two mixed states look like in the Schur basis. For the former, this is easy, as $\rho^{\otimes n}$ is just the state $Q^d(\rho)$. As a result, we can apply Schur–Weyl duality, which states that

$$U_{\text{Schur}}^d \cdot \rho^{\otimes n} \cdot (U_{\text{Schur}}^d)^\dagger = \sum_{\lambda \vdash n, \ell(\lambda) \leq d} |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda^d(\rho). \quad (10)$$

For the latter, this is a bit more challenging. To begin, let us set up some notation. Each state $|\rho\rangle$ lives inside $\mathbb{C}^d \otimes \mathbb{C}^r \cong \mathbb{C}^D$, where $D = d \cdot r$. The n d -dimensional registers will be denoted $A_1, \dots, A_n$, and the n r -dimensional registers will be denoted $B_1, \dots, B_n$, though these will often be dropped when clear from context. We will consider what $\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}$ looks like when we apply a Schur transform U_{Schur}^d to the n different $\mathbb{C}^d$ registers and a second Schur transform U_{Schur}^r to the n different $\mathbb{C}^r$ registers. We will abbreviate $U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r$ as $U_{\text{Schur}}^{\otimes 2}$. The first Schur transform will produce a basis of the form $|\lambda\rangle_Y \otimes |S\rangle_P \otimes |T\rangle_Q$, where Y is the Young diagram register, P is the symmetric group register, and Q is the unitary group register. The second Schur transform will produce a basis of the form $|\lambda'\rangle_{Y'} \otimes |S'\rangle_{P'} \otimes |T'\rangle_{Q'}$.

Our first observation is that $|\rho\rangle\langle\rho|^{\otimes n}$ is contained inside the symmetric subspace $\vee^n \mathbb{C}^D$, and so $\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}$ is in $\vee^n \mathbb{C}^D$ as well. Towards understanding the random purification channel, we now turn to characterizing $\vee^n \mathbb{C}^D$.

3.1 The symmetric subspace across two registers

We start by considering the projector onto the symmetric subspace of $(\mathbb{C}^d \otimes \mathbb{C}^r)^{\otimes n}$ in the Schur basis. [TWZ25, Lemma 2.13] gives a convenient formula for this projector, which we restate and reprove. We will need the following definition.

Definition 3.1 (Specht module EPR state). Let $\lambda \vdash n$. We write $|\text{EPR}_\lambda\rangle$ for the pure state inside $\text{Sp}_\lambda \otimes \text{Sp}_\lambda$ given by

$$|\text{EPR}_\lambda\rangle := \frac{1}{\sqrt{\dim(\lambda)}} \cdot \sum_S |S\rangle \otimes |S\rangle,$$

where the sum ranges over all SYTs of shape λ .

Lemma 3.2 (The projector onto the symmetric subspace in the Schur basis).

$$U_{\text{Schur}}^{\otimes 2} \cdot \Pi_{\text{sym}}^{n,D} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\lambda\rangle\langle\lambda\lambda|_{YY'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{PP'} \otimes I_{QQ'}. \quad (11)$$

Proof. By [Proposition 2.3](#),

$$\Pi_{\text{sym}}^{n,D} = \mathbf{E}_{\pi \sim S_n} [P_{AB}(\pi)] = \mathbf{E}_{\pi \sim S_n} [P_A(\pi) \otimes P_B(\pi)].$$

If we now Schur transform A and B, we obtain:

$$\begin{aligned}
& (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r) \cdot \Pi_{\text{sym}}^{n,D} \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^\dagger \\
&= \mathbf{E}_{\pi \sim S_n} \left[(U_{\text{Schur}}^d \cdot P_A(\pi) \cdot (U_{\text{Schur}}^d)^\dagger) \otimes (U_{\text{Schur}}^r \cdot P_B(\pi) \cdot (U_{\text{Schur}}^r)^\dagger) \right] \\
&= \mathbf{E}_{\pi \sim S_n} \left[\left(\sum_{\lambda \vdash n, \ell(\lambda) \leq d} |\lambda\rangle\langle\lambda|_{\mathcal{Y}} \otimes \kappa_\lambda(\pi)_{\mathcal{P}} \otimes (I_{\dim(V_\lambda^d)})_{\mathcal{Q}} \right) \otimes \left(\sum_{\mu \vdash n, \ell(\mu) \leq r} |\mu\rangle\langle\mu|_{\mathcal{Y}'} \otimes \kappa_\mu(\pi)_{\mathcal{P}'} \otimes (I_{\dim(V_\mu^r)})_{\mathcal{Q}'} \right) \right] \\
&= \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq d \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\rangle\langle\lambda|_{\mathcal{Y}} \otimes |\mu\rangle\langle\mu|_{\mathcal{Y}'} \otimes \mathbf{E}_{\pi \sim S_n} [\kappa_\lambda(\pi)_{\mathcal{P}} \otimes \kappa_\mu(\pi)_{\mathcal{P}'}] \otimes (I_{\dim(V_\lambda^d)})_{\mathcal{Q}} \otimes (I_{\dim(V_\mu^r)})_{\mathcal{Q}'}. \tag{12}
\end{aligned}$$

We can simplify the operator acting on $\mathcal{P}\mathcal{P}'$ using the grand orthogonality relations:

$$\begin{aligned}
\mathbf{E}_{\pi \sim S_n} [\kappa_\lambda(\pi)_{\mathcal{P}} \otimes \kappa_\mu(\pi)_{\mathcal{P}'}] &= \sum_{S, S', S'', S'''} \mathbf{E}_{\pi \sim S_n} [\kappa_\lambda(\pi)_{S, S''} \cdot \kappa_\mu(\pi)_{S', S'''}] \cdot |S\rangle\langle S''|_{\mathcal{P}} \otimes |S'\rangle\langle S'''|_{\mathcal{P}'} \\
&= \sum_{S, S', S'', S'''} \mathbf{E}_{\pi \sim S_n} [\overline{\kappa_\lambda(\pi)_{S, S''}} \cdot \kappa_\mu(\pi)_{S', S'''}] \cdot |S\rangle\langle S''|_{\mathcal{P}} \otimes |S'\rangle\langle S'''|_{\mathcal{P}'} \\
&= \sum_{S, S', S'', S'''} \frac{1}{\dim(\lambda)} \cdot \delta_{\lambda, \mu} \cdot \delta_{S, S'} \cdot \delta_{S'', S'''} \cdot |S\rangle\langle S''|_{\mathcal{P}} \otimes |S'\rangle\langle S'''|_{\mathcal{P}'} \\
&= \delta_{\lambda, \mu} \cdot |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{\mathcal{P}\mathcal{P}'}.
\end{aligned}$$

In the second equality, we have also used the fact that Young's orthogonal representation is real-valued. Plugging this back into Equation (12) yields Equation (11). $\square$

For any pure state $|u\rangle \in \mathbb{C}^D$, the state $|u\rangle\langle u|^{\otimes n} \in \mathcal{V}^n \mathbb{C}^D$. It will be useful for us to understand what this state looks like in the Schur basis as well.

Lemma 3.3 (Formula for $|u\rangle\langle u|^{\otimes n}$ in the Schur basis). *Let $|u\rangle_{A_1 B_1} \in \mathbb{C}^d \otimes \mathbb{C}^r \cong \mathbb{C}^D$. Then*

$$U_{\text{Schur}}^{\otimes 2} \cdot |u\rangle\langle u|_{AB}^{\otimes n} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger = \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq r \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\lambda\rangle\langle\mu\mu|_{\mathcal{Y}\mathcal{Y}'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{\mathcal{P}\mathcal{P}'} \otimes |u_{\lambda\lambda}\rangle\langle u_{\mu\mu}|_{\mathcal{Q}\mathcal{Q}'}, \tag{13}$$

where $|u_{\lambda\lambda}\rangle_{\mathcal{Q}\mathcal{Q}'}$ is some (unnormalized) vector in $V_\lambda^d \otimes V_\lambda^r$. Moreover, for each λ ,

$$\text{tr}_{\mathcal{Q}'}(|u_{\lambda\lambda}\rangle\langle u_{\lambda\lambda}|_{\mathcal{Q}\mathcal{Q}'}) = \dim(\lambda) \cdot \nu_\lambda^d(\sigma_u)_{\mathcal{Q}}, \tag{14}$$

where $\sigma_u = \text{tr}_{B_1}(|u\rangle\langle u|_{A_1 B_1})$.

Proof. Since $|u\rangle_{AB}^{\otimes n} \in \mathcal{V}^n \mathbb{C}^D$, we have $\Pi_{\text{sym}}^{n,D} \cdot |u\rangle_{AB}^{\otimes n} = |u\rangle_{AB}^{\otimes n}$, and thus

$$\begin{aligned}
U_{\text{Schur}}^{\otimes 2} \cdot |u\rangle_{AB}^{\otimes n} &= \left(U_{\text{Schur}}^{\otimes 2} \cdot \Pi_{\text{sym}}^{n,D} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger \right) \cdot \left(U_{\text{Schur}}^{\otimes 2} \cdot |u\rangle_{AB}^{\otimes n} \right) \\
&= \left(\sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\lambda\rangle\langle\lambda\lambda|_{\mathcal{Y}\mathcal{Y}'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{\mathcal{P}\mathcal{P}'} \otimes I_{\mathcal{Q}\mathcal{Q}'} \right) \cdot \left(U_{\text{Schur}}^{\otimes 2} \cdot |u\rangle_{AB}^{\otimes n} \right) \quad (\text{Lemma 3.2}) \\
&= \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\lambda\rangle_{\mathcal{Y}\mathcal{Y}'} \otimes |\text{EPR}_\lambda\rangle_{\mathcal{P}\mathcal{P}'} \otimes |u_{\lambda\lambda}\rangle_{\mathcal{Q}\mathcal{Q}'},
\end{aligned}$$

where $|u_{\lambda\lambda}\rangle_{\mathcal{Q}\mathcal{Q}'}$ is some (unnormalized) vector in $V_\lambda^d \otimes V_\lambda^r$. Thus,

$$U_{\text{Schur}}^{\otimes 2} \cdot |u\rangle\langle u|_{AB}^{\otimes n} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger = \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq r \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\lambda\rangle\langle\mu\mu|_{\mathcal{Y}\mathcal{Y}'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{\mathcal{P}\mathcal{P}'} \otimes |u_{\lambda\lambda}\rangle\langle u_{\mu\mu}|_{\mathcal{Q}\mathcal{Q}'}.$$

This proves Equation (13). To show Equation (14), we first note that

$$\begin{aligned} U_{\text{Schur}}^d \cdot \text{tr}_{\mathbf{B}}(|u\rangle\langle u|^{\otimes n}) \cdot (U_{\text{Schur}}^d)^\dagger &= U_{\text{Schur}}^d \cdot \sigma_u^{\otimes n} \cdot (U_{\text{Schur}}^d)^\dagger \\ &= \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes (I_{\dim(\lambda)}|_{\mathbf{P}} \otimes \nu_\lambda^d(\sigma_u)|_{\mathbf{Q}}). \end{aligned} \quad (15)$$

However, we also have

$$\begin{aligned} U_{\text{Schur}}^d \cdot \text{tr}_{\mathbf{B}}(|u\rangle\langle u|^{\otimes n}) \cdot (U_{\text{Schur}}^d)^\dagger &= \text{tr}_{\mathbf{B}} \left((U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r) \cdot |u\rangle\langle u|^{\otimes n} \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^\dagger \right) \\ &= \text{tr}_{\mathbf{B}} \left(\sum_{\lambda, \mu} |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{\mathbf{P}} \otimes |u_{\lambda\lambda}\rangle\langle u_{\mu\mu}|_{\mathbf{Q}} \right) \\ &= \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda| \otimes \left(\frac{I_{\dim(\lambda)}}{\dim(\lambda)} \right) \otimes \text{tr}_{\mathbf{Q}'}(|u_{\lambda\lambda}\rangle\langle u_{\lambda\lambda}|_{\mathbf{Q}\mathbf{Q}'}). \end{aligned} \quad (16)$$

Equation (14) then follows by comparing the matrices acting on $\mathbf{Q}$ in Equations (15) and (16). $\square$

3.2 A formula for a many-copy random purification

We now apply the results from the previous section to the random purification channel. We start by understanding the output of the purification channel, $\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}$, in the Schur basis, using the fact that this state is in the symmetric subspace. [TWZ25, Lemma 2.16] gives us the following formula.

Lemma 3.4 (Random purification formula). *Let $\rho_{\mathbf{A}_1} \in \mathbb{C}^{d \times d}$ be a mixed state, and let $|\rho_0\rangle_{\mathbf{A}_1\mathbf{B}_1} \in \mathbb{C}^d \otimes \mathbb{C}^r \cong \mathbb{C}^D$ be a fixed purification of ρ . Recall that a random purification $|\rho\rangle_{\mathbf{A}_1\mathbf{B}_1}$ is obtained from $|\rho_0\rangle$ by applying a Haar random unitary $\mathbf{U} \in U(r)$ to the purifying register, i.e. $|\rho\rangle_{\mathbf{A}_1\mathbf{B}_1} = (I_d \otimes \mathbf{U}) \cdot |\rho_0\rangle_{\mathbf{A}_1\mathbf{B}_1}$. Then*

$$U_{\text{Schur}}^{\otimes 2} \cdot \left(\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n} \right) \cdot (U_{\text{Schur}}^\dagger)^{\otimes 2} = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} \dim(\lambda) \cdot |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{\mathbf{P}} \otimes \nu_\lambda^d(\rho)_{\mathbf{Q}} \otimes \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{\mathbf{Q}'}.$$

Proof of Lemma 3.4. First, we have

$$\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n} = \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\mathbf{U}_{\mathbf{B}}^{\otimes n} \cdot |\rho_0\rangle\langle\rho_0|_{\mathbf{AB}}^{\otimes n} \cdot (\mathbf{U}_{\mathbf{B}}^{\otimes n})^\dagger].$$

We will now rewrite everything in the Schur basis. Lemma 3.3 shows us how to rewrite $|\rho_0\rangle\langle\rho_0|^{\otimes n}$:

$$U_{\text{Schur}}^{\otimes 2} \cdot |\rho_0\rangle\langle\rho_0|_{\mathbf{AB}}^{\otimes n} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger = \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq r \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{\mathbf{P}} \otimes |\rho_{0,\lambda\lambda}\rangle\langle\rho_{0,\mu\mu}|_{\mathbf{Q}\mathbf{Q}'},$$

for some unnormalized vectors $\{|\rho_{0,\lambda\lambda}\rangle\}_\lambda$. Conjugating this state by a Haar random unitary yields, in the Schur basis,

$$\begin{aligned} &U_{\text{Schur}}^{\otimes 2} \cdot \left(\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n} \right) \cdot (U_{\text{Schur}}^\dagger)^{\otimes 2} \\ &= \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq r \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{\mathbf{P}} \otimes \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\nu_\lambda^r(\mathbf{U})_{\mathbf{Q}'} \cdot |\rho_{0,\lambda\lambda}\rangle\langle\rho_{0,\mu\mu}|_{\mathbf{Q}\mathbf{Q}'} \cdot \nu_\mu^r(\mathbf{U})_{\mathbf{Q}'}^\dagger]. \end{aligned} \quad (17)$$

Let us focus on the operator in the unitary registers. This operator is a linear combination of terms of the form

$$\mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\nu_\lambda^r(\mathbf{U})_{\mathbf{Q}'} \cdot |T\rangle\langle T''|_{\mathbf{Q}} \otimes |T'\rangle\langle T'''|_{\mathbf{Q}'} \cdot \nu_\mu^r(\mathbf{U})_{\mathbf{Q}'}^\dagger] = |T\rangle\langle T''|_{\mathbf{Q}} \otimes \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\nu_\lambda^r(\mathbf{U}) \cdot |T'\rangle\langle T'''| \cdot \nu_\mu^r(\mathbf{U})^\dagger]_{\mathbf{Q}'}.$$

Here, T, T'' are SSYTs of shape λ and μ , respectively, with alphabet $[d]$. Meanwhile, T', T''' are SSYTs of shape λ and μ , respectively, with alphabet $[r]$. The operator

$$\mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\nu_\lambda^r(\mathbf{U}) \cdot |T'\rangle\langle T'''| \cdot \nu_\mu^r(\mathbf{U})^\dagger]_{\mathbf{Q}'}$$

is an intertwiner for ν_λ^r and ν_μ^r , and so by Schur's lemma is nonzero only when $\lambda = \mu$. In this case, we have

$$\begin{aligned} \mathbf{E}_{U \sim \text{Haar}} [\nu_\lambda^r(U) \cdot |T'\rangle\langle T''| \cdot \nu_\lambda^r(U)^\dagger]_{Q'} &= \text{tr} \left(\mathbf{E}_{U \sim \text{Haar}} [\nu_\lambda^r(U) \cdot |T'\rangle\langle T''| \cdot \nu_\lambda^r(U)^\dagger] \right) \cdot \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'} \\ &= \text{tr}(|T'\rangle\langle T''|) \cdot \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'}. \end{aligned}$$

Thus, twirling by a Haar random unitary maps $|T\rangle\langle T''|_Q \otimes |T'\rangle\langle T''|_{Q'}$ to

$$\delta_{\lambda,\mu} \cdot |T\rangle\langle T''|_Q \otimes \text{tr}(|T'\rangle\langle T''|) \cdot \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'} = \delta_{\lambda,\mu} \cdot \text{tr}_{Q'}(|T\rangle\langle T''|_Q \otimes |T'\rangle\langle T''|_{Q'}) \otimes \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'}.$$

Substituting this back into Equation (17) then gives:

$$(17) = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_{YY'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{PP'} \otimes \text{tr}_{Q'}(|\rho_{0,\lambda\lambda}\rangle\langle\rho_{0,\lambda\lambda}|_{QQ'}) \otimes \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'}.$$

Finally, by Lemma 3.3, we have $\text{tr}_{Q'}(|\rho_{0,\lambda\lambda}\rangle\langle\rho_{0,\lambda\lambda}|_{QQ'}) = \dim(\lambda) \cdot \nu_\lambda^d(\rho)_Q$. Plugging this into the above expression yields:

$$U_{\text{Schur}}^{\otimes 2} \cdot \left(\mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|_{AB}^{\otimes n} \right) \cdot (U_{\text{Schur}}^\dagger)^{\otimes 2} = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} \dim(\lambda) \cdot |\lambda\rangle\langle\lambda|_{YY'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\lambda|_{PP'} \otimes \nu_\lambda^d(\rho)_Q \otimes \left(\frac{I_{\dim(V_\lambda^r)}}{\dim(V_\lambda^r)} \right)_{Q'}.$$

This completes the proof. $\square$

3.3 The random purification channel

We now define the random purification channel $\Phi_{\text{Purify}}^{d,r}$.

Given n copies of ρ :

1. Apply the Schur transform U_{Schur}^d to $\rho^{\otimes n}$.
2. Perform weak Schur sampling. Letting λ be the outcome, the state collapses to

$$\rho|_\lambda := |\lambda\rangle\langle\lambda|_Y \otimes \left(\frac{I_{\dim(\lambda)}}{\dim(\lambda)} \right)_P \otimes \left(\frac{\nu_\lambda^d(\rho)}{s_\lambda^d(\rho)} \right)_Q.$$

3. Introduce a new register Y' and copy λ into it.
4. Introduce a new register P' . Discard the contents of P and place the state $|\text{EPR}_\lambda\rangle_{PP'}$ into these registers.
5. Introduce a new register Q' initialized to the maximally mixed state $I_{\dim(V_\lambda^r)}/\dim(V_\lambda^r)$.
6. Apply the inverse Schur transform $(U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^\dagger$ and output the resulting state.

Figure 8: The random purification channel $\Phi_{\text{Purify}}^{d,r}$.

By [TWZ25, Lemma 2.11], the random purification channel satisfies

$$\Phi_{\text{Purify}}^{d,r}(\rho^{\otimes n}) = \mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}.$$

Furthermore, it can be computed to δ error in time $\text{poly}(n, \log(d), \log(1/\delta))$. We now reprove these results, for completeness.

Proof of Theorem 1.1. To show $\Phi_{\text{Purify}}^{d,r}(\rho^{\otimes n})$ prepares a random purification, we track how our input state changes with each applied step. After Schur transforming, we have the state

$$\sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_{\mathbf{Y}} \otimes (I_{\dim(\lambda)})_{\mathbf{P}} \otimes \nu_{\lambda}^d(\rho)_{\mathbf{Q}}.$$

After weak Schur sampling, we obtain outcome $\lambda \vdash n$ with probability $\dim(\lambda) \cdot s_{\lambda}^d(\rho)$, and state $\rho|_{\lambda}$. Conditioned on λ , following steps 3–5, we obtain the state

$$|\lambda\lambda\rangle\langle\lambda\lambda|_{\mathbf{Y}\mathbf{Y}'} \otimes |\text{EPR}_{\lambda}\rangle\langle\text{EPR}_{\lambda}|_{\mathbf{P}\mathbf{P}'} \otimes \left(\frac{\nu_{\lambda}^d(\rho)}{s_{\lambda}^d(\rho)}\right)_{\mathbf{Q}} \otimes \left(\frac{I_{\dim(V_{\lambda}^r)}}{\dim(V_{\lambda}^r)}\right)_{\mathbf{Q}'}$$

Averaged over the outcome of weak Schur sampling, we have prepared the state

$$\sum_{\lambda \vdash n, \ell(\lambda) \leq r} \dim(\lambda) \cdot |\lambda\lambda\rangle\langle\lambda\lambda|_{\mathbf{Y}\mathbf{Y}'} \otimes |\text{EPR}_{\lambda}\rangle\langle\text{EPR}_{\lambda}|_{\mathbf{P}\mathbf{P}'} \otimes \nu_{\lambda}^d(\rho)_{\mathbf{Q}} \otimes \left(\frac{I_{\dim(V_{\lambda}^r)}}{\dim(V_{\lambda}^r)}\right)_{\mathbf{Q}'}$$

Thus, by Lemma 3.4, the output after the sixth step is

$$\Phi_{\text{Purify}}^{d,r}(\rho^{\otimes n}) = \mathbf{E}_{|\rho\rangle} |\rho\rangle\langle\rho|^{\otimes n}.$$

We now turn to efficiency. The gate complexity is dominated by the cost of the two Schur transforms, each of which can be computed to δ accuracy in diamond distance in time $\text{poly}(n, \log(d), \log(1/\delta))$ [Har05]. $\square$

Having defined the random purification channel, there are a couple of properties of it that we need to establish for our mixed state tomography algorithm. We begin by looking at its application conditioned on the Young diagram λ .

Notation 3.5. We will slightly abuse notation by allowing the input to $\Phi_{\text{Purify}}^{d,r}$ to be a state expressed in the Schur basis. In particular, when the input is $\rho|_{\lambda}$, we have

$$\Phi_{\text{Purify}}^{d,r}(\rho|_{\lambda}) = |\lambda\lambda\rangle\langle\lambda\lambda|_{\mathbf{Y}\mathbf{Y}'} \otimes |\text{EPR}_{\lambda}\rangle\langle\text{EPR}_{\lambda}|_{\mathbf{P}\mathbf{P}'} \otimes \left(\frac{\nu_{\lambda}^d(\rho)}{s_{\lambda}^d(\rho)}\right)_{\mathbf{Q}} \otimes \left(\frac{I_{\dim(V_{\lambda}^r)}}{\dim(V_{\lambda}^r)}\right)_{\mathbf{Q}'}$$

This will be convenient for us in the context of quasi-purification, where we will want to apply $\Phi_{\text{Purify}}^{d,\ell(\lambda)}$, conditioned on the outcome λ observed in weak Schur sampling.

We now turn to a couple of properties of the random purification channel that we need to establish for our mixed state tomography algorithm.

Lemma 3.6 (Partial trace of the purification channel).

$$\text{tr}_{\mathbf{Y}'\mathbf{P}'\mathbf{Q}'}(\Phi_{\text{Purify}}^{d,r}(\rho|_{\lambda})) = \rho|_{\lambda}.$$

We will also use the following formula, which shows that weak Schur sampling, when applied to $\rho^{\otimes n}$, is non-destructive, in that it does not perturb the state.

Lemma 3.7 (Average of $\rho|_{\lambda}$). $\mathbf{E}_{\lambda}[\rho|_{\lambda}] = U_{\text{Schur}}^d \cdot \rho^{\otimes n} \cdot (U_{\text{Schur}}^d)^{\dagger}$.

Proof. By Equation (10),

$$U_{\text{Schur}}^d \cdot \rho^{\otimes n} \cdot (U_{\text{Schur}}^d)^{\dagger} = \sum_{\lambda \vdash n, \ell(\lambda) \leq d} \dim(\lambda) \cdot s_{\lambda}(\rho) \cdot \rho|_{\lambda}.$$

For each λ , $\rho|_{\lambda}$ is a density matrix which lives inside the λ -irrep space. Therefore, $\mathbf{Pr}[\lambda = \lambda] = \dim(\lambda) \cdot s_{\lambda}(\rho)$. As a result,

$$\mathbf{E}_{\lambda}[\rho|_{\lambda}] = \sum_{\lambda \vdash n, \ell(\lambda) \leq d} \mathbf{Pr}[\lambda = \lambda] \cdot \rho|_{\lambda} = U_{\text{Schur}}^d \cdot \rho^{\otimes n} \cdot (U_{\text{Schur}}^d)^{\dagger}. \quad \square$$

4 Pure state tomography with unentangled measurements

In this section, we prove [Theorem 1.2](#), establishing the existence of a gate-efficient and sample-optimal pure state tomography algorithm. To the best of our knowledge, no result combining these guarantees has been explicitly stated in the literature before, although the special case when $\delta = 1/\text{poly}(d)$ appears in [\[HKOT23, Appendix C\]](#).

Throughout this section, we will assume that our pure state is represented by a system of qubits, which allows us to discuss the gate complexity of our algorithm. A consequence of this is that the overall dimension of our state d will always be a power of two. This is essentially without loss of generality: any mixed state ρ whose dimension d is not a power of two may be embedded into a Hilbert space of $\lceil \log_2 d \rceil$ qubits. This incurs only constant-factor loss, since the resulting Hilbert space has dimension $2^{\lceil \log_2 d \rceil} \leq 2d$, and since our algorithms will scale polynomially in d in both sample and gate complexity.

Our starting point is the sample-optimal pure state tomography algorithm given by Guta, Kahn, Kueng, and Tropp [\[GKKT20\]](#), which we discussed in [Section 1.1](#). Their algorithm uses the uniform POVM $\{d \cdot |u\rangle\langle u| \cdot du\}$, which is continuous and therefore cannot be implemented exactly. To resolve this, we will replace the uniform POVM with a measurement in a random basis drawn from an (approximate) t -design. These measurements can be implemented efficiently, and the resulting algorithm performs just as well so long as our application only requires using the first t moments of the uniform POVM.

The standard tomography algorithm from [Figure 2](#) performs full state tomography when given $n = O(d^3)$ copies of a generic mixed state $\sigma \in \mathbb{C}^{d \times d}$. It is well-known that this bound only requires using the first two moments of the uniform POVM (see, for example, [\[Wri16, Section 5.1\]](#)), and so it is common in the literature to see the uniform POVM in this algorithm replaced with a t -design for $t = 2$ (or $t = 3$, as is required for some applications). In our case, however, two features of our algorithm will require us to use t -designs with a much higher value of t . First, we want an improved sample complexity of $n = O(d)$ in the case when σ is a pure state, and this will require us to take $t = \Omega(\log(d))$ (as was done previously in [\[HKOT23\]](#)). Second, we want our sample complexity to have the correct dependence on the parameter δ , and this will turn out to require taking $t = \Theta(d)$. We use the following definition for an approximate unitary t -design.

Definition 4.1 (Approximate unitary t -design). A distribution $\mathcal{P}$ over $U(d)$ is an ε -approximate unitary t -design if the two mixed unitary channels

$$\mathcal{C}_t : \eta \rightarrow \mathbf{E}_{V \sim \mathcal{P}} [V^{\otimes t} \cdot \eta \cdot V^{\dagger, \otimes t}] \quad \text{and} \quad \mathcal{H}_t : \eta \rightarrow \mathbf{E}_{U \sim \text{Haar}} [U^{\otimes t} \cdot \eta \cdot U^{\dagger, \otimes t}]$$

satisfy

$$(1 - \varepsilon) \cdot \mathcal{H}_t(\eta) \preceq \mathcal{C}_t(\eta) \preceq (1 + \varepsilon) \cdot \mathcal{H}_t(\eta), \quad \text{for all PSD } \eta.$$

There is a rich literature on unitary t -designs, with many constructions offering different tradeoffs between efficiency, random seed length, and approximation. For the purposes of this section, we will use the construction from O'Donnell, Servedio, and Paredes [\[OSP23\]](#), which provides the following guarantee.

Theorem 4.2 (Existence of approximate unitary t -designs). *Let d be a power of two. For any t , there exists a distribution $\mathcal{P}$ over $U(d)$ such that:*

1. $\mathcal{P}$ is a $\frac{1}{2}$ -approximate unitary t -design.
2. Each unitary in the distribution can be constructed using a $\log_2(d)$ -qubit circuit that consists of $\text{poly}(\log_2(d) \cdot t)$ gates from a fixed discrete gate set.
3. Computing the circuit that implements a unitary from $\mathcal{P}$ takes $\text{poly}(\log_2(d) \cdot t)$ classical processing.

Proof. The proof follows by combining [\[OSP23, Theorem 2.11\]](#) with the proof of [\[HLT24, Corollary 1.2\]](#). We include an argument below for completeness.

A first observation is that it suffices to give a distribution over the unitary subgroup $SU(d)$ that satisfies the constraints of [Theorem 4.2](#). This is because any unitary U can be written as the product of an element U_0 of $SU(d)$ with a complex number z of unit norm. Then

$$U^{\otimes t} \cdot \eta \cdot U^{\dagger, \otimes t} = z^t \cdot U_0^{\otimes t} \cdot \eta \cdot U_0^{\dagger, \otimes t} \cdot \bar{z}^t = U_0^{\otimes t} \cdot \eta \cdot U_0^{\dagger, \otimes t}.$$

In particular, [OSP23, Theorem 2.11] implies that there exists some distribution $\mathcal{P}$ over $SU(d)$ that satisfies Items 2 and 3 of Theorem 4.2 and

$$\left\| \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} [\mathbf{V}^{\otimes t} \otimes \overline{\mathbf{V}}^{\otimes t}] - \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\mathbf{U}^{\otimes t} \otimes \overline{\mathbf{U}}^{\otimes t}] \right\|_{\infty} \leq \frac{1}{2} \cdot d^{-3t}.$$

This bound on the operator norm implies a bound on the diamond distance between the channels $\mathcal{C}_t$ and $\mathcal{H}_t$. Formally, Low in [Low10, Lemma 2.2.14, Item 4] gives the following inequality:

$$\|\mathcal{C}_t - \mathcal{H}_t\|_{\diamond} \leq d^t \cdot \left\| \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} [\mathbf{V}^{\otimes t} \otimes \overline{\mathbf{V}}^{\otimes t}] - \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\mathbf{U}^{\otimes t} \otimes \overline{\mathbf{U}}^{\otimes t}] \right\|_{\infty} \leq \frac{1}{2} \cdot d^{-2t}.$$

We conclude that $\mathcal{P}$ is a $\frac{1}{2}$ -approximate unitary t -design using [BHH16, Lemma 3], which states that if the channels $\mathcal{C}_t$ and $\mathcal{H}_t$ satisfy

$$\|\mathcal{C}_t - \mathcal{H}_t\|_{\diamond} \leq \varepsilon \cdot d^{-2t},$$

then $\mathcal{P}$ is an ε -approximate unitary t -design. $\square$

Given n copies of $|v\rangle\langle v| \in \mathbb{C}^{d \times d}$:

1. Let $t = 6d + 1$.
2. For each $1 \leq i \leq n$:
 - (a) Sample $\mathbf{V}$ from an approximate unitary t -design $\mathcal{P}$ that satisfies Theorem 4.2.
 - (b) Rotate the i -th copy of $|v\rangle\langle v|$ to $\mathbf{V} |v\rangle\langle v| \mathbf{V}^{\dagger}$ and measure in the computational basis to obtain the outcome $|j\rangle$ for $j \in [d]$.
 - (c) Set $\hat{\rho}_i = (d+1) \cdot \mathbf{V}^{\dagger} |j\rangle\langle j| \mathbf{V} - I_d$.
3. Let $\hat{\rho}_{\text{avg}} = \frac{1}{n} \cdot (\hat{\rho}_1 + \dots + \hat{\rho}_n)$.
4. Output $|\hat{v}\rangle\langle \hat{v}|$, where $|\hat{v}\rangle$ is the top eigenvector of $\hat{\rho}_{\text{avg}}$.

Figure 9: The [GKKT20] unentangled measurement tomography algorithm for pure states, made time-efficient.

We now analyze the algorithm of Guta, Kahn, Kueng, and Tropp [GKKT20], except with the uniform POVM measurement replaced with the approximate unitary t -design from Theorem 4.2. See Figure 9 for this algorithm, and see the discussion around Figure 2 for the original algorithm. We will follow their proof almost identically. In particular, we first prove a bound on the k -th moment.

Lemma 4.3 (Moment bounds of the time-efficient version of [GKKT20]). *Let $t = 6d + 1$, and let $\mathcal{P}$ be the approximate t -design guaranteed by Theorem 4.2. For $\mathbf{V}$ sampled from $\mathcal{P}$ and $j \in [d]$ being the outcome we obtain when we measure $\mathbf{V} |v\rangle\langle v| \mathbf{V}^{\dagger}$ in the computational basis, define*

$$\hat{\rho} = (d+1) \cdot \mathbf{V}^{\dagger} |j\rangle\langle j| \mathbf{V} - I_d.$$

It then holds for all $|z\rangle$ and all integers $k \geq 2$ that

$$\mathbf{E} [|\langle z | (\hat{\rho} - |v\rangle\langle v|) | z \rangle|^k] \leq 41 \cdot 6^{k-2} k!. \quad (18)$$

Proof. Our first observation is that

$$|\langle z | (\hat{\rho} - |v\rangle\langle v|) | z \rangle|^k \leq \|\hat{\rho} - |v\rangle\langle v|\|_{\infty}^k \leq (d+1)^k.$$

Whenever $k \geq 6d$, the Stirling approximation formula implies that $(d+1)^k \leq (k/e)^k \leq k! \leq 41 \cdot 6^{k-2} k!$, which means that the desired bound holds. Therefore, we will assume that $k \leq 6d$ for the remainder of this proof.

We define $|\mathbf{u}\rangle = \mathbf{V}^\dagger |j\rangle$ for notational brevity. Our starting point is an observation of [GKKT20] that for a fixed $|z\rangle$, one can write

$$\begin{aligned} \mathbf{s}_z &= \langle z | (\hat{\rho} - |v\rangle\langle v|) | z \rangle \\ &= \langle z | ((d+1)|\mathbf{u}\rangle\langle\mathbf{u}| - I_d - |v\rangle\langle v|) | z \rangle \\ &= (d+1) \langle \mathbf{u} | z \rangle \langle z | \mathbf{u} \rangle - (1 + |\langle v | z \rangle|^2) \\ &= (d+1) \langle \mathbf{u} | B | \mathbf{u} \rangle, \end{aligned}$$

for B equal to $|z\rangle\langle z| - \frac{1+|\langle v|z\rangle|^2}{d+1} \cdot I_d$. This allows us to upper bound $\mathbf{E}[|\mathbf{s}_z|^k]$ by

$$\mathbf{E}[|\mathbf{s}_z|^k] = \mathbf{E}[(d+1) \langle \mathbf{u} | B | \mathbf{u} \rangle]^k = (d+1)^k \mathbf{E}[\text{tr}(|\mathbf{u}\rangle\langle\mathbf{u}| \cdot B)]^k \leq (d+1)^k \mathbf{E}[\text{tr}(|\mathbf{u}\rangle\langle\mathbf{u}| \cdot |B|)]^k.$$

Here, $|B| = \sqrt{B^\dagger B}$ stands for the PSD matrix obtained from B by taking the absolute value of its eigenvalues. We now expand $|\mathbf{u}\rangle\langle\mathbf{u}|$ using $\mathbf{V}$ to obtain

$$\begin{aligned} \mathbf{E}[|\mathbf{s}_z|^k] &\leq (d+1)^k \mathbf{E}[\text{tr}(|\mathbf{u}\rangle\langle\mathbf{u}| \cdot |B|)]^k \\ &= (d+1)^k \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} \left[\sum_{j=1}^d \text{tr}(\mathbf{V}^\dagger |j\rangle\langle j| \mathbf{V} \cdot |B|) \right]^k \cdot \text{tr}(|v\rangle\langle v| \cdot \mathbf{V}^\dagger |j\rangle\langle j| \mathbf{V}) \\ &= (d+1)^k \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} \left[\sum_{j=1}^d \text{tr}(\mathbf{V}^{\dagger, \otimes k+1} \cdot |j\rangle\langle j|^{\otimes k+1} \cdot \mathbf{V}^{\otimes k+1} \cdot |B|^{\otimes k} \otimes |v\rangle\langle v|) \right] \\ &= (d+1)^k \text{tr} \left(\sum_{j=1}^d \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} [\mathbf{V}^{\dagger, \otimes k+1} \cdot |j\rangle\langle j|^{\otimes k+1} \cdot \mathbf{V}^{\otimes k+1}] \cdot |B|^{\otimes k} \otimes |v\rangle\langle v| \right). \end{aligned}$$

Let us use A to denote the first factor in the trace:

$$A = \sum_{j=1}^d \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} [\mathbf{V}^{\dagger, \otimes k+1} \cdot |j\rangle\langle j|^{\otimes k+1} \cdot \mathbf{V}^{\otimes k+1}].$$

We bound the expression above using Hölder's inequality:

$$\mathbf{E}[|\mathbf{s}_z|^k] \leq (d+1)^k \text{tr}(A \cdot |B|^{\otimes k} \otimes |v\rangle\langle v|) \leq (d+1)^k \cdot \|A\|_\infty \cdot \| |B|^{\otimes k} \otimes |v\rangle\langle v| \|_1 \leq (d+1)^k \cdot \|A\|_\infty \cdot \| |B| \|_1^k.$$

The trace norm of $|B|$ is at most

$$\| |B| \|_1 = \|B\|_1 \leq 1 + \frac{d}{d+1} \cdot (1 + |\langle v | z \rangle|^2) \leq 3.$$

We proceed to bounding the operator norm of A using the t -design property. Since we are in the $k \leq 6d$ regime, then for $t = 6d + 1$, the random unitary $\mathbf{V}$ is sampled from an approximate t -design such that $t \geq k + 1$. In particular:

$$\begin{aligned} \|A\|_\infty &\leq \sum_{j=1}^d \left\| \mathbf{E}_{\mathbf{V} \sim \mathcal{P}} [\mathbf{V}^{\dagger, \otimes k+1} \cdot |j\rangle\langle j|^{\otimes k+1} \cdot \mathbf{V}^{\otimes k+1}] \right\|_\infty \\ &\leq \sum_{j=1}^d \left(1 + \frac{1}{2} \right) \cdot \left\| \mathbf{E}_{\mathbf{U} \sim \text{Haar}} [\mathbf{U}^{\dagger, \otimes k+1} \cdot |j\rangle\langle j|^{\otimes k+1} \cdot \mathbf{U}^{\otimes k+1}] \right\|_\infty \quad (\text{Theorem 4.2}) \\ &= \sum_{j=1}^d \frac{3}{2} \cdot \left\| \frac{1}{d[k+1]} \cdot \Pi_{\text{sym}}^{k+1, d} \right\|_\infty \\ &= \frac{3d}{2 \cdot d[k+1]}. \end{aligned}$$

Recall that

$$d[k+1] = \binom{d+k}{k+1} = \frac{(d+k)\dots(d+1)d}{(k+1)!} \geq (d+1)^k \cdot \frac{d}{(k+1)!}.$$

Hence, we conclude that

$$\mathbf{E}[|s_z|^k] \leq (d+1)^k \cdot \frac{3d}{2 \cdot d[k+1]} \cdot 3^k \leq \frac{3}{2} \cdot 3^k \cdot (k+1)! \leq \frac{3}{2} \cdot 3^k \cdot (3 \cdot 2^{k-2} \cdot k!) \leq 41 \cdot 6^{k-2} k!.$$

In the third inequality we used the fact that $(k+1)! = (k+1) \cdot k! \leq 3 \cdot 2^{k-2} \cdot k!$, whenever $k \geq 2$. $\square$

Given the above moment bounds, we are ready to prove a concentration result for the operator norm of the difference of our estimator from $|v\rangle\langle v|$. For this, we will need to use a covering net argument, which we define below.

Definition 4.4 (Covering net). Let T be a set with metric $d(\cdot, \cdot)$. A subset S of T is a *covering net with radius θ* if every point x in T is at most θ away from some point in the covering net, i.e., there exists $y \in S$ such that $d(x, y) \leq \theta$.

Lemma 4.5 (Concentration bounds of the time-efficient version of [GKKT20]). *It holds that*

$$\Pr \left[\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \geq q \right] \leq 2 \cdot \exp \left((4 \log 3)d - \frac{nq^2}{704} \right).$$

Proof. We first rewrite the operator norm as

$$\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} = \max_{|z\rangle} |\langle z | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) | z \rangle|.$$

We will use a covering net argument to compute this maximum over all pure states. It is well-known that the set of d -dimensional pure states admits a covering net, since each state corresponds, via a distance-preserving mapping, to a point on the real Euclidean sphere $\mathbb{S}^{2d-1}$. In particular, for any covering net $S = \{|z_i\rangle\}_{i \in [N]}$ of radius $\theta \in [0, 1/2)$, it holds that (see, for example, [Ver18, Lemma 4.4.2]):

$$\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \leq \frac{1}{1-2\theta} \cdot \max_{i \in [N]} |\langle z_i | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) | z_i \rangle|. \quad (19)$$

Moreover, [Ver18, Corollary 4.2.11] implies that such a covering net with radius θ has size at most $(1+2/\theta)^{2d}$. We set $\theta = 1/4$, which means that $N = |S| \leq 9^{2d} = 3^{4d}$.

For any fixed $|z\rangle$, the right-hand side of Equation (19) can be written as a sum of n independent and identically distributed copies of a mean-zero random variable:

$$|\langle z | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) | z \rangle| = \left| \frac{1}{n} \sum_{i=1}^n \langle z | (\hat{\rho}_i - |v\rangle\langle v|) | z \rangle \right|.$$

In Lemma 4.3 we showed that the k -th moment of each such variable is bounded above by $41 \cdot 6^{k-2} k!$. This bound allows us to use the Bernstein inequality, whose statement is given below.

Theorem 4.6 (Bernstein inequality [FR13, Theorem 7.30]). *Consider n independent samples $s_1, \dots, s_n$ of the real random variable s , which has mean zero, and whose k -th moment is bounded by*

$$\mathbf{E}[|s|^k] \leq k! R^{k-2} \sigma^2 / 2,$$

for all integers $k \geq 2$, and some positive constants R, σ^2 . Then, for all $q > 0$,

$$\Pr \left[\left| \sum_{i=1}^n s_i \right| \geq q \right] \leq 2 \exp \left(-\frac{q^2/2}{n\sigma^2 + Rq} \right).$$

Plugging in $R = 6$ and $\sigma^2 = 82$ in the inequality above implies that for all $|z\rangle$ and $0 \leq q \leq 1$,

$$\Pr \left[|\langle z | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) |z\rangle| \geq q \right] \leq 2 \exp \left(-\frac{q^2 n^2 / 2}{82n + 6qn} \right) \leq 2 \exp \left(-\frac{nq^2}{176} \right). \quad (20)$$

Applying the union bound inequality over the covering net S of the pure states, we conclude that

$$\Pr \left[\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \geq q \right] \leq \Pr \left[\max_{i \in [N]} |\langle z_i | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) |z_i\rangle| \geq \frac{q}{2} \right] \quad (\text{Equation (19)})$$

$$\begin{aligned} &= \Pr \left[\bigvee_{i \in [N]} |\langle z_i | (\hat{\rho}_{\text{avg}} - |v\rangle\langle v|) |z_i\rangle| \geq \frac{q}{2} \right] \\ &\leq 2N \cdot \exp \left(-\frac{nq^2}{704} \right) \quad (\text{Equation (20)}) \\ &\leq 2 \cdot \exp \left((4 \log 3)d - \frac{nq^2}{704} \right). \end{aligned}$$

This completes the proof. $\square$

Theorem 4.7 (Theorem 1.2, restated). *There is an algorithm which, given*

$$n = O \left(\frac{d + \log(1/\delta)}{\varepsilon} \right)$$

copies of a pure state $|v\rangle \in \mathbb{C}^d$, outputs a pure state $|\hat{v}\rangle \in \mathbb{C}^d$ such that $|\langle \hat{v} | v \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$. Furthermore, this algorithm can be implemented in $\text{poly}(n)$ time and performs independent measurements across the copies of $|v\rangle$.

Proof. The algorithm is outlined in Figure 9. The claimed running time follows from Theorem 4.2, since the unitary t -design can be implemented using $\text{poly}(\log(d) \cdot t) = \text{poly}(d)$ quantum gates and classical processing. Moreover, the number of samples $n = \text{poly}(d, 1/\varepsilon, \log(1/\delta))$ is also polynomial in these parameters.

It remains to show that the output pure state satisfies $|\langle \hat{v} | v \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$. By setting $q = \sqrt{\varepsilon}$ and the number of samples equal to $n = O((d + \log(1/\delta))/\varepsilon)$ in the statement of Lemma 4.5, we conclude that

$$\Pr \left[\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \geq \sqrt{\varepsilon} \right] \leq \delta.$$

Let us now restrict our attention to the case when $\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \leq \sqrt{\varepsilon}$, which happens with probability at least $1 - \delta$.

Since the operator norm is unitarily invariant, the expression $\|\hat{\rho}_{\text{avg}} - |w\rangle\langle w|\|_{\infty}$ is minimized when $|w\rangle$ is the top eigenvector $|\hat{v}\rangle$ of $\hat{\rho}_{\text{avg}}$. This follows from a theorem of Mirsky (see [HJ13, Corollary 7.4.9.3]), which states that for any two Hermitian matrices A, B :

$$\|A - B\|_{\infty} \geq \|\text{spec}(A) - \text{spec}(B)\|_{\infty},$$

where on the right-hand side, the matrices $\text{spec}(A), \text{spec}(B)$ are diagonal with entries equal to the eigenvalues of A, B in nonincreasing order, respectively. By substituting $\hat{\rho}_{\text{avg}}$ for A and $|v\rangle\langle v|$ for B , we observe that

$$\|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \geq \|\text{spec}(\hat{\rho}_{\text{avg}}) - \text{spec}(|v\rangle\langle v|)\|_{\infty} = \left\| \sum_{i=1}^d \lambda_i |i\rangle\langle i| - |1\rangle\langle 1| \right\|_{\infty} = \|\hat{\rho}_{\text{avg}} - |\hat{v}\rangle\langle \hat{v}|\|_{\infty},$$

where we use λ_i for the i -th largest eigenvalue of $\hat{\rho}_{\text{avg}}$. The last equality follows from the unitary invariance of the norm. As a result, the triangle inequality implies that

$$\| |\hat{v}\rangle\langle \hat{v}| - |v\rangle\langle v| \|_{\infty} \leq \| |\hat{v}\rangle\langle \hat{v}| - \hat{\rho}_{\text{avg}} \|_{\infty} + \|\hat{\rho}_{\text{avg}} - |v\rangle\langle v|\|_{\infty} \leq 2\sqrt{\varepsilon}.$$

Finally, since $|\hat{v}\rangle\langle \hat{v}| - |v\rangle\langle v|$ only has at most two nonzero eigenvalues,

$$\| |\hat{v}\rangle\langle \hat{v}| - |v\rangle\langle v| \|_1 \leq 2 \cdot \| |\hat{v}\rangle\langle \hat{v}| - |v\rangle\langle v| \|_{\infty} \leq 4\sqrt{\varepsilon}.$$

From the relationship between the trace distance and squared overlap for pure states [Wat18, Eq. (1.186)] we conclude that

$$|\langle \hat{v} | v \rangle|^2 = 1 - \left(\frac{1}{2} \| |\hat{v}\rangle\langle \hat{v}| - |v\rangle\langle v| \|_1 \right)^2 \geq 1 - 4\varepsilon.$$

The theorem statement then follows by adjusting the parameter ε by a constant. $\square$

5 Pure state tomography with entangled measurements

In this section, we discuss and analyze two pure state tomography algorithms, both of which use the optimal choice of measurement for this task. The first is Hayashi's algorithm $\mathcal{A}_{\text{Hayashi}}$ [Hay98], which performs this measurement and outputs the naive corresponding estimator, which is itself also a pure state. The second is the algorithm $\mathcal{A}_{\text{GPS}}$ of Grier, Pashayan, and Schaeffer [GPS24], which adjusts Hayashi's estimator in order to remove its bias.

We begin with Hayashi's algorithm, where we directly show that it achieves the optimal sample complexity of $n = O((d + \log(1/\delta))/\varepsilon)$ copies for pure state tomography. This will not be used later on; we include it because $\text{Mix}(\mathcal{A}_{\text{Hayashi}})$ gives arguably the cleanest proof of sample-optimal mixed state tomography, even among results which do not achieve the optimal δ dependence.

We then analyze the moments of the Grier–Pashayan–Schaeffer estimator. In Section 6, we will study $\text{Mix}(\mathcal{A}_{\text{GPS}})$ and $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$; the moments of these mixed state estimators will follow directly from the moments of $\mathcal{A}_{\text{GPS}}$.

5.1 Hayashi's algorithm

Let us first recall Hayashi's algorithm. When performing tomography on n copies of a pure state $|\psi\rangle \in \mathbb{C}^d$, the input $|\psi\rangle^{\otimes n}$ is an element of the symmetric subspace $\vee^n \mathbb{C}^d$. This means that a pure state tomography algorithm's measurement operators need only be specified on the symmetric subspace. Motivated by this, Hayashi [Hay98] introduced the following natural pure state tomography algorithm: simply perform the POVM

$$\{d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du\} \quad (21)$$

and output the pure state $|\mathbf{v}\rangle$ that it returns (Figure 3). This is indeed a valid POVM on the symmetric subspace, as

$$\int_{|u\rangle} d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du = d[n] \cdot \mathbf{E}_{|u\rangle \sim \text{Haar}} |u\rangle\langle u|^{\otimes n} = \Pi_{\text{sym}}^{n,d}$$

due to Equation (9). Hayashi showed that $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability 99% when $n = O(d/\varepsilon)$. To see this, note that

$$\begin{aligned} \mathbf{E} |\langle \mathbf{v} | \psi \rangle|^2 &= d[n] \cdot \int_{|u\rangle} \text{tr} \left(|u\rangle\langle u|^{\otimes n} \cdot |\psi\rangle\langle \psi|^{\otimes n} \right) \cdot |\langle u | \psi \rangle|^2 \cdot du \\ &= d[n] \cdot \int_{|u\rangle} \text{tr} \left(|u\rangle\langle u|^{\otimes n+1} \cdot |\psi\rangle\langle \psi|^{\otimes n+1} \right) \cdot du \\ &= d[n] \cdot \text{tr} \left(\left(\int_{|u\rangle} |u\rangle\langle u|^{\otimes n+1} \cdot du \right) \cdot |\psi\rangle\langle \psi|^{\otimes n+1} \right) \\ &= d[n] \cdot \text{tr} \left(\left(\frac{1}{d[n+1]} \cdot \Pi_{\text{sym}}^{n+1} \right) \cdot |\psi\rangle\langle \psi|^{\otimes n+1} \right) \\ &= \frac{d[n]}{d[n+1]} \\ &= \frac{n+1}{n+d} = 1 - \frac{d-1}{n+d}. \end{aligned} \quad (22)$$

This is $1 - \varepsilon/100$ when $n = O(d/\varepsilon)$. The claim now follows from an application of Markov's inequality.

Next, we upgrade this to a high probability bound for Hayashi's estimator.

Proposition 5.1 (Hayashi's algorithm with high probability; [Proposition 1.4](#), restated). *Given n copies of a pure state $|\psi\rangle \in \mathbb{C}^d$, suppose Hayashi's algorithm returns the state $|\mathbf{v}\rangle$. Then $|\langle \mathbf{v} | \psi \rangle|^2 \geq 1 - \varepsilon$ with probability at least $1 - \delta$ when*

$$n = O\left(\frac{d + \log(1/\delta)}{\varepsilon}\right).$$

Proof. It suffices to show that the random variable $\mathbf{x} := |\langle \mathbf{v} | \psi \rangle|^2$ exhibits strong concentration about its mean. Let us write $p(x)$ for the probability density function of this random variable. Noting that $|\mathbf{v}\rangle = |u\rangle$ with measure

$$d[n] \cdot \text{tr}\left(|u\rangle\langle u|^{\otimes n} \cdot |\psi\rangle\langle\psi|^{\otimes n}\right) \cdot du = d[n] \cdot |\langle u | \psi \rangle|^{2n} \cdot du,$$

it follows that $p(x) = d[n] \cdot x^n \cdot \mu(x)$, where μ is the probability density function of $\mathbf{y} := |\langle \mathbf{u} | \psi \rangle|^2$ when $|\mathbf{u}\rangle \sim \mathbb{C}^d$ is a Haar random vector. It is known that $\mathbf{y}$ is distributed as a $\text{Beta}(1, d-1)$ random variable [[ZS00](#), Eq. (7)], and so $\mu(x) = (d-1)(1-x)^{d-2}$. As a result,

$$p(x) = (d-1) \cdot d[n] \cdot x^n (1-x)^{d-2},$$

and so $\mathbf{x}$ obeys the $\text{Beta}(n+1, d-1)$ distribution. Now we can apply [[Sko23](#), Theorem 1]⁴, which states that so long as $n+1 \geq d-1$,

$$\Pr[\mathbf{x} \leq \mathbf{E}[\mathbf{x}] - \varepsilon] \leq \exp\left(-\frac{\varepsilon^2}{2(v + \frac{c\varepsilon}{3})}\right),$$

where

$$v = \frac{(n+1)(d-1)}{(n+d)^2(n+d+1)} \leq \frac{d}{n^2} \quad \text{and} \quad c = \frac{2(n-d+2)}{(n+d)(n+d+2)} \leq \frac{2}{n}.$$

Hence, we have that

$$\Pr[\mathbf{x} \leq \mathbf{E}[\mathbf{x}] - \varepsilon] \leq \exp\left(-\frac{\varepsilon^2}{2(d/n^2 + (2/3) \cdot \varepsilon/n)}\right),$$

We will apply this bound to the case when $n = O((d + \log(1/\delta))/\varepsilon)$. For this setting of parameters, we have that $d/n^2 \leq (1/3) \cdot \varepsilon/n$; furthermore, we have that $\mathbf{E}[\mathbf{x}] \geq 1 - \varepsilon$ from [Equation \(22\)](#) above. Putting these together, for this range of n ,

$$\Pr[\mathbf{x} \leq (1 - \varepsilon) - \varepsilon] \leq \exp\left(-\frac{\varepsilon^2}{2 \cdot \varepsilon/n}\right) = \exp(-\varepsilon n/2) \leq \delta.$$

This completes the proof. $\square$

This matches the sample complexity bound for the unentangled pure state tomography algorithm $\mathcal{A}_{\text{GKKT}}$ from [Theorem 1.2](#), and so it can be used to give a mixed state tomography algorithm whose sample complexity matches [Theorem 1.3](#). However, unlike $\mathcal{A}_{\text{GKKT}}$, which can be implemented efficiently, we do not know of any efficient implementations of Hayashi's measurement, and so this only produces a mixed state learning algorithm which is sample-optimal but not necessarily efficient.

⁴We are applying [[Sko23](#), Theorem 1] with $\alpha = n+1$ and $\beta = d-1$, and we are interested in the $\alpha \geq \beta$ case of the lower-tail bound. In this case, his result states that

$$\Pr[\mathbf{x} \leq \mathbf{E}[\mathbf{x}] - \varepsilon] \leq \exp\left(-\frac{\varepsilon^2}{2(v + \frac{c\varepsilon}{3})}\right), \tag{23}$$

where

$$v = \frac{\alpha\beta}{(\alpha+\beta)^2(\alpha+\beta+1)} \quad \text{and} \quad c = \frac{2(\beta-\alpha)}{(\alpha+\beta)(\alpha+\beta+2)}.$$

There is actually a slight bug in this statement, however: since $\alpha \geq \beta$, the variable c is actually nonpositive, which means that the $c\varepsilon/3$ in the denominator of [Equation \(23\)](#) is (erroneously) nonpositive. Skorski derived this statement by reducing it to his upper-tail bound, and we can derive the correct lower-tail bound by carrying out this reduction more carefully. In particular, if we set $\mathbf{x}' := 1 - \mathbf{x}$, then $\mathbf{x}'$ is a $\text{Beta}(\beta, \alpha)$ random variable, and so using his upper-tail bound, we obtain

$$\Pr[\mathbf{x} \leq \mathbf{E}[\mathbf{x}] - \varepsilon] = \Pr[\mathbf{x}' \geq \mathbf{E}[\mathbf{x}'] + \varepsilon] \leq \exp\left(-\frac{\varepsilon^2}{2(v - \frac{c\varepsilon}{3})}\right) = \exp\left(-\frac{\varepsilon^2}{2(v + \frac{|c|\varepsilon}{3})}\right).$$

This is the corrected bound that we use above.

5.2 Grier, Pashayan, and Schaeffer's algorithm

One downside of Hayashi's algorithm is that its output $|\mathbf{v}\rangle\langle\mathbf{v}|$ is not an unbiased estimator for the true state $|\psi\rangle\langle\psi|$. This is because $|\psi\rangle\langle\psi|$ is a pure state but the expectation $\mathbf{E} |\mathbf{v}\rangle\langle\mathbf{v}|$ will necessarily be mixed. In addition, as Equation (22) demonstrates, this expectation will have poor overlap with the true state unless $n \gg d$. This limits the usefulness of Hayashi's algorithm to problems such as shadow tomography where one might hope to use significantly fewer than d copies of the state.

To address this, Grier, Pashayan, and Schaeffer [GPS24] noted that one can correct for the bias in Hayashi's algorithm via the simple modification

$$\hat{\sigma}_{\mathbf{v}} := \frac{d+n}{n} \cdot |\mathbf{v}\rangle\langle\mathbf{v}| - \frac{1}{n} \cdot I.$$

In other words, the expectation of this estimator is $\mathbf{E} \hat{\sigma}_{\mathbf{v}} = |\psi\rangle\langle\psi|$, and so it is an unbiased estimator for pure state tomography. They used this to derive bounds on the sample complexity of classical shadows when the input is a pure state.

For our applications, we will need to compute the first and second moments of $\hat{\sigma}_{\mathbf{v}}$. To do so, let us first define the following moment operators.

Definition 5.2 (Moment operator). Given integers $n \geq 1$, $d \geq 1$, and $k \geq 1$, we define the k -th moment operator to be the following operator in $(\mathbb{C}^{d \times d})^{\otimes(n+k)}$:

$$M_{\text{mom}}^{(k)} := d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \hat{\sigma}_u^{\otimes k} \cdot du.$$

Our next proposition shows that the k -th moment operator can be used to calculate the k -th moment of the Grier–Pashayan–Schaeffer algorithm. For a later application, we will state this proposition for the more general case when their algorithm is performed on a generic mixed state ψ_{sym} inside the symmetric subspace. However, for pure state tomography, it suffices to consider the case that $\psi_{\text{sym}} = |\psi\rangle\langle\psi|^{\otimes n}$ for a pure state $|\psi\rangle \in \mathbb{C}^d$.

Proposition 5.3 (Computing moments with the moment operator). Let ψ_{sym} be a mixed state on $\vee^n \mathbb{C}^d$. If $|\mathbf{v}\rangle$ is the output of Hayashi's measurement when performed on ψ_{sym} , then

$$\mathbf{E} \hat{\sigma}_{\mathbf{v}}^{\otimes k} = \text{tr}_{1\dots n}(M_{\text{mom}}^{(k)} \cdot \psi_{\text{sym}} \otimes I_d^{\otimes k}).$$

Proof. This follows by direct computation:

$$\begin{aligned} \text{tr}_{1\dots n}(M_{\text{mom}}^{(k)} \cdot \psi_{\text{sym}} \otimes I_d^{\otimes k}) &= \text{tr}_{1\dots n} \left(\left(d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \hat{\sigma}_u^{\otimes k} \cdot du \right) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes k} \right) \\ &= d[n] \cdot \int_{|u\rangle} \text{tr}(|u\rangle\langle u|^{\otimes n} \cdot \psi_{\text{sym}}) \cdot \hat{\sigma}_u^{\otimes k} \cdot du \\ &= \mathbf{E}_{|\mathbf{v}\rangle} \hat{\sigma}_{\mathbf{v}}^{\otimes k}. \quad \square \end{aligned}$$

Thus, to compute the k -th moment of this unbiased estimator, it suffices to compute the corresponding moment operator. We will do so for the first and second moments.

Lemma 5.4 (Helper lemma). Let $n \geq 1$, $d \geq 2$, and $k \geq 1$. Then

$$d[n] \cdot (d+n)^{\uparrow k} \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+k} \cdot du = (e + X_{n+k}) \cdots (e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes k}),$$

where $a^{\uparrow b} = a(a+1) \cdots (a+b-1)$ is the rising factorial.

Proof. Using $\int_{|u\rangle} |u\rangle\langle u|^{\otimes n+k} \cdot du = d[n+k]^{-1} \cdot \Pi_{\text{sym}}^{n+k}$, this is equivalent to the statement

$$\frac{d[n]}{d[n+k]} \cdot (d+n)^{\uparrow k} \cdot \Pi_{\text{sym}}^{n+k} = (e + X_{n+k}) \cdots (e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes k}).$$

This is true because

$$\begin{aligned}
& \frac{d[n]}{d[n+k]} \cdot (d+n)^{\uparrow k} \cdot \Pi_{\text{sym}}^{n+k} \\
&= \frac{d[n]}{d[n+k]} \cdot (d+n)^{\uparrow k} \cdot \left(\frac{e+X_{n+k}}{n+k} \right) \cdots \left(\frac{e+X_{n+1}}{n+1} \right) \cdot \Pi_{\text{sym}}^n \otimes I^{\otimes k} \quad (\text{by Proposition 2.6}) \\
&= \frac{d[n]}{d[n+k]} \cdot \frac{(d+n)^{\uparrow k}}{(n+1)^{\uparrow k}} \cdot (e+X_{n+k}) \cdots (e+X_{n+1}) \cdot \Pi_{\text{sym}}^n \otimes I^{\otimes k} \\
&= (e+X_{n+k}) \cdots (e+X_{n+1}) \cdot \Pi_{\text{sym}}^n \otimes I^{\otimes k}. \quad (\text{by iterating Equation (8)})
\end{aligned}$$

This completes the proof. $\square$

Lemma 5.5 (The first moment operator). *Given $n \geq 1$ and $d \geq 2$,*

$$M_{\text{mom}}^{(1)} = \frac{1}{n} \cdot X_{n+1} \cdot (\Pi_{\text{sym}}^n \otimes I).$$

Proof. We calculate using the helper lemma (Lemma 5.4):

$$\begin{aligned}
M_{\text{mom}}^{(1)} &= d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \widehat{\sigma}_u \cdot du \\
&= d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \left(\frac{d+n}{n} \cdot |u\rangle\langle u| - \frac{1}{n} \cdot I \right) \cdot du \\
&= d[n] \cdot \frac{d+n}{n} \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+1} \cdot du - d[n] \cdot \frac{1}{n} \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes I \cdot du \\
&= \frac{1}{n} \cdot (e+X_{n+1}) \cdot \Pi_{\text{sym}}^n \otimes I - \frac{1}{n} \cdot \Pi_{\text{sym}}^n \otimes I \\
&= \frac{1}{n} \cdot X_{n+1} \cdot (\Pi_{\text{sym}}^n \otimes I).
\end{aligned}$$

This completes the proof. $\square$

Lemma 5.6 (The second moment operator). *Given $n \geq 1$ and $d \geq 2$,*

$$M_{\text{mom}}^{(2)} = \frac{1}{n^2} \cdot (X_{n+2}X_{n+1} + (n+1, n+2)) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}) - \text{Lower}_{\text{mom}}.$$

where

$$\text{Lower}_{\text{mom}} = \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du.$$

Proof. We calculate:

$$\begin{aligned}
M_{\text{mom}}^{(2)} &= d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \widehat{\sigma}_u^{\otimes 2} \cdot du \\
&= d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \left(\frac{d+n}{n} \cdot |u\rangle\langle u| - \frac{1}{n} \cdot I \right)^{\otimes 2} \cdot du \\
&= d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes \left(\frac{(d+n)^2}{n^2} \cdot |u\rangle\langle u|^{\otimes 2} - \frac{d+n}{n^2} \cdot |u\rangle\langle u| \otimes I - \frac{d+n}{n^2} \cdot I \otimes |u\rangle\langle u| + \frac{1}{n^2} \cdot I \otimes I \right) \cdot du.
\end{aligned}$$

This now splits into four terms. The first term we divide further into a main term and a lower-order term

using $(d+n)^2 = (d+n)(d+n+1) - (d+n) = (d+n)^{\uparrow 2} - (d+n)$:

$$\begin{aligned}
& \frac{(d+n)^2}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du \\
&= \frac{(d+n)^{\uparrow 2}}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du - \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du \\
&= \frac{1}{n^2} \cdot (e + X_{n+2})(e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}) - \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du \\
&= \frac{1}{n^2} \cdot (e + X_{n+2})(e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}) - \text{Lower}_{\text{mom}},
\end{aligned}$$

using the helper lemma ([Lemma 5.4](#)). The second term is (negative)

$$\frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+1} \otimes I \cdot du = \frac{1}{n^2} \cdot (e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}).$$

The third term is identical to the second term, except with its $(n+1)$ -st and $(n+2)$ -nd registers exchanged. Hence, we can write it as (negative)

$$\begin{aligned}
& (n+1, n+2) \cdot \left(\frac{1}{n^2} \cdot (e + X_{n+1}) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}) \right) \cdot (n+1, n+2) \\
&= \frac{1}{n^2} \cdot (e + (n+1, n+2) \cdot X_{n+1} \cdot (n+1, n+2)) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}),
\end{aligned}$$

where we have used here that $(n+1, n+2)$ commutes with $(\Pi_{\text{sym}}^n \otimes I^{\otimes 2})$. Finally, the fourth term is

$$\frac{1}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \otimes I^{\otimes 2} \cdot du = \frac{1}{n^2} \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2}).$$

Comparing all four terms (and excluding the lower-order term from the first term), we see that they are all an element of the symmetric group algebra times $(1/n^2) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2})$. Combining these prefactors, we get

$$\begin{aligned}
& (e + X_{n+2})(e + X_{n+1}) - (e + X_{n+1}) - (e + (n+1, n+2) \cdot X_{n+1} \cdot (n+1, n+2)) + e \\
&= X_{n+2} + X_{n+2}X_{n+1} - (n+1, n+2) \cdot X_{n+1} \cdot (n+1, n+2) \\
&= (n+1, n+2) + X_{n+2}X_{n+1},
\end{aligned}$$

where the final step uses the fact that $(n+1, n+2) \cdot X_{n+1} \cdot (n+1, n+2)$ contains all the swaps in X_{n+2} except $(n+1, n+2)$. Multiplying this by $(1/n^2) \cdot (\Pi_{\text{sym}}^n \otimes I^{\otimes 2})$ and subtracting off the lower-order term completes the proof. $\square$

Now we use these operators to compute the first and second moments of the Grier–Pashayan–Schaeffer algorithm. As before, we will compute these in the general case when their algorithm is performed on a mixed state ψ_{sym} inside the symmetric subspace.

Lemma 5.7 (First moment). *Let ψ_{sym} be a mixed state on $\vee^n \mathbb{C}^d$. If $|v\rangle$ is the output of Hayashi’s measurement when performed on ψ_{sym} , then*

$$\mathbf{E}[\widehat{\sigma}_v] = (\psi_{\text{sym}})_1.$$

Recall our notation for partial trace: $(\psi_{\text{sym}})_1 = \text{tr}_{2\dots n}(\psi_{\text{sym}})$.

Proof. By [Proposition 5.3](#),

$$\begin{aligned}
\mathbf{E}[\widehat{\sigma}_{\mathbf{v}}] &= \text{tr}_{1\dots n}(M_{\text{mom}}^{(1)} \cdot \psi_{\text{sym}} \otimes I_d) \\
&= \text{tr}_{1\dots n} \left(\left(\frac{1}{n} \cdot X_{n+1} \cdot (\Pi_{\text{sym}}^n \otimes I) \right) \cdot \psi_{\text{sym}} \otimes I_d \right) \quad (\text{by Lemma 5.5}) \\
&= \frac{1}{n} \cdot \text{tr}_{1\dots n}(X_{n+1} \cdot \psi_{\text{sym}} \otimes I_d) \\
&= \frac{1}{n} \cdot \sum_{i=1}^n \text{tr}_{1\dots n}((i, n+1) \cdot \psi_{\text{sym}} \otimes I_d) \\
&= \frac{1}{n} \cdot \sum_{i=1}^n \text{tr}_{1\dots n}((1, n+1) \cdot \psi_{\text{sym}} \otimes I_d) \quad (\text{because } \psi_{\text{sym}} \in \mathbb{V}^n \mathbb{C}^d) \\
&= (\psi_{\text{sym}})_1. \quad (\text{by Proposition 2.4})
\end{aligned}$$

This completes the proof. $\square$

Lemma 5.8 (Second moment). *Let ψ_{sym} be a mixed state on $\mathbb{V}^n \mathbb{C}^d$. If $|\mathbf{v}\rangle$ is the output of Hayashi's measurement when performed on ψ_{sym} , then*

$$\mathbf{E}[\widehat{\sigma}_{\mathbf{v}} \otimes \widehat{\sigma}_{\mathbf{v}}] = \frac{n-1}{n} \cdot (\psi_{\text{sym}})_{1,2} + \frac{1}{n} \cdot ((\psi_{\text{sym}})_1 \otimes I + I \otimes (\psi_{\text{sym}})_1) \cdot \text{SWAP} + \frac{1}{n^2} \cdot \text{SWAP} - \text{Lower}_{\psi_{\text{sym}}},$$

where $\text{Lower}_{\psi_{\text{sym}}} \in \text{SoS}(d)$.

Proof. By [Proposition 5.3](#),

$$\mathbf{E}[\widehat{\sigma}_{\mathbf{v}} \otimes \widehat{\sigma}_{\mathbf{v}}] = \text{tr}_{1\dots n}(M_{\text{mom}}^{(2)} \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2}). \quad (24)$$

From [Lemma 5.6](#), we have that

$$M_{\text{mom}}^{(2)} = \frac{1}{n^2} \cdot (X_{n+2}X_{n+1} + (n+1, n+2)) \cdot (\Pi_{\text{sym}}^n \otimes I_d^{\otimes 2}) - \text{Lower}_{\text{mom}},$$

where

$$\text{Lower}_{\text{mom}} = \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} |u\rangle\langle u|^{\otimes n+2} \cdot du.$$

Let us begin by calculating the main terms. Plugging these in to [Equation \(24\)](#) above, we get

$$\begin{aligned}
&\text{tr}_{1\dots n} \left(\left(\frac{1}{n^2} \cdot (X_{n+2}X_{n+1} + (n+1, n+2)) \cdot (\Pi_{\text{sym}}^n \otimes I_d^{\otimes 2}) \right) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \\
&= \text{tr}_{1\dots n} \left(\left(\frac{1}{n^2} \cdot (X_{n+2}X_{n+1} + (n+1, n+2)) \right) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right), \quad (25)
\end{aligned}$$

because $\psi_{\text{sym}} \in \mathbb{V}^n \mathbb{C}^d$. Now we calculate the product of the Jucys–Murphy elements:

$$\begin{aligned}
X_{n+2}X_{n+1} &= \sum_{i=1}^{n+1} (i, n+2) \sum_{j=1}^n (j, n+1) \\
&= \sum_{i=1}^n (i, n+2)(i, n+1) + \sum_{j=1}^n (n+1, n+2)(j, n+1) + \sum_{i \neq j=1}^n (i, n+2)(j, n+1) \\
&= \sum_{i=1}^n (i, n+1, n+2) + \sum_{j=1}^n (j, n+2, n+1) + \sum_{i \neq j=1}^n (i, n+2)(j, n+1) \\
&= \sum_{i=1}^n (i, n+1)(n+1, n+2) + \sum_{j=1}^n (j, n+2)(n+1, n+2) + \sum_{i \neq j=1}^n (i, n+2)(j, n+1).
\end{aligned}$$

We will plug these sums back into Equation (25) above one at a time. First,

$$\begin{aligned}
& \frac{1}{n^2} \cdot \sum_{i=1}^n \text{tr}_{1\dots n} \left((i, n+1)(n+1, n+2) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \\
&= \frac{1}{n^2} \cdot \sum_{i=1}^n \text{tr}_{1\dots n} \left((i, n+1) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \cdot \text{SWAP} \\
&= \frac{1}{n^2} \cdot \sum_{i=1}^n \text{tr}_{1\dots n} \left((1, n+1) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \cdot \text{SWAP} && (\text{because } \psi_{\text{sym}} \in \mathbb{V}^n \mathbb{C}^d) \\
&= \frac{1}{n} \cdot (\psi_{\text{sym}})_1 \otimes I \cdot \text{SWAP}. && (\text{by Proposition 2.4})
\end{aligned}$$

A similar calculation shows that

$$\frac{1}{n^2} \cdot \sum_{j=1}^n \text{tr}_{1\dots n} \left((j, n+2)(n+1, n+2) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) = \frac{1}{n} \cdot I \otimes (\psi_{\text{sym}})_1 \cdot \text{SWAP}.$$

Next,

$$\begin{aligned}
& \frac{1}{n^2} \cdot \sum_{i \neq j=1}^n \text{tr}_{1\dots n} \left((i, n+2)(j, n+1) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \\
&= \frac{1}{n^2} \cdot \sum_{i \neq j=1}^n \text{tr}_{1\dots n} \left((1, n+2)(2, n+1) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) && (\text{because } \psi_{\text{sym}} \in \mathbb{V}^n \mathbb{C}^d) \\
&= \frac{n-1}{n} \cdot (\psi_{\text{sym}})_{1,2}. && (\text{by Proposition 2.4})
\end{aligned}$$

The final term which occurs in Equation (25) is

$$\frac{1}{n^2} \cdot \text{tr}_{1\dots n} \left((n+1, n+2) \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) = \frac{1}{n^2} \cdot \text{SWAP}.$$

This accounts for all the main terms in the lemma statement. The only remaining term we have to account for is the lower order term, which comes from

$$\begin{aligned}
\text{tr}_{1\dots n} \left(\text{Lower}_{\text{mom}} \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) &= \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} \text{tr}_{1\dots n} \left(|u\rangle\langle u|^{\otimes n+2} \cdot \psi_{\text{sym}} \otimes I_d^{\otimes 2} \right) \cdot du \\
&= \frac{d+n}{n^2} \cdot d[n] \cdot \int_{|u\rangle} \text{tr} \left(|u\rangle\langle u|^{\otimes n} \cdot \psi_{\text{sym}} \right) \cdot |u\rangle\langle u|^{\otimes 2} \cdot du,
\end{aligned}$$

which is a nonnegative linear combination of terms of the form $|u\rangle\langle u|^{\otimes 2}$ and is therefore in $\text{SoS}(d)$. This completes the proof. $\square$

Now we specialize these results to the case of $\psi_{\text{sym}} = |\psi\rangle\langle\psi|^{\otimes n}$ which occurs in pure state tomography. The following is an immediate corollary of Lemmas 5.7 and 5.8.

Corollary 5.9 (Moments of Grier–Pashayan–Schaeffer; Theorem 1.5, restated). *If $|\mathbf{v}\rangle$ is the output of Hayashi’s measurement on n copies of $|\psi\rangle \in \mathbb{C}^d$, then $\widehat{\sigma}_{\mathbf{v}}$ is an unbiased estimator for $|\psi\rangle\langle\psi|$, i.e. $\mathbf{E} \widehat{\sigma}_{\mathbf{v}} = |\psi\rangle\langle\psi|$. In addition,*

$$\mathbf{E}[\widehat{\sigma}_{\mathbf{v}} \otimes \widehat{\sigma}_{\mathbf{v}}] = \frac{n-1}{n} \cdot |\psi\rangle\langle\psi|^{\otimes 2} + \frac{1}{n} \cdot (|\psi\rangle\langle\psi| \otimes I_d + I_d \otimes |\psi\rangle\langle\psi|) \cdot \text{SWAP} + \frac{1}{n^2} \cdot \text{SWAP} - \text{Lower}_{\psi},$$

where $\text{Lower}_{\psi} \in \text{SoS}(d)$.

This matches the second moment of our algorithm from Theorem 1.9 in the case of pure states, as $\ell(\boldsymbol{\lambda}) = 1$ always for pure states.

6 A new unbiased estimator for mixed state tomography

In this section, we apply our reduction to the Grier–Pashayan–Schaeffer algorithm to obtain a new unbiased estimator for mixed state tomography. We will use the moments of $\mathcal{A}_{\text{GPS}}$ from [Section 5.2](#) ([Theorem 1.5](#)) to derive the moments of the resulting mixed state tomography algorithm.

We begin by applying the generic reduction to get $\text{Mix}(\mathcal{A}_{\text{GPS}})$ and prove [Theorem 1.7](#). In [Section 6.2](#), we obtain the improved result of [Theorem 1.9](#) that is required for some of our applications. We do this by analyzing the algorithm $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ that results from quasi-purification.

6.1 Warmup: direct reduction to pure state tomography

First, let us apply our reduction to the Grier–Pashayan–Schaeffer algorithm. This results in the following mixed state tomography algorithm.

Given n copies of ρ :

1. First apply $\Phi_{\text{Purify}}^{d,r}$ to produce n copies of a random purification $|\rho\rangle_{\text{AB}} \in \mathbb{C}^d \otimes \mathbb{C}^r$.
2. Apply the Grier–Pashayan–Schaeffer algorithm to learn an estimate $\hat{\sigma}_{\mathbf{v}}$ of $|\rho\rangle\langle\rho|$.
3. Set $(\hat{\rho}_{\mathbf{v}})_{\text{A}} = \text{tr}_{\text{B}}((\hat{\sigma}_{\mathbf{v}})_{\text{AB}})$ of ρ . Output $\hat{\rho}_{\mathbf{v}}$.

Figure 10: The mixed state tomography algorithm $\text{Mix}(\mathcal{A}_{\text{GPS}})$; [Figure 4](#), restated.

The main result of this section is a formula for the first and second moments of the output $\hat{\rho}_{\mathbf{v}}$.

Theorem 6.1 (Moments of $\text{Mix}(\mathcal{A}_{\text{GPS}})$; [Theorem 1.7](#), restated). *Let $\hat{\rho}_{\mathbf{v}}$ be the output of $\text{Mix}(\mathcal{A}_{\text{GPS}})$ in [Figure 10](#) when run on n copies of a rank- r state $\rho \in \mathbb{C}^{d \times d}$. Then $\hat{\rho}_{\mathbf{v}}$ is an unbiased estimator for ρ with second moment*

$$\mathbf{E}[\hat{\rho}_{\mathbf{v}} \otimes \hat{\rho}_{\mathbf{v}}] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{r}{n^2} \cdot \text{SWAP} - \text{Lower}_{\rho},$$

where $\text{Lower}_{\rho} \in \text{SoS}(d)$.

Proof. We will compute the first and second moments, conditioned on the outcome of the random purification $|\rho\rangle$ in [Step 1](#) of the algorithm. It will turn out that the moments don't depend on the purification, so then we will be done. First, some notation: we will write A for the $\mathbb{C}^d$ register of $|\rho\rangle$ and B for its $\mathbb{C}^r$ register. We can calculate the first moment using [Corollary 5.9](#):

$$\mathbf{E}_{\mathbf{v}}[\hat{\rho}_{\mathbf{v}}] = \mathbf{E}_{\mathbf{v}}[\text{tr}_{\text{B}}((\hat{\sigma}_{\mathbf{v}})_{\text{AB}})] = \text{tr}_{\text{B}}(|\rho\rangle\langle\rho|_{\text{AB}}) = \rho.$$

As this holds for any $|\rho\rangle$, it also holds on average for a random $|\rho\rangle$. Thus, $\hat{\rho}_{\mathbf{v}}$ is an unbiased estimator for ρ . Now for the second moment: here, we are looking at two copies of each register, A_1, B_1 , and A_2, B_2 . Conditioned on the purification $|\rho\rangle$, we have that

$$\mathbf{E}_{\mathbf{v}}[\hat{\rho}_{\mathbf{v}} \otimes \hat{\rho}_{\mathbf{v}}] = \mathbf{E}_{\mathbf{v}}[\text{tr}_{\text{B}_1}((\hat{\sigma}_{\mathbf{v}})_{\text{A}_1\text{B}_1}) \otimes \text{tr}_{\text{B}_2}((\hat{\sigma}_{\mathbf{v}})_{\text{A}_2\text{B}_2})] = \mathbf{E}_{\mathbf{v}}[\text{tr}_{\text{B}_1\text{B}_2}(\hat{\sigma}_{\mathbf{v}} \otimes \hat{\sigma}_{\mathbf{v}})] = \text{tr}_{\text{B}_1\text{B}_2}(\mathbf{E}_{\mathbf{v}}[\hat{\sigma}_{\mathbf{v}} \otimes \hat{\sigma}_{\mathbf{v}}]).$$

To calculate the expectation in the partial trace, we can appeal to [Corollary 5.9](#), which states that

$$\mathbf{E}_{\mathbf{v}}[\hat{\sigma}_{\mathbf{v}} \otimes \hat{\sigma}_{\mathbf{v}}] = \frac{n-1}{n} \cdot |\rho\rangle\langle\rho|^{\otimes 2} + \frac{1}{n} \cdot (|\rho\rangle\langle\rho| \otimes I_{\text{A}_2\text{B}_2} + I_{\text{A}_1\text{B}_1} \otimes |\rho\rangle\langle\rho|) \cdot \text{SWAP}_{\text{AB}} + \frac{1}{n^2} \cdot \text{SWAP}_{\text{AB}} - \text{Lower}_{|\rho\rangle},$$

where $\text{Lower}_{|\rho\rangle} \in \text{SoS}(D)$ for $D = d \cdot r$. Here SWAP_{AB} is the operator which swaps A_1B_1 with A_2B_2 ; note that by [Proposition 2.3](#) it factorizes as

$$\text{SWAP}_{\text{AB}} = \text{SWAP}_{\text{A}} \otimes \text{SWAP}_{\text{B}}.$$

Now we compute the partial trace of this expression term-by-term. First,

$$\text{tr}_{\mathbf{B}_1\mathbf{B}_2}(|\rho\rangle\langle\rho|^{\otimes 2}) = \text{tr}_{\mathbf{B}_1}(|\rho\rangle\langle\rho|_{\mathbf{A}_1\mathbf{B}_1}) \otimes \text{tr}_{\mathbf{B}_2}(|\rho\rangle\langle\rho|_{\mathbf{A}_2\mathbf{B}_2}) = \rho_{\mathbf{A}_1} \otimes \rho_{\mathbf{A}_2}.$$

For the second term, we have that

$$\text{tr}_{\mathbf{B}_1\mathbf{B}_2}(|\rho\rangle\langle\rho|_{\mathbf{A}_1\mathbf{B}_1} \otimes I_{\mathbf{A}_2\mathbf{B}_2} \cdot \text{SWAP}_{\mathbf{AB}}) = (\rho_{\mathbf{A}_1} \otimes I_{\mathbf{A}_2}) \cdot \text{SWAP}_{\mathbf{A}},$$

by [Proposition 2.5](#). For the third term, a similar calculation shows that

$$\text{tr}_{\mathbf{B}_1\mathbf{B}_2}(I_{\mathbf{A}_1\mathbf{B}_1} \otimes |\rho\rangle\langle\rho|_{\mathbf{A}_2\mathbf{B}_2} \cdot \text{SWAP}_{\mathbf{AB}}) = (I_{\mathbf{A}_1} \otimes \rho_{\mathbf{A}_2}) \cdot \text{SWAP}_{\mathbf{A}}.$$

For the fourth term, we have that

$$\text{tr}_{\mathbf{B}_1\mathbf{B}_2}(\text{SWAP}_{\mathbf{AB}}) = \text{tr}_{\mathbf{B}_1\mathbf{B}_2}(\text{SWAP}_{\mathbf{A}} \otimes \text{SWAP}_{\mathbf{B}}) = \text{SWAP}_{\mathbf{A}} \cdot \text{tr}(\text{SWAP}_{\mathbf{B}}) = \text{SWAP}_{\mathbf{A}} \cdot r.$$

Finally, for the lower-order term, since $\text{Lower}_{|\rho\rangle} \in \text{SoS}(D)$, $\text{tr}_{\mathbf{B}_1\mathbf{B}_2}(\text{Lower}_{|\rho\rangle}) \in \text{SoS}(d)$ due to [Proposition 2.1](#).

Putting everything together, we have

$$\mathbf{E}[\widehat{\rho}_{\mathbf{v}} \otimes \widehat{\rho}_{\mathbf{v}}] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{r}{n^2} \cdot \text{SWAP} - \text{tr}_{\mathbf{B}_1\mathbf{B}_2}(\text{Lower}_{|\rho\rangle}).$$

As this holds for any $|\rho\rangle$, it also holds on average for a random $|\rho\rangle$. This completes the proof. $\square$

The second moment of $\text{Mix}(\mathcal{A}_{\text{GPS}})$ nearly matches the bound in [Theorem 1.9](#), except the $\mathbf{E}[\ell(\boldsymbol{\lambda})]$ in that expression is replaced by a larger factor of r here. As mentioned before, the factor of $\mathbf{E}[\ell(\boldsymbol{\lambda})]$ is crucial for achieving our sample complexities for the applications of shadow tomography and tomography with limited entanglement. We will be able to achieve the improved second moment bound from [Theorem 1.9](#) using quasi-purification, which we explain in [Section 6.2](#) below.

6.2 Improving the reduction by only quasi-purifying

Now we give our unbiased estimator $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ which improves upon the $\text{Mix}(\mathcal{A}_{\text{GPS}})$ algorithm from [Figure 10](#) above. To motivate our construction, let us take a new perspective on $\text{Mix}(\mathcal{A}_{\text{GPS}})$ that does not treat the random purification channel as a black-box. Previously, we viewed the random purification channel as a way to go from n copies of a mixed state $\rho^{\otimes n}$ to a random purification $|\rho\rangle\langle\rho|^{\otimes n}$, which is then fed into the Grier–Pashayan–Schaeffer algorithm. However, under the perspective of [Section 3](#), the random purification channel performs weak Schur sampling on $\rho^{\otimes n}$, resulting in a Young diagram $\boldsymbol{\lambda} \vdash n$, and the state collapses to $\rho|_{\boldsymbol{\lambda}}$. The purification channel $\Phi_{\text{Purify}}^{d,r}$ is then applied, resulting in the state

$$\Phi_{\text{Purify}}^{d,r}(\rho|_{\boldsymbol{\lambda}}) = |\boldsymbol{\lambda}\boldsymbol{\lambda}\rangle\langle\boldsymbol{\lambda}\boldsymbol{\lambda}|_{\mathbf{Y}\mathbf{Y}'} \otimes |\text{EPR}_{\boldsymbol{\lambda}}\rangle\langle\text{EPR}_{\boldsymbol{\lambda}}|_{\mathbf{P}\mathbf{P}'} \otimes \left(\frac{\nu_{\boldsymbol{\lambda}}^d(\rho)}{s_{\boldsymbol{\lambda}}^d(\rho)}\right)_{\mathbf{Q}} \otimes \left(\frac{I_{\dim(V_{\boldsymbol{\lambda}}^r)}}{\dim(V_{\boldsymbol{\lambda}}^r)}\right)_{\mathbf{Q}'}. \quad (26)$$

Then the Grier–Pashayan–Schaeffer algorithm is applied to *this* state. At first blush, this perspective looks a bit strange, because their algorithm typically takes as input a state of the form $|\psi\rangle\langle\psi|^{\otimes n}$, but the state in [Equation \(26\)](#) is clearly not of this form. But this state *is* at least in the symmetric subspace $\vee^n(\mathbb{C}^d \otimes \mathbb{C}^r)$ due to [Lemma 3.2](#), and so applying the Grier–Pashayan–Schaeffer algorithm to this state is at the very least a well-defined operation. In addition, we have seen in [Section 5.2](#) that it is possible to analyze their algorithm for general states drawn from the symmetric subspace, rather than just states of the form $|\psi\rangle\langle\psi|^{\otimes n}$.

From this perspective, there is no particular reason why, for each $\boldsymbol{\lambda}$, we must purify our state using a purification register $\mathbb{C}^r$ of the same dimension r . Instead, for each $\boldsymbol{\lambda}$, our algorithm will purify to the smallest dimension possible for the purification register, which is $\ell(\boldsymbol{\lambda})$. (That $\ell(\boldsymbol{\lambda})$ is the smallest possible purification dimension comes from the fact that the purifying irrep register $V_{\boldsymbol{\lambda}}^r$ in [Equation \(26\)](#) is only well-defined when $r \geq \ell(\boldsymbol{\lambda})$.) Intuitively, reducing the size of the purification register reduces the size of the Hilbert space that the Grier–Pashayan–Schaeffer algorithm needs to search over, reducing the number of copies needed.

Given n copies of ρ :

1. Apply the Schur transform U_{Schur}^d to $\rho^{\otimes n}$.
2. Perform weak Schur sampling. Letting λ be the outcome, the state collapses to $\rho|_{\lambda}$. Set $\ell := \ell(\lambda)$.
3. Apply the purification channel to compute $\Phi_{\text{Purify}}^{d,\ell}(\rho|_{\lambda})$.
4. Apply an inverse Schur transform to both registers, i.e. the operation $(U_{\text{Schur}}^d \otimes U_{\text{Schur}}^{\ell})^{\dagger}$. Write

$$\tau_{\lambda}(\rho) := (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^{\ell})^{\dagger} \cdot \Phi_{\text{Purify}}^{d,\ell}(\rho|_{\lambda}) \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^{\ell})$$

for the resulting state.

5. The state $\tau_{\lambda}(\rho)$ is an element of $\vee^n(\mathbb{C}^d \otimes \mathbb{C}^{\ell}) \cong \vee^n \mathbb{C}^D$, for $D = d \cdot \ell$. Apply the Grier–Pashayan–Schaeffer algorithm to learn an estimate $\hat{\sigma}_v^{\lambda}$.
6. Set $\hat{\rho}_v^{\lambda} = \text{tr}_2(\hat{\sigma}_v^{\lambda})$. Output $\hat{\rho}_v^{\lambda}$.

Figure 11: Our improved mixed state tomography algorithm $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$.

Before calculating any moments, let us first consider the intermediate state $\tau_{\lambda}(\rho)$. It has n registers of type $\mathbb{C}^d$, which we will refer to as the $A_1, \dots, A_n$ registers, and n registers of type $\mathbb{C}^{\ell}$, which we will refer to as the $B_1, \dots, B_n$ registers. We will use the following helper lemma.

Lemma 6.2 (Partial trace helper lemma). *For any $1 \leq k \leq n$,*

$$\mathbf{E}_{\lambda} \left[\text{tr}_{A_{k+1} \dots A_n, B_1 \dots B_n}(\tau_{\lambda}(\rho)) \right] = \rho^{\otimes k}.$$

Proof. Let us first calculate the partial trace for a fixed λ .

$$\begin{aligned} \text{tr}_{A_{k+1} \dots A_n, B_1 \dots B_n}(\tau_{\lambda}(\rho)) &= \text{tr}_{A_{k+1} \dots A_n, B_1 \dots B_n} \left((U_{\text{Schur}}^d \otimes U_{\text{Schur}}^{\ell})^{\dagger} \cdot \Phi_{\text{Purify}}^{d,\ell}(\rho|_{\lambda}) \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^{\ell}) \right) \\ &= \text{tr}_{A_{k+1} \dots A_n, B_1 \dots B_n} \left((U_{\text{Schur}}^d \otimes I)^{\dagger} \cdot \Phi_{\text{Purify}}^{d,\ell}(\rho|_{\lambda}) \cdot (U_{\text{Schur}}^d \otimes I) \right) \\ &\quad \text{(by the cyclic property of the partial trace)} \\ &= \text{tr}_{A_{k+1} \dots A_n} \left((U_{\text{Schur}}^d)^{\dagger} \cdot \text{tr}_B(\Phi_{\text{Purify}}^{d,\ell}(\rho|_{\lambda})) \cdot U_{\text{Schur}}^d \right) \\ &= \text{tr}_{A_{k+1} \dots A_n} \left((U_{\text{Schur}}^d)^{\dagger} \cdot \rho|_{\lambda} \cdot U_{\text{Schur}}^d \right). \end{aligned} \quad \text{(by Lemma 3.6)}$$

Averaging over λ gives us

$$\begin{aligned} \mathbf{E}_{\lambda} \left[\text{tr}_{A_{k+1} \dots A_n, B_1, \dots, B_n}(\tau_{\lambda}(\rho)) \right] &= \mathbf{E}_{\lambda} \left[\text{tr}_{A_{k+1} \dots A_n} \left((U_{\text{Schur}}^d)^{\dagger} \cdot \rho|_{\lambda} \cdot U_{\text{Schur}}^d \right) \right] \\ &= \text{tr}_{A_{k+1} \dots A_n} \left((U_{\text{Schur}}^d)^{\dagger} \cdot \mathbf{E}_{\lambda}[\rho|_{\lambda}] \cdot U_{\text{Schur}}^d \right) \\ &= \text{tr}_{A_{k+1} \dots A_n}(\rho^{\otimes n}) \quad \text{(by Lemma 3.7)} \\ &= \rho^{\otimes k}. \end{aligned}$$

This completes the proof. $\square$

Now we calculate the first and second moments of our estimator. We begin by showing that it does indeed give an unbiased estimator of ρ .

Theorem 6.3 (First moment of $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$). *Let $\hat{\rho}_v^{\lambda}$ be the output of our mixed state tomography algorithm. Then*

$$\mathbf{E}[\hat{\rho}_v^{\lambda}] = \rho.$$

Proof. First, let us calculate the expectation of $\hat{\rho}_v^\lambda$ conditioned on the result λ of weak Schur sampling (which, in turn, conditions on the values ℓ and D). This estimate is the result of applying the Grier–Pashayan–Schaeffer algorithm to the state $\tau_\lambda(\rho)$. Then by [Lemma 5.7](#),

$$\mathbf{E}[\hat{\sigma}_v^\lambda \mid \lambda] = \text{tr}_{\substack{A_2 \dots A_n \\ B_2 \dots B_n}}(\tau_\lambda(\rho)).$$

Then because $\hat{\rho}_v^\lambda$ is the result of tracing out $\hat{\sigma}_v^\lambda$'s purifying register, we have

$$\mathbf{E}[\hat{\rho}_v^\lambda \mid \lambda] = \text{tr}_{\substack{A_2 \dots A_n \\ B_1 \dots B_n}}(\tau_\lambda(\rho)).$$

Note the addition of the B_1 register to the partial trace. Averaging over λ gives us

$$\mathbf{E}[\hat{\rho}_v^\lambda] = \mathbf{E}_\lambda \left[\text{tr}_{\substack{A_2 \dots A_n \\ B_1 \dots B_n}}(\tau_\lambda(\rho)) \right] = \rho,$$

by [Lemma 6.2](#). This completes the proof. $\square$

Next, we compute its second moment.

Theorem 6.4 (Second moment of $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$). *Let $\hat{\rho}_v^\lambda$ be the output of our mixed state tomography algorithm. Then*

$$\mathbf{E}[\hat{\rho}_v^\lambda \otimes \hat{\rho}_v^\lambda] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \cdot \text{SWAP} - \text{Lower}_\rho,$$

where $\text{Lower}_\rho \in \text{SoS}(d)$.

Proof. Following the proof of the first moment case, we will calculate the expectation of $\hat{\rho}_v^\lambda \otimes \hat{\rho}_v^\lambda$ conditioned on the result λ of weak Schur sampling. By [Lemma 5.8](#),

$$\begin{aligned} \mathbf{E}[\hat{\sigma}_v^\lambda \otimes \hat{\sigma}_v^\lambda \mid \lambda] &= \frac{n-1}{n} \cdot \text{tr}_{\substack{A_3 \dots A_n \\ B_3 \dots B_n}}(\tau_\lambda(\rho)) && \text{(term 1)} \\ &+ \frac{1}{n} \cdot (\text{tr}_{\substack{A_2 \dots A_n \\ B_2 \dots B_n}}(\tau_\lambda(\rho)) \otimes I_D) \cdot \text{SWAP}_{AB} && \text{(term 2)} \\ &+ \frac{1}{n} \cdot (I_D \otimes (\text{tr}_{\substack{A_2 \dots A_n \\ B_2 \dots B_n}}(\tau_\lambda(\rho))) \cdot \text{SWAP}_{AB} && \text{(term 3)} \\ &+ \frac{1}{n^2} \cdot \text{SWAP}_{AB} && \text{(term 4)} \\ &- \text{Lower}_{\tau_\lambda(\rho)}, && \text{(term 5)} \end{aligned}$$

where $\text{Lower}_{\tau_\lambda(\rho)} \in \text{SoS}(D)$. Next, to compute $\mathbf{E}[\hat{\rho}_v^\lambda \otimes \hat{\rho}_v^\lambda \mid \lambda]$, we trace out both purifying registers. For term 1, this results in

$$\frac{n-1}{n} \cdot \text{tr}_{\substack{A_3 \dots A_n \\ B_1 \dots B_n}}(\tau_\lambda(\rho)).$$

For term 2, this results in

$$\text{tr}_B \left(\frac{1}{n} \cdot (\text{tr}_{\substack{A_2 \dots A_n \\ B_2 \dots B_n}}(\tau_\lambda(\rho)) \otimes I_D) \cdot \text{SWAP}_{AB} \right) = \frac{1}{n} \cdot (\text{tr}_{\substack{A_2 \dots A_n \\ B_1 \dots B_n}}(\tau_\lambda(\rho)) \otimes I_d) \cdot \text{SWAP}_A,$$

by [Proposition 2.5](#). Term 3 results in a similar expression. For term 4, we have

$$\text{tr}_B \left(\frac{1}{n^2} \cdot \text{SWAP}_{AB} \right) = \text{tr} \left(\frac{1}{n^2} \cdot \text{SWAP}_B \right) \cdot \text{SWAP}_A = \frac{\ell(\lambda)}{n^2} \cdot \text{SWAP}_A,$$

by [Proposition 2.3](#). Finally, for term 5, $\text{Lower}'_{\lambda} = \text{tr}_B(\text{Lower}_{\tau_{\lambda}(\rho)})$ is in $\text{SoS}(d)$, by [Proposition 2.1](#). In total, we have

$$\begin{aligned}
\mathbf{E}[\hat{\rho}_v^{\lambda} \otimes \hat{\rho}_v^{\lambda} \mid \lambda] &= \frac{n-1}{n} \cdot \text{tr}_{A_3 \dots A_n, B_1 \dots B_n}(\tau_{\lambda}(\rho)) && \text{(term 1)} \\
&+ \frac{1}{n} \cdot (\text{tr}_{A_2 \dots A_n, B_1 \dots B_n}(\tau_{\lambda}(\rho)) \otimes I_d) \cdot \text{SWAP}_A && \text{(term 2)} \\
&+ \frac{1}{n} \cdot (I_d \otimes (\text{tr}_{A_2 \dots A_n, B_2 \dots B_n}(\tau_{\lambda}(\rho)))) \cdot \text{SWAP}_A && \text{(term 3)} \\
&+ \frac{\ell(\lambda)}{n^2} \cdot \text{SWAP}_A && \text{(term 4)} \\
&- \text{Lower}'_{\lambda}. && \text{(term 5)}
\end{aligned}$$

Now we take the expectation of this with respect to λ . By [Lemma 6.2](#), this is

$$\mathbf{E}[\hat{\rho}_v^{\lambda} \otimes \hat{\rho}_v^{\lambda}] = \frac{n-1}{n} \cdot \rho^{\otimes 2} + \frac{1}{n} \cdot (\rho \otimes I_d + I_d \otimes \rho) \cdot \text{SWAP} + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \cdot \text{SWAP} - \mathbf{E}[\text{Lower}'_{\lambda}].$$

Note that $\mathbf{E}_{\lambda}[\text{Lower}'_{\lambda}] \in \text{SoS}(d)$, as it is a convex combination of matrices in $\text{SoS}(d)$. This completes the proof. $\square$

Acknowledgments

A.P. is supported by DARPA under Agreement No. HR00112020023. J.S. and J.W. are supported by the NSF CAREER award CCF-233971. E.T. is supported by the Miller Institute for Basic Research in Science, University of California, Berkeley.

References

- [ABDY22] Srinivasan Arunachalam, Sergey Bravyi, Arkopal Dutt, and Theodore Yoder. Optimal algorithms for learning quantum phase states. In *Proceedings of the 18th Conference on the Theory of Quantum Computation, Communication and Cryptography*, 2022.
- [BCH05] Dave Bacon, Isaac Chuang, and Aram Harrow. The quantum Schur transform: I. efficient qudit circuits. In *Proceedings of the 18th Annual ACM-SIAM Symposium on Discrete Algorithms*, 2005.
- [BCS⁺25] Jacob Beckey, Luke Coffman, Ariel Shlosberg, Louis Schatzki, and Felix Leditzky. Product testing with single-copy measurements. Technical report, arXiv:2510.07820, 2025.
- [Bel75] Viacheslav Belavkin. Optimal multiple quantum statistical hypothesis testing. *Stochastics: An International Journal of Probability and Stochastic Processes*, 1(1-4):315–345, 1975.
- [BHH16] Fernando Brandao, Aram Harrow, and Michał Horodecki. Local random quantum circuits are approximate polynomial-designs. *Communications in Mathematical Physics*, 346(2):397–434, 2016.
- [BK02] Howard Barnum and Emanuel Knill. Reversing quantum dynamics with near-optimal quantum and classical fidelity. *Journal of Mathematical Physics*, 43(5):2097–2106, 2002.
- [Bog07] V.I. Bogachev. *Measure Theory*. Mathematics and Statistics. Springer Berlin Heidelberg, 2007.
- [CCHL22] Sitan Chen, Jordan Cotler, Hsin-Yuan Huang, and Jerry Li. Exponential separations between learning with and without quantum memory. In *Proceedings of the 62nd Annual IEEE Symposium on Foundations of Computer Science*, pages 574–585, 2022.

- [CHL⁺23] Sitan Chen, Brice Huang, Jerry Li, Allen Liu, and Mark Sellke. When does adaptivity help for quantum state learning? In *Proceedings of the 64th Annual IEEE Symposium on Foundations of Computer Science*, pages 391–404, 2023.
- [CLL24a] Sitan Chen, Jerry Li, and Allen Liu. Optimal high-precision shadow estimation. In *27th Conference on Quantum Information Processing*, 2024.
- [CLL24b] Sitan Chen, Jerry Li, and Allen Liu. An optimal tradeoff between entanglement and copy complexity for state tomography. In *Proceedings of the 56th Annual ACM Symposium on Theory of Computing*, pages 1331–1342, 2024.
- [CWZ24] Kean Chen, Qisheng Wang, and Zhicheng Zhang. Local test for unitarily invariant properties of bipartite quantum states. Technical report, arXiv:2404.04599, 2024.
- [DDGG20] Rafał Demkowicz-Dobrzański, Wojciech Górecki, and Mădălin Guță. Multi-parameter estimation beyond quantum fisher information. *Journal of Physics A: Mathematical and Theoretical*, 53(36):363001, 2020.
- [FR13] Simon Foucart and Holger Rauhut. *A Mathematical Introduction to Compressive Sensing*. Springer, 2013.
- [GKKT20] Madalin Guță, Jonas Kahn, Richard Kueng, and Joel A Tropp. Fast state tomography with optimal error bounds. *Journal of Physics A: Mathematical and Theoretical*, 53(20):204001, 2020.
- [GLM24] Daniel Grier, Sihan Liu, and Gaurav Mahajan. Improved classical shadows from local symmetries in the Schur basis. Technical report, arXiv:2405.09525, 2024.
- [GPS24] Daniel Grier, Hakop Pashayan, and Luke Schaeffer. Sample-optimal classical shadows for pure states. *Quantum*, 8:1373, 2024.
- [Har05] Aram Harrow. *Applications of coherent classical communication and the Schur transform to quantum information theory*. PhD thesis, Massachusetts Institute of Technology, 2005.
- [Har13] Aram Harrow. The church of the symmetric subspace. Technical report, arXiv:1308.6595, 2013.
- [Har23] Aram Harrow. Approximate orthogonality of permutation operators, with application to quantum information. *Letters in Mathematical Physics*, 114(1):1, 2023.
- [Hay98] Masahito Hayashi. Asymptotic estimation theory for a finite-dimensional pure state model. *Journal of Physics A: Mathematical and General*, 31(20):4633, 1998.
- [Hel67] Carl Helstrom. Minimum mean-squared error of estimates in quantum statistics. *Physics letters A*, 25(2):101–102, 1967.
- [HHJ⁺16] Jeongwan Haah, Aram Harrow, Zhengfeng Ji, Xiaodi Wu, and Nengkun Yu. Sample-optimal tomography of quantum states. In *Proceedings of the 48th Annual ACM Symposium on Theory of Computing*, August 2016. Preprint.
- [HJ13] Roger Horn and Charles Johnson. *Matrix analysis*. Cambridge University Press, 2nd edition, 2013.
- [HKOT23] Jeongwan Haah, Robin Kothari, Ryan O’Donnell, and Ewin Tang. Query-optimal estimation of unitary channels in diamond distance. In *Proceedings of the 64th Annual IEEE Symposium on Foundations of Computer Science*, pages 363–390, 2023.
- [HKP20] Hsin-Yuan Huang, Richard Kueng, and John Preskill. Predicting many properties of a quantum system from very few measurements. *Nature Physics*, 16(10):1050–1057, 2020.
- [HLT24] Jeongwan Haah, Yunchao Liu, and Xinyu Tan. Efficient approximate unitary designs from random pauli rotations. In *Proceedings of the 65th Annual IEEE Symposium on Foundations of Computer Science*, pages 463–475, 2024.

- [Hol79] Alexander Holevo. On asymptotically optimal hypothesis testing in quantum statistics. *Theory of Probability & Its Applications*, 23(2):411–415, 1979.
- [Hol11] Alexander Holevo. *Probabilistic and statistical aspects of quantum theory*. Springer Science & Business Media, 2011.
- [HW94] Paul Hausladen and William Wootters. A ‘pretty good’ measurement for distinguishing quantum states. *Journal of Modern Optics*, 41(12):2385–2390, 1994.
- [KRT14] Richard Kueng, Holger Rauhut, and Ulrich Terstiege. Low rank matrix recovery from rank one measurements. Technical report, arXiv:1410.6913, 2014.
- [Low10] Andrew Richard Low. *Pseudo-randomness and Learning in Quantum Computation*. PhD thesis, University of Bristol, 2010.
- [OSP23] Ryan O’Donnell, Rocco A Servedio, and Pedro Paredes. Explicit orthogonal and unitary designs. In *Proceedings of the 64th Annual IEEE Symposium on Foundations of Computer Science*, pages 1240–1260, 2023.
- [OW15] Ryan O’Donnell and John Wright. Quantum spectrum testing. In *Proceedings of the 47th Annual ACM Symposium on Theory of Computing*, 2015.
- [OW16] Ryan O’Donnell and John Wright. Efficient quantum tomography. In *Proceedings of the 48th Annual ACM Symposium on Theory of Computing*, 2016.
- [OW17] Ryan O’Donnell and John Wright. Efficient quantum tomography II. In *Proceedings of the 49th Annual ACM Symposium on Theory of Computing*, 2017.
- [PSW25] Angelos Pelecanos, Jack Spilecki, and John Wright. The debiased Keyl’s algorithm: a new unbiased estimator for full state tomography. Manuscript, 2025.
- [Sko23] Maciej Skorski. Bernstein-type bounds for beta distribution. *Modern Stochastics: Theory and Applications*, 10(2):211–228, 2023.
- [SSW25] Thilo Scharnhorst, Jack Spilecki, and John Wright. Optimal lower bounds for quantum state tomography. Manuscript, 2025.
- [SW22] Mehdi Soleimanifar and John Wright. Testing matrix product states. In *Proceedings of the 33rd Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 1679–1701, 2022.
- [TWZ25] Ewin Tang, John Wright, and Mark Zhandry. Conjugate queries can help. Technical report, arXiv:2510.07622, 2025.
- [Ver18] Roman Vershynin. *High-dimensional probability: an introduction with applications in data science*. Cambridge University Press, 2018.
- [Wat18] John Watrous. *The theory of quantum information*. Cambridge University Press, 2018.
- [Wri16] John Wright. *How to learn a quantum state*. PhD thesis, Carnegie Mellon University, 2016.
- [Yue23] Henry Yuen. An improved sample complexity lower bound for (fidelity) quantum state tomography. *Quantum*, 7:890, 2023.
- [ZC25] Sisi Zhou and Senrui Chen. Randomized measurements for multi-parameter quantum metrology. Technical report, arXiv:2502.03536, 2025.
- [ZS00] Karol Zyczkowski and Hans-Jürgen Sommers. Truncations of random unitary matrices. *Journal of Physics A: Mathematical and General*, 33(10):2045–2057, March 2000.
- [ZS01] Karol Zyczkowski and Hans-Jürgen Sommers. Induced measures in the space of mixed quantum states. *Journal of Physics A: Mathematical and General*, 34(35):7111–7125, August 2001.

A Viewing our algorithms as pretty good measurements

Any tomography algorithm—whether it performs multiple measurements, introduces ancillary systems, or applies intermediate channels—is equivalent to another tomography algorithm that performs a (possibly complicated) POVM directly on the input state $\rho^{\otimes n}$. This raises the question: what POVMs are our algorithms performing? In this section, we show that $\mathbf{Mix}(\mathcal{A}_{\text{GPS}})$ implements a pretty good measurement (PGM) over a natural distribution known as the *Hilbert–Schmidt measure* (defined below). Furthermore, we show that $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$ is only slightly more elaborate: it performs a PGM over one of several such distributions, conditioned on the outcome of weak Schur sampling.

Ours is not the first tomography algorithm which can be viewed as performing a PGM on the input. Indeed, Hayashi’s pure state tomography algorithm [Hay98] can be viewed as performing a PGM over a distribution induced by a Haar random unit vector. In addition, one of the two mixed state tomography algorithms in [HHJ⁺16] performs a PGM over a particular distribution on mixed states, and the other, like $\mathbf{Mix}^+(\mathcal{A}_{\text{GPS}})$, performs a PGM conditioned on the outcome of weak Schur sampling. To our knowledge, however, the particular choice of PGM we study, based on the Hilbert–Schmidt measure, has not appeared in the literature prior to our work.

A.1 PGM preliminaries

A PGM is defined in terms of a set of states $\{\rho_i\}$ and a prior probability distribution on these states $\{\alpha_i\}$. The corresponding PGM has measurement operators $\{M_i\}$ with

$$M_i := N^{-1/2} \cdot \alpha_i \rho_i \cdot N^{-1/2},$$

where $N := \sum_i \alpha_i \rho_i$. Here, $N^{-1/2}$ is the *Moore–Penrose pseudoinverse* of $N^{1/2}$, i.e. the inverse restricted to the support of $N^{1/2}$.

The PGM was originally introduced for its application to the problem of *quantum state discrimination* [Bel75, Hol79, HW94]. In this problem, we are given one copy of a state ρ drawn randomly from the set $\{\rho_i\}$, with ρ_i sampled with probability α_i , and asked to identify the state. The pretty good measurement is *pretty good* at this task. In particular, let P_{PGM} be the success probability of the algorithm that first measures ρ using the PGM, obtaining outcome i , and then outputs ρ_i . Then $P_{\text{PGM}} \geq P_{\text{OPT}}^2$, where P_{OPT} is the optimal success probability across all possible measurement schemes [BK02].

A.2 The PGM over the Hilbert–Schmidt measure

Our PGM will have measurement outcomes indexed by mixed states $\sigma \in \mathbb{C}^{d \times d}$. The state indexed by σ is the n -fold product $\sigma^{\otimes n}$. We will need to define a measure on mixed states to state the probability associated to this state.

Definition A.1 (Rank- r Hilbert–Schmidt measure). The *rank- r Hilbert–Schmidt measure*, denoted $\mu_{\text{HS}}^{d,r}$, is the measure on mixed states $\sigma \in \mathbb{C}^{d \times d}$ induced by the Haar measure on pure states $|u\rangle_{\mathbf{A}_1 \mathbf{B}_1} \in \mathbb{C}^d \otimes \mathbb{C}^r$, obtained by setting $\sigma_{\mathbf{A}_1} = \text{tr}_{\mathbf{B}_1}(|u\rangle\langle u|_{\mathbf{A}_1 \mathbf{B}_1})$.⁵

Equivalently, the rank- r Hilbert–Schmidt measure is the *pushforward* of the Haar measure on pure states under the map $\text{tr}_{\mathbf{B}_1}$. Since $\text{tr}_{\mathbf{B}_1}$ is measurable, for every measurable g on the space of $d \times d$ density matrices we have the change-of-variables formula

$$\int_{\text{tr}_2(X)} g(\sigma) \cdot d\mu_{\text{HS}}(\sigma) = \int_X g(\text{tr}_2(|u\rangle\langle u|)) \cdot du, \quad (27)$$

whenever one side is well-defined [Bog07, Theorem 3.6.1].

For our PGM, the prior probability of $\sigma^{\otimes n}$ will be $d\mu_{\text{HS}}^{d,r}(\sigma)$. The measurement operators of the PGM are then given by

$$M_{\text{HS}}^{d,r}(\sigma) := N^{-1/2} \cdot \sigma^{\otimes n} \cdot d\mu_{\text{HS}}^{d,r}(\sigma) \cdot N^{-1/2}. \quad (28)$$

⁵This measure is also known as the *induced measure* in the literature, and, in the $r = d$ special case, coincides with the *Hilbert–Schmidt measure* [ZS01]

We now compute S . Letting $|u\rangle_{AB} \in \mathbb{C}^d \otimes \mathbb{C}^r$ and $(\sigma_u)_A := \text{tr}_B(|u\rangle\langle u|_{AB})$, we have

$$N = \int_{\sigma} \sigma^{\otimes n} \cdot d\mu_{\text{HS}}^{d,r}(\sigma) = \int_{|u\rangle} \sigma_u^{\otimes n} \cdot du = \text{tr}_{B_1 \dots B_n} \left(\int_{|u\rangle} |u\rangle\langle u|^{\otimes n} \cdot du \right) = \text{tr}_{B_1 \dots B_n} \left(\frac{1}{D[n]} \cdot \Pi_{\text{sym}}^{n,D} \right). \quad (29)$$

Here, $D := d \cdot r$ and $D[n] = \dim(\vee^n \mathbb{C}^D)$. However, by [Lemma 3.2](#), we know

$$(U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r) \cdot \Pi_{\text{sym}}^{n,D} \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^{\dagger} = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_{Y Y'} \otimes |\text{EPR}_{\lambda}\rangle\langle\text{EPR}_{\lambda}|_{P P'} \otimes I_{Q Q'}.$$

Tracing out the B registers (which now correspond to $Y' P' Q'$) and plugging the result back into [Equation \(29\)](#) gives us an expression for N in the Schur basis:

$$U_{\text{Schur}}^d \cdot N \cdot (U_{\text{Schur}}^d)^{\dagger} = \frac{1}{D[n]} \cdot \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_Y \otimes \left(\frac{I_{\dim(\lambda)}}{\dim(\lambda)} \right)_P \otimes \left(\dim(V_{\lambda}^r) \cdot I_{\dim(V_{\lambda}^r)} \right)_Q. \quad (30)$$

Moreover, $\sigma^{\otimes n}$ can also be expressed in the Schur basis as

$$U_{\text{Schur}}^d \cdot \sigma^{\otimes n} \cdot (U_{\text{Schur}}^d)^{\dagger} = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} |\lambda\rangle\langle\lambda|_Y \otimes (I_{\dim(\lambda)})_P \otimes (\nu_{\lambda}^d(\sigma))_Q. \quad (31)$$

Combining [Equations \(28\)](#), [\(30\)](#) and [\(31\)](#), we get:

$$U_{\text{Schur}}^d \cdot M_{\text{HS}}^{d,r}(\sigma) \cdot (U_{\text{Schur}}^d)^{\dagger} = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} \frac{D[n] \cdot \dim(\lambda)}{\dim(V_{\lambda}^r)} \cdot |\lambda\rangle\langle\lambda|_Y \otimes (I_{\dim(\lambda)})_P \otimes \nu_{\lambda}^d(\sigma)_Q \cdot d\mu_{\text{HS}}^{d,r}(\sigma). \quad (32)$$

We summarize this construction with the following definition.

Definition A.2. The *pretty good measurement over the rank- r Hilbert–Schmidt measure* is a PGM with operators labeled by mixed states $\sigma \in \mathbb{C}^{d \times d}$. The state corresponding to σ is $\sigma^{\otimes n}$, and the corresponding probability is $d\mu_{\text{HS}}^{d,r}(\sigma)$. The resulting measurement operators are given by [Equation \(32\)](#).

A.3 Viewing $\text{Mix}(\mathcal{A}_{\text{GPS}})$ as a PGM

In this section, we show that the measurement performed by $\text{Mix}(\mathcal{A}_{\text{GPS}})$ is equivalent to the PGM over the rank- r Hilbert–Schmidt measure. This PGM is also the measurement performed by $\text{Mix}(\mathcal{A}_{\text{Hayashi}})$, as $\text{Mix}(\mathcal{A}_{\text{GPS}})$ and $\text{Mix}(\mathcal{A}_{\text{Hayashi}})$ differ only in their post-processing of the measurement outcome.

The algorithm $\text{Mix}(\mathcal{A}_{\text{GPS}})$, described in [Figure 4](#), first applies $\Phi_{\text{Purify}}^{d,r}$ to the input state, and then applies Hayashi’s measurement. Recall that this is the POVM with measurement operators:

$$\{D[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du : |u\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r\}.$$

Finally, the algorithm proceeds by processing $\sigma_u = \text{tr}_2(|u\rangle\langle u|)$. Let $X \subseteq \mathbb{C}^{d \times d}$ be a subset of mixed states (measurable with respect to the Hilbert–Schmidt measure), and let $Y \subseteq \mathbb{C}^D$ be the preimage of X under tracing out the second register, i.e. $X = \text{tr}_2(Y)$ (since X is measurable, Y is measurable with respect to the Haar measure). The probability that we obtain a $\sigma \in X$ after purifying, measuring, and tracing out a given, generic, input $\psi \in \mathbb{C}^{d^n \times d^n}$ (i.e. not necessarily in the symmetric subspace) is

$$\mathbf{Pr}_{\text{Mix}(\mathcal{A}_{\text{GPS}})}[\sigma \in X | \psi] = \text{tr} \left(\Phi_{\text{Purify}}^{d,r}(\psi) \cdot \int_{|u\rangle \in Y} D[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du \right).$$

On the other hand, the probability that our PGM obtains $\sigma \in X$, given input ψ , is

$$\mathbf{Pr}_{\text{PGM}}[\sigma \in X | \psi] = \int_{\sigma \in X} \text{tr}(\psi \cdot M_{\text{HS}}^{d,r}(\sigma)).$$

Proposition A.3. For any input ψ and any measurable set $X \subseteq \mathbb{C}^{d \times d}$ (measured with $\mu_{\text{HS}}^{d,r}$),

$$\Pr_{\text{Mix}(\mathcal{A}_{\text{GPS}})}[\sigma \in X|\psi] = \Pr_{\text{PGM}}[\sigma \in X|\psi].$$

Thus, $\text{Mix}(\mathcal{A}_{\text{GPS}})$ implements the pretty good measurement over the rank- r Hilbert–Schmidt measure.

Before proving this statement, it will be useful to note that we can decompose $\Phi_{\text{Purify}}^{d,r}$ into Kraus operators $\{K_{\lambda ST}\}$ as $\Phi_{\text{Purify}}^{d,r}(\psi) = \sum_{\lambda, S, T} K_{\lambda ST} \cdot \psi \cdot K_{\lambda ST}^\dagger$ for any generic state ψ written in the Schur basis of $(\mathbb{C}^d)^{\otimes n}$, with

$$K_{\lambda ST} := |\lambda\lambda\rangle_{YY'} \langle \lambda|_Y \otimes |\text{EPR}_\lambda\rangle_{PP'} \langle S|_P \otimes \frac{|T\rangle_{Q'}}{\sqrt{\dim(V_\lambda^r)}} \otimes (I_{\dim(V_\lambda^d)})_Q. \quad (33)$$

Here, $\lambda \vdash n$ and $\ell(\lambda) \leq r$; S is an SYT of shape λ ; T is an SSYT of shape λ and alphabet $[r]$. We can consider the adjoint $\Phi_{\text{Purify}}^{d,r,\dagger}$, which can be expanded into Kraus operators as $\Phi_{\text{Purify}}^{d,r,\dagger}(\varphi) = \sum_{\lambda, S, T} K_{\lambda ST}^\dagger \cdot \varphi \cdot K_{\lambda ST}$, where φ is a state written in the Schur basis of $(\mathbb{C}^D)^{\otimes n}$.

Lemma A.4. For any $|u\rangle \in \mathbb{C}^d \otimes \mathbb{C}^r$, we have

$$\Phi_{\text{Purify}}^{d,r,\dagger} \left((U_{\text{Schur}}^{\otimes 2}) \cdot |u\rangle\langle u|^{\otimes n} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger \right) = \sum_{\lambda \vdash n, \ell(\lambda) \leq r} \frac{\dim(\lambda)}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda^d(\sigma_u).$$

Proof. By Lemma 3.3, we have

$$(U_{\text{Schur}}^{\otimes 2}) \cdot |u\rangle\langle u|^{\otimes n} \cdot (U_{\text{Schur}}^{\otimes 2})^\dagger = \sum_{\substack{\lambda \vdash n, \ell(\lambda) \leq r \\ \mu \vdash n, \ell(\mu) \leq r}} |\lambda\lambda\rangle\langle\mu\mu|_{YY'} \otimes |\text{EPR}_\lambda\rangle\langle\text{EPR}_\mu|_{PP'} \otimes |u_{\lambda\lambda}\rangle\langle u_{\mu\mu}|_{QQ'},$$

for some vectors $|u_{\lambda\lambda}\rangle_{QQ'}$ which satisfy $\text{tr}_{Q'}(|u_{\lambda\lambda}\rangle\langle u_{\lambda\lambda}|_{QQ'}) = \dim(\lambda) \cdot \nu_\lambda^d(\sigma_u)$, with $\sigma_u = \text{tr}_2(|u\rangle\langle u|)$. We then have

$$\begin{aligned} & \Phi_{\text{Purify}}^{d,r,\dagger} \left((U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r) \cdot |u\rangle\langle u|^{\otimes n} \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^\dagger \right) \\ &= \sum_{\lambda, S, T} K_{\lambda ST}^\dagger \cdot \left((U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r) \cdot |u\rangle\langle u|^{\otimes n} \cdot (U_{\text{Schur}}^d \otimes U_{\text{Schur}}^r)^\dagger \right) \cdot K_{\lambda ST} \\ &= \sum_{\lambda, S, T} K_{\lambda ST}^\dagger \cdot \left(\sum_{\sigma, \tau} |\sigma\sigma\rangle\langle\tau\tau|_{YY'} \otimes |\text{EPR}_\sigma\rangle\langle\text{EPR}_\tau|_{PP'} \otimes |u_{\sigma\sigma}\rangle\langle u_{\tau\tau}|_{QQ'} \right) \cdot K_{\lambda ST} \\ &= \sum_{\lambda, S, T} \frac{1}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda|_Y \otimes |S\rangle\langle S|_P \otimes (\langle T|_{Q'} \cdot |u_{\lambda\lambda}\rangle\langle u_{\lambda\lambda}|_{QQ'} \cdot |T\rangle_{Q'}) \quad (\text{Equation (33)}) \\ &= \sum_{\lambda} \frac{1}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \text{tr}_{Q'}(|u_{\lambda\lambda}\rangle\langle u_{\lambda\lambda}|_{QQ'}) \\ &= \sum_{\lambda} \frac{\dim(\lambda)}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda^d(\sigma_u). \end{aligned}$$

This completes the proof. $\square$

Proof of Proposition A.3. Note that

$$\begin{aligned}
& \int_{|u\rangle \in Y} D[n] \cdot \Phi_{\text{Purify}}^{d,r,\dagger}(|u\rangle\langle u|^{\otimes n}) \cdot du \\
&= (U_{\text{Schur}}^d)^\dagger \cdot \left(\sum_{\lambda \vdash n, \ell(\lambda) \leq r} \frac{D[n] \cdot \dim(\lambda)}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \int_{|u\rangle \in Y} \nu_\lambda^d(\sigma_u) \cdot du \right) \cdot U_{\text{Schur}}^d \quad (\text{Lemma A.4}) \\
&= (U_{\text{Schur}}^d)^\dagger \cdot \left(\sum_{\lambda \vdash n, \ell(\lambda) \leq r} \frac{D[n] \cdot \dim(\lambda)}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \int_{\sigma \in X} \nu_\lambda^d(\sigma) \cdot d\mu_{\text{HS}}^{d,r}(\sigma) \right) \cdot U_{\text{Schur}}^d \quad (\text{Equation (27)}) \\
&= \int_{\sigma \in X} (U_{\text{Schur}}^d)^\dagger \cdot \left(\sum_{\lambda \vdash n, \ell(\lambda) \leq r} \frac{D[n] \cdot \dim(\lambda)}{\dim(V_\lambda^r)} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda^d(\sigma) \cdot d\mu_{\text{HS}}^{d,r}(\sigma) \right) \cdot U_{\text{Schur}}^d \\
&= \int_{\sigma \in X} M_{\text{HS}}^{d,r}(\sigma). \quad (\text{Equation (32)})
\end{aligned}$$

Thus,

$$\begin{aligned}
\mathbf{Pr}_{\text{Mix}(\mathcal{A}_{\text{GPS}})}[\sigma \in X | \psi] &= \text{tr}_{\text{AB}} \left(\Phi_{\text{Purify}}^{d,r}(\psi) \cdot \int_{|u\rangle \in Y} D[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du \right) \\
&= \text{tr}_{\text{A}} \left(\psi \cdot \int_{|u\rangle \in Y} D[n] \cdot \Phi_{\text{Purify}}^{d,r,\dagger}(|u\rangle\langle u|^{\otimes n}) \cdot du \right) \\
&= \text{tr}_{\text{A}} \left(\psi \cdot \int_{\sigma \in X} M_{\text{HS}}^{d,r}(\sigma) \right) \\
&= \mathbf{Pr}_{\text{PGM}}[\sigma \in X | \psi]. \quad \square
\end{aligned}$$

Viewing $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ as a PGM. Note that $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$, described in Figure 11, can be regarded as the following two-step process: first, weak Schur sample $\rho^{\otimes n}$, to obtain $\lambda \vdash n$ and a state $\rho|_\lambda$; second, apply $\text{Mix}(\mathcal{A}_{\text{GPS}})$ to $\rho|_\lambda$ with r set to $\ell(\lambda)$. As a consequence of Proposition A.3, we therefore have the following result.

Corollary A.5. $\text{Mix}^+(\mathcal{A}_{\text{GPS}})$ implements weak Schur sampling, and, conditioned on the Young diagram λ observed, followed by a pretty good measurement over the rank- $\ell(\lambda)$ Hilbert–Schmidt measure.