

The debiased Keyl’s algorithm: a new unbiased estimator for full state tomography

Angelos Pelecanos*

Jack Spilecki*

John Wright*

Abstract

In the problem of quantum state tomography, one is given n copies of an unknown rank- r mixed state $\rho \in \mathbb{C}^{d \times d}$ and asked to produce an estimator $\hat{\rho}$ of ρ which is close to it. One of the most basic and useful properties a statistical estimator can have is that of being *unbiased*, meaning that $\hat{\rho}$ is equal to the true state ρ in expectation. Unfortunately, of the three estimators for tomography which are known to be either sample-optimal or almost sample-optimal, namely Keyl’s algorithm [Key06] and the two tomography algorithms from [HHJ⁺16], none of them produce unbiased estimators. This situation has recently been highlighted as a bottleneck in several important tomographic applications [CLL24b, CLL24a].

In this work, we present the *debiased Keyl’s algorithm*, the first estimator for full state tomography which is both unbiased and sample-optimal. Our algorithm is the result of taking Keyl’s algorithm and carefully modifying it in order to correct for its bias. In addition, we derive an explicit formula for the second moment of our estimator which is simple and easy to use. Using it, we show the following applications.

- We show that $n = O(rd/\varepsilon^2)$ copies are sufficient to learn a rank- r mixed state $\rho \in \mathbb{C}^{d \times d}$ to trace distance error ε . This gives a new proof of the main result of [OW16] and matches the lower bound of [SSW25].
- We then improve this result and show that $n = O(rd/\varepsilon^2)$ copies are sufficient to learn to error ε in the more challenging Bures distance. This is optimal due to the lower bound of [Yue23] and improves on the best known bound of $n = \tilde{O}(rd/\varepsilon^2)$ from prior work [HHJ⁺16, OW17, FO24].
- We consider the scenario of full state tomography with limited entanglement, in which one is given n copies of ρ but only allowed to perform entangled measurements across k of them at a time. We show that

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

copies of ρ suffice to learn ρ to trace distance error ε in this scenario. This improves on the prior work of [CLL24b] and matches their lower bound which holds for $k \leq 1/\varepsilon^c$, where c is a constant.

- In the problem of shadow tomography, one is given m observables $O_1, \dots, O_m$, and the goal is to estimate each quantity $\text{tr}(O_i \cdot \rho)$ up to ε error. If $\|O_i\|_\infty \leq 1$ for all i , we show that $O(\log(m)/\varepsilon^2)$ copies are sufficient to solve this task in the “high accuracy regime” when $\varepsilon = O(1/d)$. This improves a result of [CLL24a], who showed this bound holds when $\varepsilon = O(1/d^{12})$. More generally, we show that if $\text{tr}(O_i^2) \leq F$ for all i , then

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

copies suffice to solve this task. This improves on a result of [GLM24], who showed a bound of $n = O(\log(m) \cdot \sqrt{rF}/\varepsilon^2)$, and disproves a conjecture of [CLL24a].

- Our final application is to the field of quantum metrology. In quantum metrology, one is provided copies of a mixed state ρ_θ parameterized by a vector $\theta \in \mathbb{R}^n$, and the goal is to estimate the parameter θ . An algorithm is *locally unbiased* at a parameter θ^* if, roughly, it is unbiased given ρ_θ when θ is in a small neighborhood around θ^* . The quantum Cramér–Rao bound states that the mean squared error matrix (MSEM) of any locally unbiased algorithm is lower bounded by the inverse of the Quantum Fisher Information (QFI) matrix. We give a locally unbiased algorithm whose MSEM is *upper bounded* by twice the inverse of the QFI matrix in the asymptotic limit of large n , which is optimal.

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

Contents

I	Introduction	4
1	Introduction	4
1.1	Debiasing tomography algorithms	5
1.1.1	The single copy case	5
1.1.2	The pure state case	6
1.1.3	Keyl's algorithm	6
1.1.4	The debiased Keyl's algorithm	7
1.2	Application 1: full state tomography in trace distance	10
1.3	Application 2: full state tomography in Bures distance	11
1.4	Application 3: tomography with limited entanglement	12
1.5	Application 4: shadow tomography	13
1.6	Application 5: quantum metrology	15
1.7	Technical overview	17
1.7.1	The single copy case	18
1.7.2	The two copy case	20
1.7.3	The general copy case	22
1.8	Open problems	23
1.9	Paper organization	24
2	Preliminaries	25
2.1	Quantum distances	25
2.2	Learning with high probability and amplification	26
2.3	On low-rank tomography algorithms	26
2.4	Representation theory	27
2.4.1	Basics	27
2.4.2	Partitions and Young diagrams	28
2.4.3	The irreducible representations of S_n	29
2.4.4	The irreducible representations of $U(d)$	31
2.4.5	Schur-Weyl duality	33
2.4.6	Weak Schur sampling	35
3	The debiased Keyl's algorithm	35
3.1	Keyl's algorithm	35
3.2	The debiased Keyl's algorithm	36
3.3	A family of unbiased estimators	37
II	Applications	42
4	Full state tomography in trace distance	42
5	Full state tomography in Bures distance	44
5.1	Full-rank tomography	45
5.2	Low-rank tomography	47
5.3	Learning bipartite pure states	52
6	Tomography with limited entanglement	53
7	Shadow tomography	55
8	Quantum metrology	57

III	The Clebsch-Gordan transform	62
9	The Clebsch-Gordan transform	62
10	Clebsch-Gordan coefficients	62
10.1	Clebsch-Gordan insertion	65
10.2	One-step coefficients	68
10.3	Two-step coefficients	70
11	Constructing the Schur transform	72
12	Dual Clebsch-Gordan transform	73
IV	The moments of the debiased Keyl’s estimator	75
13	The first moment	75
13.1	Technical lemmas concerning one-step Clebsch-Gordan coefficients	75
13.2	Proof of the first moment	77
14	The second moment	81
14.1	Strategy overview	81
14.2	Block-diagonal expressions in the Schur basis	85
14.3	Technical lemmas concerning two-step Clebsch-Gordan coefficients	88
14.4	Main Lemma	94
14.5	Proof of the Diagonal Expression Lemma	97
14.6	Deferred proofs	101
14.6.1	Proof of Theorem 14.13	101
14.6.2	Proof of Theorem 14.14	104
15	Large-d expansion of the second moment	111
V	Appendix	116
A	Obtaining Young’s orthogonal basis	116

Part I

Introduction

1 Introduction

A statistical estimator $\widehat{X}$ for a quantity X is said to be *unbiased* if $\mathbf{E}[\widehat{X}] = X$. Unbiased estimators have many desirable properties, and for many quantities of interest in statistics, the gold standard estimator is unbiased. As Voinov and Nikulin put it in *Unbiased Estimators and their Applications* [VN12, VN96], “Unbiasedness is one of the most important properties of statistical estimators”. Illustrating the abundance of unbiased estimators in statistics, they collect roughly a thousand (!) such estimators for different statistical parameters.

This work is about the statistical problem of estimating a mixed quantum state $\rho \in \mathbb{C}^{d \times d}$ given n identical copies of ρ . This task, known as *quantum tomography*, is of fundamental importance to both theory and practice, and it has received significant attention by the research community in recent years. A variety of different tomography algorithms have been studied in the literature, and many of these do actually give unbiased estimators for ρ . Two well-known examples are the “uniform POVM tomography algorithm” from [Wri16, GKKT20] and the “textbook” Pauli tomography algorithm, which is covered in [Wri24a]. In addition to these, we also know of several estimators which achieve the optimal sample complexity of $n = O(d^2)$ copies for full state tomography [HHJ⁺16, OW16]. However, we do not know of any estimators which are simultaneously unbiased *and* sample-optimal. The reason is that most existing unbiased estimators for full state tomography measure each copy of ρ independently, but it is known that any sample-optimal tomography algorithm must use entangled measurements across many copies of ρ [CHL⁺23]. On the flip side, measurements which are entangled across many copies of ρ are often mathematically complex and difficult to understand, and even computing the bias of known sample-optimal tomography algorithms appears to be quite challenging.

It is important to understand whether sample-optimal estimators can also be unbiased if we want to understand the information theoretic limits of quantum state learning. This issue was highlighted in two recent works of Chen, Li, and Liu [CLL24b, CLL24a], who studied two problems in quantum tomography for which having an unbiased estimator would be extremely useful: full state tomography using measurements with limited entanglement and shadow tomography. However, since we do not know of any sample-optimal estimators which are unbiased, they had to develop sophisticated mathematical techniques for analyzing the known biased, sample-optimal tomography algorithms as if they were unbiased. In the course of this work, they suggested that perhaps there may not even *exist* any unbiased, sample-optimal tomography algorithms, saying that “unbiasedness seems specific to incoherent measurements”.

In this work, we counter this intuition by developing an estimator for full state which is both unbiased and sample-optimal. Our estimator is based on an estimator for full state tomography introduced by Keyl in [Key06] and shown to be sample-optimal in [OW16]. Keyl’s algorithm produces a biased estimator for ρ , but we show that a simple modification to this algorithm can correct for the bias and give an unbiased estimator. With this “debiased Keyl’s algorithm” in hand, we show a number of interesting applications to full state tomography. For some of these applications, we are able to derive conceptually cleaner proofs of existing results. For other applications, we are able to use the debiased Keyl’s algorithm to give new results which were not known before, such as the first sample-optimal bound for learning mixed states with fidelity error. Prior to our work, there were three other sample-optimal (or approximately sample-optimal) algorithms to choose from when performing full state tomography [Key06, HHJ⁺16], but we believe that our results suggest that the debiased Keyl’s algorithm may be the “right” fully entangled tomography algorithm to study in the future, and we expect that it will be useful in resolving several other remaining open problems in the study of full state tomography. Below, we describe several known unbiased estimators for full state tomography in order to motivate our construction of the debiased Keyl’s algorithm. Following that, we survey the applications of this new algorithm.

1.1 Debiasing tomography algorithms

1.1.1 The single copy case

Suppose we only have a single copy of ρ , and we want to perform full state tomography on it. Exactly which measurement is the best to perform is perhaps not clear, but a natural thing to do is to measure ρ in *some* orthonormal basis, as in the following algorithm.

1. Pick an orthonormal basis $|u_1\rangle, \dots, |u_d\rangle \in \mathbb{C}^d$.
2. Measure ρ in this basis. Let $\mathbf{i} \in [d]$ be the outcome.
3. Output $\hat{\rho}_0 = |u_{\mathbf{i}}\rangle\langle u_{\mathbf{i}}|$.

How well this algorithm performs depends on which basis we decide to measure in. For simplicity, let us assume that ρ is diagonal in the standard basis, meaning that we can write $\rho = \sum_{i=1}^d \alpha_i \cdot |i\rangle\langle i|$. If the basis we measure in happens to also be the standard basis, it is easy to see that

$$\mathbf{E}[\hat{\rho}_0] = \alpha_1 \cdot |1\rangle\langle 1| + \dots + \alpha_d \cdot |d\rangle\langle d| = \rho,$$

and so $\hat{\rho}_0$ is an unbiased estimator for ρ . On the other hand, suppose the basis we measure in forms a mutually unbiased basis with the standard basis, meaning that $|\langle i|u_j\rangle|^2 = 1/d$ for all i and j . This is the case if, for example, $\{|u_i\rangle\}_{i \in [d]}$ is the Fourier basis. Then we will observe each $i \in [d]$ with equal probability, and so

$$\mathbf{E}[\hat{\rho}_0] = \frac{1}{d} \cdot |u_1\rangle\langle u_1| + \dots + \frac{1}{d} \cdot |u_d\rangle\langle u_d| = \frac{I}{d},$$

giving the maximally mixed state in expectation, regardless of what the state ρ is. We can view the first case as giving us “all signal” and the second case as giving us “all noise”.

It is natural to randomize the choice of basis in order to avoid accidentally falling into the “all noise” case all of the time. Doing so involves measuring according to a random basis $|u_1\rangle, \dots, |u_d\rangle \in \mathbb{C}^d$ sampled from the Haar measure. Such a basis will still be *almost* mutually unbiased with the standard basis, in the sense that we expect $|\langle i|u_j\rangle|^2$ to be close to $1/d$ for most values of i and j , but the fact that it is not *exactly* mutually unbiased means that we can still hope for a small amount of signal to leak through. In this case, our measurement, over the randomness in the choice of basis, is equivalent to performing the *uniform POVM*, the continuous POVM given by $\{d \cdot |u\rangle\langle u| \cdot du\}$. This means that we can rewrite the algorithm as follows.

1. Measure ρ with the uniform POVM $\{d \cdot |u\rangle\langle u| \cdot du\}$. Let $\mathbf{u}$ be the outcome.
2. Output $\hat{\rho}_0 = |\mathbf{u}\rangle\langle \mathbf{u}|$.

The expectation of this estimator can be computed using standard techniques; the calculation can be found, for example, in [Wri16, Page 115]:

$$\mathbf{E}[\hat{\rho}_0] = \left(\frac{1}{d+1}\right) \cdot \rho + \left(\frac{d}{d+1}\right) \cdot \frac{I}{d}.$$

We can see that the first term is a small “signal” term and the second term is a large “noise” term. Moreover, the weights on these two terms are fixed constants independent of the state ρ . It is therefore straightforward to correct for this bias, and doing so results in the following estimator:

$$\hat{\rho} = (d+1) \cdot \hat{\rho}_0 - I = (d+1) \cdot |\mathbf{u}\rangle\langle \mathbf{u}| - I.$$

Indeed, this satisfies $\mathbf{E}[\hat{\rho}] = \rho$ and is therefore an unbiased estimator. This single copy estimator was independently introduced by Krishnamurthy and Wright [Wri16, Section 5.1] and Guta et al. [GKKT20] and forms the basis of an optimal unentangled measurement tomography algorithm; we will discuss this algorithm further in [Section 1.4](#) below.

1.1.2 The pure state case

Suppose we are given n copies of a rank-one mixed state ρ , i.e. $\rho = |v\rangle\langle v|$ for some pure state $|v\rangle \in \mathbb{C}^d$. In this case, there is a well-known tomography algorithm due to Hayashi [Hay98] which is given as follows.

1. Measure $\rho^{\otimes n}$ with the POVM $\{d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du\}$, where $d[n] = \binom{d+n-1}{n}$ is the dimension of the symmetric subspace $\vee^n \mathbb{C}^d$. Let $|u\rangle$ be the outcome.
2. Output $\hat{\rho}_0 = |u\rangle\langle u|$.

The POVM $\{d[n] \cdot |u\rangle\langle u|^{\otimes n} \cdot du\}$ can be viewed as a natural generalization of the uniform POVM from above to the n -fold symmetric subspace. It is well-known that this is a sample-optimal algorithm for pure state tomography. In particular, when $n = O(d/\varepsilon^2)$, this estimator satisfies $D_{\text{tr}}(\rho, \hat{\rho}_0) \leq \varepsilon$ with probability 99% [Wri24b]. However, $\hat{\rho}_0$ is *not* an unbiased estimator of ρ . This is because ρ is a pure state, whereas $\mathbf{E}[\hat{\rho}_0]$ will certainly be a mixed state with rank strictly larger than 1 (unless $\hat{\rho}_0 = \rho$ with probability 1, which is impossible). The expectation of this estimator can be computed using standard techniques; [GPS24, Lemma 13] give the following expression:

$$\mathbf{E}[\hat{\rho}_0] = \left(\frac{n}{d+n}\right) \cdot \rho + \left(\frac{d}{d+n}\right) \cdot \frac{I}{d}.$$

Again, the first term corresponds to the “signal” and the second term corresponds to the “noise”. Note that the coefficient on the “signal” term now improves as n grows larger, and approaches 1 as n approaches ∞ . Since the coefficients are fixed constants independent of the state ρ , it is straightforward to correct for this bias, resulting in the following unbiased estimator:

$$\hat{\rho} = \left(\frac{d+n}{n}\right) \cdot \hat{\rho}_0 - \frac{1}{n} \cdot I = \left(\frac{d+n}{n}\right) \cdot |u\rangle\langle u| - \frac{1}{n} \cdot I.$$

This estimator was introduced by Grier, Pashayan, and Schaeffer [GPS24] with applications to the shadow tomography problem; we will discuss this algorithm further in Section 1.5.

1.1.3 Keyl’s algorithm

Our goal is to perform a similar debiasing for a fully entangled tomography algorithm. There are three reasonable candidates to select from: Keyl’s algorithm [Key06] and the two pretty good measurement-inspired algorithms from Haah et al. [HHJ⁺16]. Of these, two are known to achieve optimal sample complexities of $n = O(d^2/\varepsilon^2)$ for full state tomography: Keyl’s algorithm (due to [OW16]) and the first of the two Haah et al. algorithms (due to [OW15a]). Of these two, we have decided to study Keyl’s algorithm, for two reasons: first, it is the more mathematically elegant of the two, and we believe this will make its bias easier to analyze. Second, as Aram Harrow once suggested to the third author, it seems plausible that it uses the *optimal* entangled measurement for full state tomography [Har16], though it is not clear how to exactly formulate what it means for it to be optimal.

Keyl’s algorithm is a sophisticated representation theoretic algorithm which exploits the symmetries in the state $\rho^{\otimes n}$. In phase one of the algorithm, it performs a measurement on $\rho^{\otimes n}$ known as *weak Schur sampling* in order to estimate the eigenvalues of ρ . Weak Schur sampling returns as a measurement outcome a *Young diagram* $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$, which is a vector of nonnegative integers satisfying $\lambda_1 \geq \dots \geq \lambda_d$ and $\lambda_1 + \dots + \lambda_d = n$. Young diagrams are represented pictorially by arranging boxes into rows, as in Figure 1. From this picture, the Young diagram looks like a (sorted) histogram of n samples drawn from some discrete distribution over d items, and coincidentally, ρ ’s eigenvalues $\alpha = (\alpha_1, \dots, \alpha_d)$ form a probability distribution over the set of integers $\{1, \dots, d\}$. It is therefore natural to wonder if the *normalized Young diagram* $\lambda/n = (\lambda_1/n, \dots, \lambda_d/n)$ provides an accurate approximation of the spectrum α , and this question was answered in the affirmative by independent works of Alicki, Rudnicki, and Sadowski [ARS88] and Keyl and Werner [KW01]; these works showed that as $n \rightarrow \infty$, $\lambda/n \rightarrow \alpha$ with probability 1. This estimator is known both as the *empirical Young diagram algorithm* and also as the *Keyl-Werner algorithm*, and there exists a number of follow-up works showing stronger and stronger bounds on the number of copies it requires to estimate α [HM02, CM06, OW16, OW17]. The final two of these works show that λ/n is close to α once

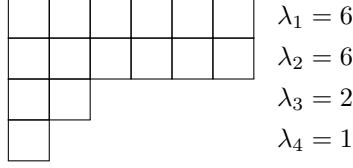


Figure 1: The Young diagram corresponding to the partition $\lambda = (6, 6, 2, 1)$. In this case, $n = 15$ and $d = 4$. This resembles the histogram, sorted from highest to lowest and then turned on its side, of $n = 15$ samples $\mathbf{x}_1, \dots, \mathbf{x}_{15}$ drawn from a discrete distribution $p = (p_1, p_2, p_3, p_4)$ whose probability values, when sorted from highest to lowest, are approximately $6/15$, $6/15$, $2/15$, and $1/15$.

$n = O(d^2)$ copies are used in total variation distance, Hellinger distance, KL divergence, and chi-squared divergence; furthermore, a lower bound from [OW15b] shows that the empirical Young diagram algorithm requires $n = \Omega(d^2)$ copies to produce a good estimate in any of these distances. To date, it remains the best known algorithm for estimating the spectrum of a mixed state. (That said, we don't yet have matching lower bounds, suggesting the tantalizing possibility that there might exist an even better spectrum estimation algorithm.)

Having received the measurement outcome λ , the state $\rho^{\otimes n}$ partially collapses. In phase two of Keyl's algorithm, it performs a further measurement on this collapsed state in order to learn the eigenvectors of ρ . This measurement is based on a strategy of "rotated highest weight vectors"; we will describe it in detail in Section 3.1 below once we have established the relevant representation theoretic background. It produces a unitary matrix $U \in U(d)$ as a measurement outcome which is meant to be a good proxy for the change of basis which rotates the standard basis into ρ 's eigenbasis. Setting $\hat{\alpha}_0 = \lambda/n$ for our approximation of the spectrum, the estimate that Keyl's algorithm outputs is $\hat{\rho}_0 = U \cdot \hat{\alpha}_0 \cdot U^\dagger$.

Keyl's algorithm produces a high quality estimate of the state ρ . It is known, for example, that $D_{\text{tr}}(\rho, \hat{\rho}_0) \leq \varepsilon$ with high probability once $n = O(d^2/\varepsilon^2)$, and furthermore that this is optimal due to a matching lower bound of Haah et al. [HHJ⁺16]. However, it is certainly not unbiased. For example, when $n = 1$, it turns out that Keyl's algorithm is equivalent to the biased single-copy tomography algorithm from Section 1.1.1. Furthermore, when n is allowed to be any value but ρ is restricted to be a pure state, it turns out that Keyl's algorithm is equivalent to the biased pure state tomography algorithm from Section 1.1.2 due to Hayashi. However, unlike in these special cases, the bias of Keyl's algorithm is in general quite difficult to calculate, and so there is no obvious method for debiasing it. Compounding this issue, the bias of Keyl's algorithm actually depends on the underlying state ρ , and so there cannot be just a single correction factor which applies in every case.

1.1.4 The debiased Keyl's algorithm

To debias Keyl's algorithm, we draw inspiration from the debiased algorithms in Sections 1.1.1 and 1.1.2. In both of these cases, to convert the biased estimator $\hat{\rho}_0$ into an unbiased estimator $\hat{\rho}$, we modify the eigenvalues of $\hat{\rho}_0$ but leave the eigenvectors intact. This suggests that to debias Keyl's algorithm, we will want an estimator of the form $\rho = U \cdot \hat{\alpha} \cdot U^\dagger$, where U is our original change of basis but $\hat{\alpha}$ is a new estimate of the spectrum which results from modifying the original estimate $\hat{\alpha}_0$. To define this modification, we first introduce a transformation on Young diagrams.

Definition 1.1 (Young diagram box donation). Let $\lambda = (\lambda_1, \dots, \lambda_d)$ be a Young diagram. Then $\text{donate}(\lambda) \in \mathbb{Z}^d$ is the vector defined as

$$\text{donate}(\lambda)_i = \lambda_i - \sum_{j=1}^{i-1} 1[\lambda_j > \lambda_i] + \sum_{j=i+1}^d 1[\lambda_i > \lambda_j].$$

Intuitively, each row of the Young diagram λ "donates" a box to each row which starts off longer than it and receives a box from each row which starts off shorter than it. We include a diagram of $\text{donate}(\lambda)$ in Figure 2.

To our knowledge, this is a new transformation on Young diagrams, and we could find no similar transformation in the literature. Note that this transformation only shifts boxes around and does not introduce

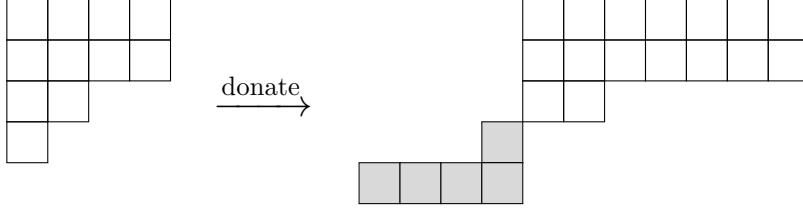


Figure 2: Let $d = 5$. On the left is the Young diagram $\lambda = (4, 4, 2, 1, 0)$. On the right is $\text{donate}(\lambda) = (7, 7, 2, -1, -4)$. The gray-colored boxes correspond to rows with a negative length. Note that $\text{donate}(\lambda)$ is not necessarily a Young diagram, since it may have negative entries.

any new boxes, and so if $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ and $\lambda' = \text{donate}(\lambda)$, then $\lambda'_1 + \dots + \lambda'_d = \lambda_1 + \dots + \lambda_d = n$. With this definition in place, we can define our debiased Keyl's algorithm.

Definition 1.2 (Debiased Keyl's algorithm). Suppose we are given n copies of a mixed state $\rho \in \mathbb{C}^{d \times d}$. The *debiased Keyl's algorithm* works as follows.

1. Measure $\rho^{\otimes n}$ as in Keyl's algorithm, producing a Young diagram $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ and a unitary matrix $U \in U(d)$.
2. Set $\lambda^\dagger = \text{donate}(\lambda)$ and $\hat{\alpha} = \lambda^\dagger/n$. Output $\hat{\rho} = U \cdot \hat{\alpha} \cdot U^\dagger$.

To gain intuition for this algorithm, let us define $|u_i\rangle := U \cdot |i\rangle$ for each $1 \leq i \leq d$. Then

$$\hat{\rho}_0 = \sum_{i=1}^d (\lambda_i/n) \cdot |u_i\rangle\langle u_i| \quad \text{and} \quad \hat{\rho} = \sum_{i=1}^d (\lambda_i^\dagger/n) \cdot |u_i\rangle\langle u_i|$$

For example, when $n = 1$ and $\lambda = (1, 0, \dots, 0)$, we have

$$\begin{aligned} \hat{\rho}_0 &= |u_1\rangle\langle u_1| \\ \hat{\rho} &= d \cdot |u_1\rangle\langle u_1| - (|u_2\rangle\langle u_2| + \dots + |u_d\rangle\langle u_d|) = (d+1) \cdot |u_1\rangle\langle u_1| - I, \end{aligned}$$

as in [Section 1.1.1](#). Next, when ρ is a pure state, then $\lambda = (n, 0, \dots, 0)$ always, and

$$\begin{aligned} \hat{\rho}_0 &= \binom{n}{n} \cdot |u_1\rangle\langle u_1| = |u_1\rangle\langle u_1| \\ \hat{\rho} &= \left(\frac{n+d-1}{n}\right) \cdot |u_1\rangle\langle u_1| - \frac{1}{n} \cdot (|u_2\rangle\langle u_2| + \dots + |u_d\rangle\langle u_d|) = \left(\frac{n+d}{n}\right) \cdot |u_1\rangle\langle u_1| - \frac{1}{n} \cdot I, \end{aligned}$$

as in [Section 1.1.2](#). More generally, Keyl's algorithm is designed to ensure that $|u_1\rangle$, the eigenvector corresponding to the highest eigenvalue of $\hat{\rho}_0$, is the most accurate of all the eigenvectors, whereas $|u_d\rangle$ is the least accurate. This means that $|u_1\rangle\langle u_1|$ will be the term with the highest “signal” and $|u_d\rangle\langle u_d|$ will be the term with the highest “noise”. Thus, the noisier terms will donate boxes to each of the terms with more signal, and this “rich gets richer” approach will boost the signal while depressing the noise. The next theorem is our first main result, which shows that the Young diagram box donation transformation results in an estimator in which all of the noise is precisely cancelled out.

Theorem 1.3 (Debiased Keyl's algorithm is an unbiased estimator). *Let $\hat{\rho}$ be the output of the debiased Keyl's algorithm when run on $\rho^{\otimes n}$. Then $\mathbf{E}[\hat{\rho}] = \rho$.*

Having computed the first moment of $\hat{\rho}$, we next want to compute an expression for its second moment $\mathbf{E}[\hat{\rho} \otimes \hat{\rho}]$. The second moment is useful because it helps us understand the extent to which $\hat{\rho}$ varies about its expectation. We do so in the following theorem, which is our second main result.

Theorem 1.4 (Second moment of the debiased Keyl’s algorithm). *Let $\hat{\rho}$ be the output of the debiased Keyl’s algorithm when run on $\rho^{\otimes n}$. Then*

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

Here, Lower_ρ refers to the “lower order terms” and has the following characterization: it can be written as a positive linear combination of terms of the form $(P \otimes P) \cdot \text{SWAP}$, where each $P \in \mathbb{C}^{d \times d}$ is a Hermitian, positive semidefinite matrix.

To understand this expression, the first term, which can be viewed as the “main term”, is essentially equal to $\rho \otimes \rho$, which is what we would expect to see if $\hat{\rho}$ were always equal to its expectation. The second and third terms are the “error terms” and have the largest contribution to the variance of $\hat{\rho}$. As for the final term, it turns out that we can expand the second moment of $\hat{\rho}$ as a polynomial in the variable d^{-1} , yielding an expression of the form

$$\mathbf{E}[\hat{\rho} \otimes \hat{\rho}] = A_\rho + B_\rho \cdot \frac{1}{d} + C_\rho \cdot \frac{1}{d^2} + \dots,$$

where A_ρ, B_ρ, C_ρ , and so forth are matrix-valued coefficients which depend on the state ρ and the number of copies n but not on the dimension d . As we will explain in [Section 15](#) below, the highest order term A_ρ is exactly the first three terms in [Equation \(1\)](#), and so the remaining low-order terms $B/d + C/d^2 + \dots$ are given by $-\text{Lower}_\rho$. In the case when $d \gg n$ these terms are small enough that they can effectively be discarded. (Indeed, as we discuss in [Section 15](#), there are situations where it is possible to embed ρ into a larger space of dimension $D \gg d$ in order to artificially inflate its dimension and make these terms arbitrarily small.)

However, it turns out that for our applications, we can disregard these lower order terms for a simpler reason, which is that their contribution is negative in the cases we care about. In one of our applications, for example, we have to upper bound expressions of the form

$$\mathbf{E}[\hat{\rho}_{i,j}]^2 = \mathbf{E}[\hat{\rho}_{i,j} \cdot \hat{\rho}_{j,i}] = \mathbf{E}[\langle i | \hat{\rho} | j \rangle \cdot \langle j | \hat{\rho} | i \rangle] = \langle ij | \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}] \cdot | ji \rangle.$$

The contribution to this expression from the lower order terms is $-\langle ij | \cdot \text{Lower}_\rho \cdot | ji \rangle$. From our characterization of Lower_ρ , we can expand $\langle ij | \cdot \text{Lower}_\rho \cdot | ji \rangle$ as a positive linear combination of terms of the form

$$\langle ij | \cdot (P \otimes P) \cdot \text{SWAP} \cdot | ji \rangle = \langle ij | \cdot (P \otimes P) \cdot | ij \rangle = P_{ii} \cdot P_{jj} \geq 0,$$

where the last step uses the fact that P is positive semidefinite. Hence, $-\langle ij | \cdot \text{Lower}_\rho \cdot | ji \rangle$ is nonpositive, and we are safe to discard it. In another of our applications, we have to upper bound expressions of the form

$$\mathbf{E} \text{tr}(O \cdot \hat{\rho})^2 = \mathbf{E} \text{tr}(O \otimes O \cdot \hat{\rho} \otimes \hat{\rho}) = \text{tr}(O \otimes O \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}]),$$

where $O \in \mathbb{C}^{d \times d}$ is an observable. The contribution to this expression from the lower order terms is $-\text{tr}(O \otimes O \cdot \text{Lower}_\rho)$; we can expand $\text{tr}(O \otimes O \cdot \text{Lower}_\rho)$ as a positive linear combination of the form

$$\begin{aligned} \text{tr}(O \otimes O \cdot (P \otimes P) \cdot \text{SWAP}) &= \text{tr}((OP) \otimes (OP) \cdot \text{SWAP}) \\ &= \text{tr}((OP)^2) = \text{tr}(OPOP) = \text{tr}((P^{1/2}OP^{1/2})^2) \geq 0, \end{aligned}$$

where we used the fact that P is positive semidefinite and therefore has a square root. Hence, we again have that $-\text{tr}(O \otimes O \cdot \text{Lower}_\rho)$ is nonpositive, and we are safe to discard it. The bottom line is that when we apply [Theorem 1.4](#), we only have to worry about the first three terms, not the lower order terms, and these are easy to compute and bound.

In fact, we show that there is a whole family of unbiased estimators which one can define based on the measurements in Keyl’s algorithm, of which our debiased Keyl’s algorithm is just one among many. Among these, we show in [Section 3.3](#) that our debiased Keyl’s algorithm is the *minimum variance unbiased estimator*. Note that this is only among those estimators which use Keyl’s measurement and does not preclude an unbiased estimator with even smaller variance which is based on a different measurement.

Discovering the estimator. Before moving on to applications of our debiased Keyl’s algorithm, let us first describe how we discovered the debiased Keyl’s algorithm and proved [Theorems 1.3](#) and [1.4](#). Recall that when ρ is a pure state, the observed λ is always equal to $(n, 0, \dots, 0)$, and Keyl’s algorithm is equivalent to the biased pure state tomography algorithm from [Section 1.1.2](#). As mentioned above, the correction we use to remove the bias in this case is the same as in our pure state tomography algorithm, and it is not too hard to see that this is essentially the *only* way to correct the bias in this case. But this means we are forced to correct the bias in this way whenever we receive the Young diagram $\lambda = (n, 0, \dots, 0)$, because there is a chance that ρ is a pure state, and we have to ensure that our estimator is unbiased in this case. When $n = 2$ and $d = 2$, say, there are only two possible Young diagrams $(2, 0)$ and $(1, 1)$, and the fact that we have already determined what to do in the first case severely limits the possible corrections that we can apply in the second case, which makes finding the right correction much easier. Applying similar logic to other small values of n and d , we solved enough special cases that we were eventually able to spot a pattern that generalized to all values of our parameters. Perhaps the biggest difficulty in this process was deciding on the best way to compute the bias of Keyl’s algorithm; after trying various approaches, we settled on an approach which uses Clebsch-Gordan coefficients. Getting into what exactly Clebsch-Gordan coefficients are is perhaps a bit too technical for this introduction, and we will defer introducing them until [Section 10](#). Suffice to say, they appear to be especially well suited for computing expectations which arise from Keyl’s measurements, although they can be more than a little cumbersome to work with.

Proving [Theorem 1.4](#) was significantly more difficult than discovering the debiased Keyl’s algorithm and is the most technically involved part of this paper. The reason it was so difficult is that we also use Clebsch-Gordan coefficients to compute the variance of $\hat{\rho}$, and going from our large expressions involving the Clebsch-Gordan coefficients to the nice expression in [Theorem 1.4](#) involves some miraculous cancellations which we admit to not having a great understanding of. Proving [Theorem 1.4](#) then required us first having a guess for the precise form of the variance, and only once we had the right guess could we reverse engineer a proof to give us the right calculations. Developing a guess for the variance, in turn, required spending significant amounts of time and effort on numerical simulations, in which we pored over command line outputs and performed numerology in a lengthy and difficult process that we likened to the office jobs in the television show *Severance*. However, given that the main terms in [Theorem 1.4](#) are so nice, and the remaining terms are negative and so can be ignored, we hope that there is an alternative means of computing variances in which these nice expressions fall out naturally, and we leave this for future work.

1.2 Application 1: full state tomography in trace distance

Our first application is to the problem of full state tomography in trace distance. Trace distance is perhaps the most fundamental distance measure between two quantum states, quantifying the extent to which they can be distinguished by a quantum measurement. O’Donnell and Wright showed that $n = O(rd/\varepsilon^2)$ copies suffice for full state tomography of rank- r states in trace distance [[OW16](#)]. In the case when $r = d$, Haah et al. [[HHJ+16](#)] showed a matching lower bound of $n = \Omega(d^2/\varepsilon^2)$ copies, and concurrent work of Scharnhorst, Spilecki, and Wright [[SSW25](#)] has extended this to a matching lower bound of $n = \Omega(rd/\varepsilon^2)$ copies for all values of r . We show that the upper bound of O’Donnell and Wright also holds for the debiased Keyl’s algorithm.

Theorem 1.5 (Trace distance tomography). *Let ρ be rank r , and let $\hat{\rho}$ be the output of the debiased Keyl’s algorithm when run on $\rho^{\otimes n}$. Then $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability when $n = O(rd/\varepsilon^2)$.*

In fact, this result follows immediately from the prior result of O’Donnell and Wright. They showed that if $\hat{\rho}_0$ is the output of Keyl’s algorithm, then $D_{\text{tr}}(\rho, \hat{\rho}_0) \leq \varepsilon$ with high probability. Next, it is easy to see that $D_{\text{tr}}(\hat{\rho}_0, \hat{\rho}) \leq d^2/n$ because $\|\lambda - \lambda^\dagger\|_1 \leq d^2$ always. Thus, by the triangle inequality,

$$D_{\text{tr}}(\rho, \hat{\rho}_0) \leq D_{\text{tr}}(\rho, \hat{\rho}) + D_{\text{tr}}(\hat{\rho}_0, \hat{\rho}) \leq \varepsilon + d^2/n \leq \varepsilon + \varepsilon^2 = O(\varepsilon),$$

where the second-to-last inequality uses the fact that $n = O(d^2/\varepsilon^2)$. However, we are able to give an entirely new proof of this fact which does not rely on this prior result of O’Donnell and Wright. Our proof begins by bounding the expected trace distance between ρ and $\hat{\rho}$ as follows:

$$\mathbf{E}[D_{\text{tr}}(\rho, \hat{\rho})] = \frac{1}{2} \mathbf{E} \|\rho - \hat{\rho}\|_1 \leq \frac{1}{2} \mathbf{E} \sqrt{d} \cdot \|\rho - \hat{\rho}\|_2 \leq \frac{1}{2} \sqrt{d} \cdot \sqrt{\mathbf{E} \|\rho - \hat{\rho}\|_2^2},$$

where the last step uses Jensen's inequality and concavity of the square root function. Now, because $\hat{\rho}$ is an unbiased estimator for ρ , we have $\mathbf{E}[\hat{\rho}] = \rho$, and so $\mathbf{E} \|\rho - \hat{\rho}\|_2^2$ is just the variance of $\hat{\rho}$. This means

$$\begin{aligned} \mathbf{E} \|\rho - \hat{\rho}\|_2^2 &= \mathbf{Var}[\hat{\rho}] = \mathbf{E}[\text{tr}(\hat{\rho}^2)] - \text{tr}(\mathbf{E}[\hat{\rho}]^2) = \mathbf{E}[\text{tr}(\hat{\rho}^2)] - \text{tr}(\rho^2) = \mathbf{E}[\text{tr}(\hat{\rho} \otimes \hat{\rho} \cdot \text{SWAP})] - \text{tr}(\rho^2) \\ &= \text{tr}(\mathbf{E}[\hat{\rho} \otimes \hat{\rho}] \cdot \text{SWAP}) - \text{tr}(\rho^2). \end{aligned}$$

Now we simply substitute our expression for the second moment of $\hat{\rho}$ from [Theorem 1.4](#), and from there bounding the variance is relatively straightforward. In comparison to the tomography proof of [\[OW16\]](#), every step in this proof feels “canonical”, whereas the proof of [\[OW16\]](#) requires a few lucky tricks (for example, the “ $r \geq 2 - \frac{1}{r}$ ” step always seemed like a stroke of luck to the third author) to work.

1.3 Application 2: full state tomography in Bures distance

We also give an application to *fidelity* or *Bures distance* tomography. Bures distance is a more challenging distance metric than trace distance which requires learning the small eigenspace of ρ to high accuracy, and as a result we still lack truly optimal bounds for Bures distance tomography. Haah et al. [\[HHJ+16\]](#) showed that $\tilde{O}(rd/\varepsilon^2)$ samples suffice to learn a rank r mixed state ρ to ε error in Bures distance (equivalently, $\tilde{O}(rd/\delta)$ samples suffice to learn to δ infidelity error). This result has subsequently been reproved in [\[OW17, FO24\]](#), the former of which showed that Keyl's algorithm itself attains this bound. We improve on all of these results and show that Bures distance tomography can be achieved “without the tilde”.

Theorem 1.6 (Bures distance tomography). *Given a rank r mixed state $\rho \in \mathbb{C}^{d \times d}$, $n = O(rd/\varepsilon^2)$ samples suffice to output an estimate $\hat{\rho}$ such that $D_B(\rho, \hat{\rho}) \leq \varepsilon$ with high probability.*

This is optimal due to a matching lower bound of $n = \Omega(rd/\varepsilon^2)$ due to Yuen [\[Yue23\]](#). Indeed, this lower bound also follows from the $n = \Omega(rd/\varepsilon^2)$ lower bound for trace distance tomography from above [\[SSW25\]](#) combined with the fact that Bures distance tomography is at least as hard as trace distance tomography. It is the first time an optimal Bures distance bound without any logs has been shown for either entangled or unentangled measurements.

For intuition, let us sketch the proof of the general $n = O(d^2/\varepsilon^2)$ copy upper bound from [Theorem 1.6](#), which holds when the rank $r = d$. We actually rely on a slight variant of the debiased Keyl's algorithm in order to handle the case when the state ρ is poorly conditioned (i.e. when it has extremely small eigenvalues). To begin, we bound

$$\mathbf{E}[D_B(\rho, \hat{\rho})] \leq \mathbf{E}[\sqrt{D_{\chi^2}(\hat{\rho} \parallel \rho)}] \leq \sqrt{\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho)]},$$

where $D_{\chi^2}(\hat{\rho} \parallel \rho)$ is the *Bures χ^2 -divergence* between $\hat{\rho}$ and ρ . Assuming without loss of generality that ρ is diagonal in the standard basis, i.e. $\rho = \sum_{i=1}^d \alpha_i \cdot |i\rangle\langle i|$, we can express the Bures χ^2 -divergence as

$$\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho)] = \mathbf{E} \left[\sum_{i,j=1}^d \frac{2}{\alpha_i + \alpha_j} \cdot |\hat{\rho}_{ij} - \rho_{ij}|^2 \right] = \sum_{i,j=1}^d \frac{2}{\alpha_i + \alpha_j} \cdot \mathbf{E}[|\hat{\rho}_{ij} - \rho_{ij}|^2].$$

Now, if $\hat{\rho}$ is the output of the debiased Keyl's algorithm, then it is an unbiased estimator for ρ , and so $\mathbf{E}[|\hat{\rho}_{ij} - \rho_{ij}|^2]$ is just the variance $\mathbf{Var}[\hat{\rho}_{ij}]$. To calculate this, we can use our bound on the second moment of $\hat{\rho}$ from [Theorem 1.4](#), which allows us to compute the bound $\mathbf{Var}[\hat{\rho}_{ij}] \leq (\alpha_i + \alpha_j)/n + d/n^2$. Plugging this in above, we have

$$\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho)] \leq \sum_{i,j=1}^d \frac{2}{\alpha_i + \alpha_j} \cdot \left(\frac{\alpha_i + \alpha_j}{n} + \frac{d}{n^2} \right) = \sum_{i,j=1}^d \left(\frac{2}{n} + \frac{2d}{(\alpha_i + \alpha_j)n^2} \right) \leq \frac{2d^2}{n} + \frac{d^3}{\min_i \{\alpha_i\} \cdot n^2}.$$

If we know that ρ is well-conditioned so that each $\alpha_i \geq d/n$, then we can bound this final expression by $3d^2/n$, which is $O(\varepsilon^2)$ when $n = O(d^2/\varepsilon^2)$. This implies a bound of $O(\varepsilon)$ on $\mathbf{E}[D_B(\rho, \hat{\rho})]$, as desired. To ensure that ρ is sufficiently well-conditioned, we can always apply the depolarizing channel to each copy of ρ with noise rate d^2/n prior to running the debiased Keyl's algorithm in order to prepare identical copies of the state

$$\rho' = \left(1 - \frac{d^2}{n}\right) \cdot \rho + \frac{d^2}{n} \cdot \frac{I}{d},$$

which has all eigenvalues at least d/n . We show that this ρ' is still close enough to ρ that it suffices to learn ρ' instead of ρ . One final wrinkle in this proof is that $D_B(\rho, \hat{\rho})$ is only well-defined with $\hat{\rho}$ positive semi-definite, but the debiased Keyl's algorithm in general can output states with negative eigenvalues. However, we show that because ρ' is sufficiently well-conditioned and we are using $n = O(d^2/\varepsilon^2)$ copies, the estimator $\hat{\rho}$ is an actual density matrix (and is therefore positive semidefinite) with high probability. This proof handles the case when $r = d$; to handle the case of small r , we will require a different argument that treats the small eigenvalues of ρ with more care.

As an application of our Bures distance result, we show how to efficiently learn a bipartite pure state given a bound on its Schmidt rank. We thank Benjamin Lovitz for suggesting this application.

Theorem 1.7 (Learning bipartite pure states). *Let $\mathcal{H}_A$ and $\mathcal{H}_B$ be two Hilbert spaces of dimension d_A and d_B , respectively. Then $n = O(r(d_A + d_B)/\varepsilon^2)$ copies of a pure state $|\psi\rangle_{AB}$ with Schmidt rank r are sufficient to learn it to Bures distance ε with high probability.*

1.4 Application 3: tomography with limited entanglement

Suppose we want to perform tomography on ρ , but we can only perform entangled measurements across k copies of ρ at a time. This question is motivated by the fact that implementing large, entangled measurements can be challenging in practice; indeed, one often simply does not have enough qubits to store all n copies of ρ at once! In this case, the natural strategy is to do the following.

1. Assuming for simplicity that n is a multiple of k , divide the n copies of ρ into $n' := n/k$ disjoint blocks.
2. Run a k -copy tomography algorithm on each block. Let $\hat{\rho}_1, \dots, \hat{\rho}_{n'}$ be the resulting estimators.
3. Output $\hat{\rho} = \frac{1}{n'} \cdot (\hat{\rho}_1 + \dots + \hat{\rho}_{n'})$.

The estimators $\hat{\rho}_1, \dots, \hat{\rho}_{n'}$ are independent and identically distributed, which in particular means that they have the same expectation $\mathbf{E}[\hat{\rho}_1] = \dots = \mathbf{E}[\hat{\rho}_{n'}]$. By linearity of expectation, the averaged estimator $\hat{\rho}$ has the same expectation as well. However, even though it has the same expectation, it will concentrate around this expectation much more strongly than the individual estimators do. This can be seen by considering its variance and recalling that the variance of independent random variables is equal to the sum of their variances:

$$\mathbf{Var}[\hat{\rho}] = \frac{1}{(n')^2} \cdot \mathbf{Var}[\hat{\rho}_1 + \dots + \hat{\rho}_{n'}] = \frac{1}{(n')^2} \cdot (\mathbf{Var}[\hat{\rho}_1] + \dots + \mathbf{Var}[\hat{\rho}_{n'}]) = \frac{1}{(n')} \cdot \mathbf{Var}[\hat{\rho}_{n'}].$$

If the tomography algorithm that is used produces an unbiased estimator, then this means that the estimator $\hat{\rho}$ will concentrate well around ρ when n' is large. As an example, in the case when $k = 1$, the measurements must be unentangled across all n copies of ρ , and the natural tomography algorithm to use is the single-copy unbiased estimator from [Section 1.1.1](#). In this case, the averaged estimator $\hat{\rho}$ is known as the “uniform POVM tomography algorithm” and satisfies $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability once $n = O(d^3/\varepsilon^2)$, which was shown independently by Krishnamurthy and Wright [[Wri16](#), Section 5.1] and Guta et al. [[GKKT20](#)].

However, if the tomography algorithm that is used does *not* produce an unbiased estimator, then although $\hat{\rho}$ will still concentrate well around the expectation of the $\hat{\rho}_i$'s, this will no longer be useful if this expectation is far from ρ . This was the issue faced by Chen, Li, and Liu in [[CLL24b](#)], who studied the case when the algorithm is allowed to make entangled measurements across $k > 1$ copies at a time. In spite of the fact that they did not have an unbiased, entangled tomography algorithm at their disposal, they were still able to prove nearly matching upper and lower bounds for this problem. In particular, when $k \leq \min\{d^2, (\sqrt{d}/\varepsilon)^c\}$, where c is some small constant, they showed that $n = \tilde{O}(d^3/(\sqrt{k}\varepsilon^2))$ copies of ρ suffice to learn ρ to error ε in trace distance. Interestingly, this suggests that increasing the amount of entanglement k helps only until $k = d^2$, and further increases in k do not improve the copy complexity. They also showed that $n = \Omega(d^3/(\sqrt{k}\varepsilon^2))$ copies are necessary to learn ρ to error ε in trace distance if one performs even adaptively chosen measurements across k copies of ρ at a time, at least when $k \leq 1/\varepsilon^c$, where c is some small constant. Finally, they suggested that these restrictions on the parameter k in both their upper and lower bounds are artifacts of their proof techniques, and they conjectured that the optimal number of copies for this task is $n = \tilde{\Theta}(d^3/(\sqrt{k}\varepsilon^2))$ when $k \leq d^2$.

For our third application of the debiased Keyl’s algorithm, we show that their conjectured upper bound is attainable even “without the tilde”.

Theorem 1.8 (Tomography with limited entanglement). *Let $n = n' \cdot k$. Let $\hat{\rho}_1, \dots, \hat{\rho}_{n'}$ be the outputs of the debiased Keyl’s algorithm when run in n' independent copies of $\rho^{\otimes k}$. Set $\hat{\rho} = (\hat{\rho}_1 + \dots + \hat{\rho}_{n'})/n'$. Then $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability when*

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

By the lower bound of Chen, Li, and Liu [CLL24b], this bound is optimal, at least for all $k \leq 1/\varepsilon^c$. As they conjecture, we believe that our bound is optimal for all k , and we leave this question for future work.

1.5 Application 4: shadow tomography

Given m observables $O_1, \dots, O_m \in \mathbb{C}^{d \times d}$, *shadow tomography* refers to the task of producing estimates of the observable values $\text{tr}(O_1 \cdot \rho), \dots, \text{tr}(O_m \cdot \rho)$ given access to copies of ρ . First introduced by Aaronson in [Aar18], this problem occupies a central place in the field of quantum learning theory. To date, the best algorithm for general bounded observables (satisfying $\|O_i\|_\infty \leq 1$) is due to Bădescu and O’Donnell [BO24], who showed that $n = \tilde{O}(\log^2(m) \log(d)/\varepsilon^4)$ copies suffice (see also the work of Bene Watts and Bostanci [BWB24], who reproved this bound with a different algorithm). On the flip side, the only known lower bound is due to Aaronson, who showed that $n = \Omega(\log(m)/\varepsilon^2)$ copies are necessary.

One natural approach for the shadow tomography problem is the “plug-in estimator”: simply run a full state tomography algorithm to produce an estimate $\hat{\rho}$ of ρ and “plug it in” to compute the observable values $\hat{o}_1 := \text{tr}(O_1 \cdot \hat{\rho}), \dots, \hat{o}_m := \text{tr}(O_m \cdot \hat{\rho})$. For example, if $\hat{\rho}$ is the output of Keyl’s algorithm on n copies of ρ , then each $\hat{o}_i$ will be within ε of the true observable value $\text{tr}(O_i \cdot \rho)$ once $n = O(d^2/\varepsilon^2)$, no matter how many observables m there are. This is because with this many samples, Keyl’s algorithm produces an estimate $\hat{\rho}$ with $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability, and being ε close in trace distance is equivalent to every observable value being correct within ε . However, if we use a smaller number of copies $n \ll d^2/\varepsilon^2$, then each $\hat{o}_i$ might be far from the corresponding $\text{tr}(O_i \cdot \rho)$, as Keyl’s algorithm does not give guarantees in this regime. Keyl’s algorithm fails in the “sub- d^2/ε^2 regime” for a more fundamental reason, which is that the estimator $\hat{\rho}$ it produces is heavily biased away from ρ in this regime, which means that $\text{tr}(O_i \cdot \hat{\rho})$ will be far from $\text{tr}(O_i \cdot \rho)$ for some observables O_i , at least in expectation. This suggests using an unbiased estimator for ρ in place of Keyl’s algorithm; in this case, $\hat{o}_1, \dots, \hat{o}_m$ will be unbiased estimators for $\text{tr}(O_1 \cdot \rho), \dots, \text{tr}(O_m \cdot \rho)$, respectively, and they may approximate these values with fewer copies of ρ than it takes for $\hat{\rho}$ to approximate all of ρ . We note that this “plug-in estimator” approach is one method for solving the more general *oblivious shadow tomography* problem, in which one has to perform all of one’s measurements on $\rho^{\otimes n}$ prior to learning the observables $O_1, \dots, O_m$ to be estimated. (The algorithm of Bădescu and O’Donnell, for example, is not oblivious, as the measurements it performs depend on the observables $O_1, \dots, O_m$.)

Huang, Kueng, and Preskill [HKP20] studied this “plug-in estimator” approach for the case when $\hat{\rho}$ is the output of the uniform POVM tomography algorithm on n' copies of ρ . In this case, they showed that each $\text{tr}(O_i \cdot \hat{\rho})$ is an unbiased estimator for $\text{tr}(O_i \cdot \rho)$ with variance $O(\text{tr}(O_i^2)/n') \leq O(F/n')$, where $F = \max\{\text{tr}(O_1^2), \dots, \text{tr}(O_m^2)\}$. By setting $n' = O(F/\varepsilon^2)$, each $\text{tr}(O_i \cdot \hat{\rho})$ will have variance $O(\varepsilon^2)$ and will therefore be equal to $\text{tr}(O_i \cdot \rho) \pm O(\varepsilon)$ with high probability. Suppose we then repeat this process k times, generating observable estimates $\hat{o}_1^1, \dots, \hat{o}_m^1$ through $\hat{o}_1^k, \dots, \hat{o}_m^k$. If we then set $\hat{o}_i = \text{median}\{\hat{o}_i^1, \dots, \hat{o}_i^k\}$ for all $1 \leq i \leq m$, then $\hat{o}_i$ will be equal to $\text{tr}(O_i \cdot \rho) \pm \varepsilon$ for all $1 \leq i \leq m$ with high probability, once $k = O(\log(m))$. In total, this uses $n = n' \cdot k = O(\log(m)F/\varepsilon^2)$ copies of ρ .¹ Note that if each O_i is bounded, with $\|O_i\|_\infty \leq 1$, then F can be as large as d , in which case this approach requires $n = O(\log(m)d/\varepsilon^2)$ copies of ρ , which is better than the $O(d^2/\varepsilon^2)$ required by Keyl’s algorithm but far from the copy complexities of the best algorithms. If, on the other hand, the observables are all rank-1, then $F \leq 1$, and so this algorithm

¹We note that the step of taking the median of multiple observable values is necessitated because Huang, Kueng, and Preskill also studied a variant of the uniform POVM tomography algorithm where the uniform POVM is replaced by a more computationally efficient measurement based on Clifford unitaries. Indeed, Helsen, Helsen, and Walter [HW23] showed that if $\hat{\rho}$ is the output of the uniform POVM tomography algorithm on $n = O(\log(m)F/\varepsilon^2)$ copies of ρ , then $\text{tr}(O_i \cdot \hat{\rho}) = \text{tr}(O_i \cdot \rho) \pm \varepsilon$ for all $1 \leq i \leq m$, no medians required.

only needs $n = O(\log(m)/\varepsilon^2)$ copies of ρ . This is the same number of copies as Aaronson’s lower bound from above, although strictly speaking Aaronson’s lower bound holds against general and not just rank-1 observables and therefore does not apply. Huang, Kueng, and Preskill called their algorithm the “classical shadows algorithm”, and as a result the more broad problem of oblivious shadow tomography is sometimes called the “classical shadows problem”.

Grier, Pashayan, and Schaeffer [GPS24] studied this “plug-in estimator” approach for the special case when ρ is a pure state and $\hat{\rho}$ is the output of the bias-corrected pure state algorithm from Section 1.1.2. Following a similar analysis as [HKP20], they showed that

$$n = O\left(\log(m) \cdot \left(\frac{\sqrt{F}}{\varepsilon} + \frac{1}{\varepsilon^2}\right)\right)$$

copies of ρ suffice for shadow tomography. This sample complexity bound splits into two regimes, based on whether ε is large or small: when $\varepsilon \geq 1/\sqrt{F}$, their sample complexity scales as $O(\log(m)\sqrt{F}/\varepsilon)$, and when $\varepsilon \leq 1/\sqrt{F}$, their sample complexity scales as $O(\log(m)/\varepsilon^2)$. Interestingly, in the latter regime of small ε , sometimes known as the “high accuracy regime” [CLL24a], their bound matches the lower bound of $n = \Omega(\log(m)/\varepsilon^2)$ copies due to Aaronson, although again here the Aaronson bound does not necessarily apply since his lower bound is shown for general states ρ and not just for pure states. However, their sample complexity is optimal for the classical shadows problem, as Huang, Kueng, and Preskill [HKP20] showed that

$$n = \Omega\left(\log(m) \cdot \left(\frac{\sqrt{d}}{\varepsilon} + \frac{1}{\varepsilon^2}\right)\right)$$

copies are necessary for the classical shadows problem, even in the special case when ρ is pure.

Subsequent work has attempted to generalize the Grier, Pashayan, and Schaeffer result to the case when ρ is a general mixed state. For example, the work of Grier, Liu, and Mahajan [GLM24] gave an algorithm for the classical shadows problem which uses $n = O(\log(m) \cdot \sqrt{rF}/\varepsilon^2)$ copies in the case when ρ is rank r . For small values of r , this improves on the sample complexity of the result of [HKP20] by a factor of $\sqrt{F}$; however, it does not achieve a sample complexity of $O(\log(m)/\varepsilon^2)$ in the high accuracy regime of small ε . On the flip side, Chen, Li, and Liu [CLL24a] focused exclusively on the high accuracy regime and showed that in the case when ρ is allowed to have arbitrary rank, $n = O(\log(m)/\varepsilon^2)$ copies are sufficient for the classical shadows problem, at least in the high accuracy regime of $\varepsilon = O(1/d^{12})$. As their result holds for general states and observables, the $n = \Omega(\log(m)/\varepsilon^2)$ lower bound of Aaronson finally applies, meaning that theirs is the first provably optimal algorithm for the general shadow tomography problem, for any regime of ε . This implies, perhaps surprisingly, that in the high accuracy regime, knowing the observables in advance does not actually help when performing shadow tomography. More broadly, they conjectured that

$$n = \Theta\left(\log(m) \cdot \left(\frac{d}{\varepsilon} + \frac{1}{\varepsilon^2}\right)\right)$$

should be the optimal sample complexity for the classical shadows problem on general states and observables.

We study the “plug-in estimator” approach in the case when $\hat{\rho}$ is the output of the debiased KeyL’s algorithm. Our main result improves on the sample complexity of [GLM24] in the low accuracy regime, improves the threshold for the high accuracy regime in [CLL24a] from $\varepsilon = O(1/d^{12})$ to $\varepsilon = O(1/d)$, and even refutes the conjecture of [CLL24a] by improving on their conjectured optimal bound.

Theorem 1.9 (Classical shadows). *The classical shadows task for states of rank r and observables $O_1, \dots, O_m$ with $\text{tr}(O_i^2) \leq F$ can be solved using*

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

copies of ρ .

To understand our sample complexity bound, let us consider two special cases. In the $r = 1$ case of pure states, $r = 1$, we have that $F^{1/2}/\varepsilon$ is less than $F^{2/3}/\varepsilon^{4/3} = (F^{1/2}/\varepsilon)^{4/3}$ unless $F^{1/2}/\varepsilon < 1$, in which case the

first term is trivial and can be discarded. This means the “min” term is just $\sqrt{F}/\varepsilon$ and the overall sample complexity is

$$n = O\left(\log(m) \cdot \left(\frac{\sqrt{F}}{\varepsilon} + \frac{1}{\varepsilon^2}\right)\right),$$

matching the bound of Grier, Pashayan, and Schaeffer [GPS24]. On the flip side, in the fully general case of $r = F = d$, the bound becomes

$$n = O\left(\log(m) \cdot \left(\min\left\{\frac{d}{\varepsilon}, \frac{d^{2/3}}{\varepsilon^{4/3}}\right\} + \frac{1}{\varepsilon^2}\right)\right).$$

Here, the “min” term is equal to $d^{2/3}/\varepsilon^{4/3}$ for $\varepsilon \geq 1/d$ and d/ε for $\varepsilon \leq 1/d$. But in the latter case, the $1/\varepsilon^2$ term dominates both of these terms, meaning that we can simplify the sample complexity to be

$$n = O\left(\log(m) \cdot \left(\frac{d^{2/3}}{\varepsilon^{4/3}} + \frac{1}{\varepsilon^2}\right)\right),$$

which improves on the conjectured upper bound of [CLL24a] in the low accuracy regime. We conjecture that our upper bound is optimal, and leave the task of proving matching lower bounds to future work.

1.6 Application 5: quantum metrology

Quantum metrology is the study of optimal measurements for estimating parameters of quantum systems. A parameterized quantum state is a density matrix ρ_θ which depends smoothly on some parameters $\theta = (\theta_1, \dots, \theta_m) \in \Theta \subseteq \mathbb{R}^m$. There is a fixed state ρ_{θ^*} whose parameters $\theta^* \in \Theta$ one is assumed to already know. Now, one is given a state ρ_θ , where θ is “in the neighborhood” of θ^* , and the goal is to estimate the parameters θ . This models, for example, a common strategy for learning an unknown parameterized quantum channel Φ_θ . To do this, one might initialize a starting “probe state” ρ_{probe} , send it through the channel Φ_θ , and measure the resulting state ρ_θ to learn θ . The better a measurement one uses, the better an estimate of Φ_θ one can produce. Furthermore, one is assumed to already have prior knowledge that θ is close to some parameter vector θ^* , and one is allowed to use this knowledge to design one’s measurements.

To formalize this question, we focus on *locally unbiased estimators*. As we have already seen, an estimator is an algorithm which takes as input a copy of ρ_θ and produces a (randomly distributed) estimate $\hat{\theta}$ of θ , and it is unbiased if $\mathbf{E}[\hat{\theta}] = \theta$. Being locally unbiased is a weaker condition: it requires only that

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

The first of these conditions states that $\hat{\theta}$ is indeed an unbiased estimator, but just on the state ρ_{θ^*} . The second of these conditions states that $\hat{\theta}$ remains unbiased, at least up to first order, in a small neighborhood around θ^* . Given a locally unbiased estimator, the figure of merit we associate with it is its *mean squared error matrix (MSEM)* $V \in \mathbb{R}^{m \times m}$, defined as

$$V_{ij} = \mathbf{E}[(\hat{\theta}_i - \theta_i^*) \cdot (\hat{\theta}_j - \theta_j^*) \mid \rho_{\theta^*}],$$

where $\hat{\theta}$ is the output of the estimator on the state ρ_{θ^*} .

The most well-known lower bound on V is in terms of the *Quantum Fisher Information (QFI)* matrix $\mathcal{F}$. To define $\mathcal{F}$, we first define the *symmetric logarithm derivative* matrices $L_1, \dots, L_m$ implicitly via the equations

$$\frac{\partial}{\partial \theta_i} \rho_\theta|_{\theta=\theta^*} = \frac{1}{2} L_i \rho_{\theta^*} + \frac{1}{2} \rho_{\theta^*} L_i, \quad \text{for all } 1 \leq i \leq m.$$

These equations do not actually uniquely specify $L_1, \dots, L_m$, but any matrices which satisfy these equations will lead to the same QFI matrix $\mathcal{F}$. Given these, we define

$$\mathcal{F}_{ij} = \frac{1}{2} \text{tr}(\rho_{\theta^*} L_i L_j) + \frac{1}{2} \text{tr}(\rho_{\theta^*} L_j L_i).$$

The *quantum Cramér–Rao bound (QCRB)*, shown by Helstrom in [Hel67], states that

$$V \succeq \mathcal{F}^{-1}.$$

This places a lower bound on the variance of any locally unbiased estimator for θ .

Is the QCRB aqually achievable? The answer is yes, when $m = 1$, given asymptotically many copies of ρ_θ . To explain this, suppose we are given n copies of ρ_θ rather than one. This can be viewed as being given one copy of the parameterized state $\rho_\theta^n := \rho_\theta^{\otimes n}$. If $\mathcal{F}$ is the QFI matrix of the ρ_θ 's and $\mathcal{F}_n$ is the QFI matrix of the ρ_θ^n 's, they are related via the equality $\mathcal{F}_n = n \cdot \mathcal{F}$. Hence, if V_n is the MSEM of a locally unbiased estimator for the ρ_θ^n 's, then the QCRB implies that

$$V_n \succeq \mathcal{F}_n^{-1} = \frac{1}{n} \cdot \mathcal{F}^{-1}, \quad (2)$$

or, equivalently, $n \cdot V_n \succeq \mathcal{F}^{-1}$. We say that the QCRB is achievable asymptotically if there is a sequence of such locally unbiased estimators, one for each value n , such that $n \cdot V_n \rightarrow \mathcal{F}^{-1}$ as $n \rightarrow \infty$. In this case, V_n is equal to $(1/n) \cdot \mathcal{F}^{-1}$, plus other terms of order $o(1/n)$ which, albeit possibly large for small n , disappear as n increases. It was shown by Braunstein and Caves [BC94] that the QCRB is indeed achievable asymptotically in the $m = 1$ case of a single unknown parameter.

When $m > 1$, the problem is more challenging, and there are examples of multiparameter estimation problems for which the QCRB is not achievable, even asymptotically, the key issue being the noncommutativity of different quantum measurements. To simplify the algorithm's task, let us introduce a positive semidefinite cost matrix $C \in \mathbb{R}^{m \times m}$ and assign to the estimator the cost

$$\text{tr}(C \cdot V) = \sum_{i,j=1}^m C_{ij} \cdot \mathbf{E}[(\hat{\theta}_i - \theta_i^*) \cdot (\hat{\theta}_j - \theta_j^*) \mid \rho_{\theta^*}].$$

Given this cost matrix, the QCRB implies that

$$\text{tr}(C \cdot V) \geq \text{tr}(C \cdot \mathcal{F}^{-1}) \quad (3)$$

for any locally unbiased estimator. Achieving the QCRB in Equation (2) is equivalent to achieving the bound in Equation (3) for all cost matrices C simultaneously; hence, achieving the bound in Equation (3) for just a single cost matrix C is potentially an easier problem, as perhaps one can hand-tailor the locally unbiased estimator for this particular cost matrix. As it turns out, the QCRB with a specific cost matrix C is still not achievable, even asymptotically.

However, there is a slight strengthening of Equation (3) which is achievable asymptotically. In [Hol11], Holevo introduced a quantity $\text{Hol}(C)$ and showed that $\text{tr}(C \cdot V) \geq \text{Hol}(C)$, an inequality known as the *Holevo Cramér–Rao bound (HCRB)*. For our purposes, it is not important what $\text{Hol}(C)$ is exactly, only that it satisfies two properties. The first property is that

$$2 \cdot \text{tr}(C \cdot \mathcal{F}^{-1}) \geq \text{Hol}(C) \geq \text{tr}(C \cdot \mathcal{F}^{-1}),$$

and so the HCRB is only stronger than the QCRB by at most a factor of 2. The first inequality was shown by [ATD19, CSDV19] and is known to be tight in some cases; see [DDGG20, Section 3.1.1] for a simple example. The second property we will need is that if $\text{Hol}_n(C)$ is the Holevo bound corresponding to the n copy states $\rho_\theta^n = \rho_\theta^{\otimes n}$, then, similarly to the QFI matrices, we have $\text{Hol}(C) = n \cdot \text{Hol}_n(C)$. Thus, the HCRB implies that

$$\text{tr}(C \cdot V_n) \geq \text{Hol}_n(C) = \frac{1}{n} \cdot \text{Hol}(C).$$

We say that the HCRB is achievable asymptotically if $n \cdot \text{tr}(C \cdot V_n) \rightarrow \text{Hol}(C)$ as $n \rightarrow \infty$. A line of work studying quantum local asymptotic normality [KG09, YFG13, YCH19] has developed estimators which do achieve the HCRB asymptotically. This means that in the setting of asymptotically many copies, the HCRB is tight.

This paper is about designing unbiased estimators for full state tomography, and naively it seems like having good unbiased estimators for quantum tomography would help in creating good locally unbiased

estimators for quantum metrology. Motivated by this reasoning, a recent work of Zhou and Chen [ZC25] studied a natural locally unbiased estimator for quantum metrology which first computes the $n = 1$ copy unbiased estimator $\hat{\rho}$ from Section 1.1.1 and then performs classical postprocessing on the outcome. They showed that in the case when the parameterized states ρ_θ are all rank-one, the MSEM V of their estimator satisfies

$$V = \frac{4(d+1)}{d+2} \cdot \mathcal{F}^{-1} \preceq 4 \cdot \mathcal{F}^{-1}. \quad (4)$$

As $V \succeq \mathcal{F}^{-1}$, this shows that V is within a factor of 4 of being optimal; they referred to such estimators whose MSEM is within a constant factor of $\mathcal{F}^{-1}$ as *near optimal*. Now, it was already known that the HCRB was attainable for rank one states ρ_θ due to the work of Matsumoto [Mat02], even in the non-asymptotic regime where only one copy of ρ_θ is provided. However, Zhou and Chen’s estimator has the advantage that because it provides an approximation to $\mathcal{F}^{-1}$ rather than just the $\text{Hol}(C)$, it is oblivious to the cost matrix C , in the sense that the measurement is independent of the cost matrix C (in fact, it is independent of the parameterized state family ρ_θ altogether), and so one can measure first and be provided C later. (This is similar to the sense in which classical shadows algorithms are oblivious.) They also generalized their results to hold in the case when the state family ρ_θ is low rank rather than just rank one, but their guarantees degrade as ρ_{θ^*} becomes more and more poorly conditioned.

What about when more than one copy of ρ_θ is available? As we have better unbiased estimators for larger n , one might hope that we can use these to design better locally unbiased estimators for quantum metrology. To study this, we consider the locally unbiased estimator we get when we replace the unbiased estimator Zhou and Chen use with the n -copy debiased Keyl’s algorithm, but keep the classical post-processing the same. We show the following result.

Theorem 1.10 (Quantum metrology). *Let ρ_θ be a parameterized state family with probe state*

$$\rho_{\theta^*} = \sum_{i=1}^d \alpha_i \cdot |v_i\rangle\langle v_i|,$$

where we assume without loss of generality that the eigenvalues are sorted so that $\alpha_1 \geq \dots \geq \alpha_r > 0$ and $\alpha_{r+1} = \dots = \alpha_d = 0$ (so that ρ_{θ^} is rank r). Then for every n , there is a locally unbiased estimator which takes as input n copies of ρ_θ and whose MSEM V_n has the following property. When $n = r/(\varepsilon \cdot \alpha_r)$,*

$$V_n \preceq (1 + \varepsilon) \cdot \frac{2}{n} \cdot \mathcal{F}^{-1}.$$

Hence, as $n \rightarrow \infty$, we have $n \cdot V_n \rightarrow 2 \cdot \mathcal{F}^{-1}$, and so this estimator achieves twice the QCRB asymptotically.

As a result, for asymptotically large n and any cost matrix C , we have that $n \cdot \text{tr}(C \cdot V_n) \rightarrow 2 \cdot \text{tr}(C \cdot \mathcal{F}^{-1}) \leq 2 \cdot \text{Hol}(C)$. This implies the surprising fact that, in the asymptotic limit, knowing the cost matrix C ahead of time only buys you at most a factor of 2 in the final cost of the estimator. This is the optimal result of this form that one could hope for. To see this, recall that there are examples in which $\text{Hol}(C) = 2 \cdot \text{tr}(C \cdot \mathcal{F}^{-1})$. In these cases, we have

$$\text{tr}(C \cdot V_n) \geq \text{Hol}_n(C) = \frac{1}{n} \cdot \text{Hol}(C) = \frac{1}{n} \cdot 2 \cdot \text{tr}(C \cdot \mathcal{F}^{-1}).$$

Hence, $n \cdot \text{tr}(C \cdot V_n) \geq 2 \cdot \text{tr}(C \cdot \mathcal{F}^{-1})$ for all n in such cases, and so one cannot have $n \cdot V_n \preceq \beta \cdot \mathcal{F}^{-1}$ for any constant $\beta < 2$. Finally, we note that when ρ_{θ^*} is rank $r = 1$, then $\alpha_1 = 1$, and so if we apply Theorem 1.10 with $n = 1$, and therefore with $\varepsilon = 1$, we get $V_1 \preceq 4 \cdot \mathcal{F}^{-1}$, which recovers Zhou and Chen’s bound from Equation (4).

1.7 Technical overview

The key technical challenge that we overcome in this work is understanding how to compute exact expressions for the first and second moments $\mathbf{E}[\hat{\rho}]$ and $\mathbf{E}[\hat{\rho} \otimes \hat{\rho}]$ of our representation-theoretic estimators. What makes this difficult is that to do so, one must understand the probability distributions on the Young diagram λ and

the unitary U produced by Keyl's measurement, but this distribution involves complicated representation-theoretic formulas which are unwieldy to manipulate. Indeed, although the two works [OW16, HHJ⁺16] were able to analyze three separate entangled tomography algorithms and show that they are either sample-optimal or almost sample-optimal, neither work was able to derive expressions for even the first moment of their estimators.

1.7.1 The single copy case

To illustrate our techniques, let us consider the simplest possible case, when we are given only a single copy of ρ . In this case, it turns out that Keyl's measurement always returns the Young diagram $\lambda = \square$. In addition, the unitary U it returns has probability density function

$$\Pr[U = U \mid \rho] = d \cdot \langle 1 | U^\dagger \rho U | 1 \rangle \cdot dU. \quad (5)$$

Since $\lambda^\dagger = (d, -1, \dots, -1)$, our debiased Keyl's algorithm will output the estimator

$$\hat{\rho} = d \cdot |\mathbf{u}_1\rangle\langle\mathbf{u}_1| - |\mathbf{u}_2\rangle\langle\mathbf{u}_2| - \dots - |\mathbf{u}_d\rangle\langle\mathbf{u}_d|,$$

where $|\mathbf{u}_i\rangle := U \cdot |i\rangle$. Intuitively, $|\mathbf{u}_1\rangle$ should be a rough approximation of the largest eigenvector of ρ , or should at least have decent overlap with ρ 's larger eigenvectors. This is because the PDF in Equation (5) scales with $\langle 1 | U^\dagger \rho U | 1 \rangle$, and is therefore maximized for any unitary U such that $U \cdot |1\rangle$ is the largest eigenvector of ρ . As we have already pointed out in Section 1.1.4, this is equivalent to the uniform POVM tomography algorithm from Section 1.1.1.

How to compute the expectation $\mathbf{E}[\hat{\rho}]$ of this estimator? It turns out that for the uniform POVM tomography algorithm, there is a slick, two-line proof that $\mathbf{E}[\hat{\rho}] = \rho$ [Wri16, Section 5.1]. Instead, we will describe a significantly more cumbersome version of this argument, as it is this version that we were able to extend to the case of larger n . To begin, we can model a generic tomography algorithm by a POVM $M = \{M_{\hat{\rho}}\}_{\hat{\rho}}$, so that the algorithm measures $\rho^{\otimes n}$ with M , receives an outcome $\hat{\rho}$, and then outputs $\hat{\rho}$ as its estimator. Then the expectation can be written as

$$\mathbf{E}[\hat{\rho}] = \sum_{\hat{\rho}} \text{tr}(M_{\hat{\rho}} \cdot \rho^{\otimes n}) \cdot \hat{\rho} = \text{tr}_{[n]} \left(\left(\sum_{\hat{\rho}} M_{\hat{\rho}} \otimes \hat{\rho} \right) \cdot \rho^{\otimes n} \otimes I \right) =: \text{tr}_{[n]}(M_{\text{avg}} \cdot \rho^{\otimes n} \otimes I).$$

It therefore suffices to compute a formula for the matrix M_{avg} . When $n = 1$ and the algorithm is the debiased Keyl's algorithm, this matrix is

$$M_{\text{avg}} = \int_U \left(d \cdot U |1\rangle\langle 1| U^\dagger \cdot dU \right) \otimes \sum_{i=1}^d \lambda_i^\dagger \cdot U |i\rangle\langle i| U^\dagger = d \cdot \sum_{i=1}^d \lambda_i^\dagger \cdot \int_U U^{\otimes 2} |1, i\rangle\langle 1, i| U^{\dagger, \otimes 2} \cdot dU.$$

The symmetric and antisymmetric subspaces. As a first step, let us compute a formula for the integral

$$\int_U U^{\otimes 2} |1, i\rangle\langle 1, i| U^{\dagger, \otimes 2} \cdot dU. \quad (6)$$

This is relatively simple in the case when $i = 1$, because the vector $|1, 1\rangle$ is an element of the *symmetric subspace*

$$V_{\square} := \{|\psi\rangle \in (\mathbb{C}^d)^{\otimes 2} \mid \text{SWAP} \cdot |\psi\rangle = |\psi\rangle\} = \text{span}\{|v\rangle \otimes |v\rangle \mid |v\rangle \in \mathbb{C}^d\}.$$

The symmetric subspace is an irreducible representation of the unitary group, which implies that for any vector $|\psi_{\square}\rangle \in V_{\square}$,

$$\int_U U^{\otimes 2} |\psi_{\square}\rangle\langle\psi_{\square}| U^{\dagger, \otimes 2} \cdot dU = \frac{\Pi_{\square}}{\dim(V_{\square})}, \quad (7)$$

where $\Pi_{\square}$ is the projector onto $V_{\square}$. In particular, this applies for $|\psi_{\square}\rangle = |1, 1\rangle$, which gives us a formula for Equation (6) when $i = 1$.

When $i > 1$, however, $|1, i\rangle\langle 1, i|$ is not entirely contained inside the symmetric subspace. Instead, it is split between the symmetric subspace and the *antisymmetric subspace*, defined as

$$V_{\square} := \{|\psi\rangle \in (\mathbb{C}^d)^{\otimes 2} \mid \text{SWAP} \cdot |\psi\rangle = -|\psi\rangle\}.$$

Like the symmetric subspace, the antisymmetric subspace is an irreducible representation of the unitary group, which implies that for any vector $|\psi_{\square}\rangle \in V_{\square}$,

$$\int_U U^{\otimes 2} |\psi_{\square}\rangle \langle \psi_{\square}| U^{\dagger, \otimes 2} \cdot dU = \frac{\Pi_{\square}}{\dim(V_{\square})}, \quad (8)$$

where $\Pi_{\square}$ is the projector onto $V_{\square}$. In addition, because the antisymmetric subspace is non-isomorphic to the symmetric subspace, a result in representation theory known as Schur's lemma states that for any vectors $|\psi_{\square\square}\rangle \in V_{\square\square}$ and $|\psi_{\square}\rangle \in V_{\square}$,

$$\int_U U^{\otimes 2} |\psi_{\square\square}\rangle \langle \psi_{\square}| U^{\dagger, \otimes 2} \cdot dU = 0. \quad (9)$$

Hence, if we could simply split $|1, i\rangle$ into its symmetric and antisymmetric components, we could use these formulas to evaluate Equation (6).

To this end, let us define the vectors

$$|\boxed{i}\boxed{i}\rangle := |i, i\rangle, \quad |\boxed{i}\boxed{j}\rangle := \frac{1}{\sqrt{2}} |i, j\rangle + \frac{1}{\sqrt{2}} |j, i\rangle, \quad \text{and} \quad \left| \begin{smallmatrix} \boxed{i} \\ \boxed{j} \end{smallmatrix} \right\rangle := \frac{1}{\sqrt{2}} |i, j\rangle - \frac{1}{\sqrt{2}} |j, i\rangle, \quad \text{for } i < j.$$

Note that the $|\boxed{i}\boxed{j}\rangle$ vectors are elements of the symmetric subspace, and the $\left| \begin{smallmatrix} \boxed{i} \\ \boxed{j} \end{smallmatrix} \right\rangle$ vectors are elements of the antisymmetric subspace. In addition, they are orthogonal to each other by inspection, and in fact they actually form an orthonormal basis for all of $(\mathbb{C}^d)^{\otimes 2}$, as we can write the standard basis vectors as linear combinations of these vectors:

$$|i, i\rangle = |\boxed{i}\boxed{i}\rangle, \quad |i, j\rangle = \frac{1}{\sqrt{2}} |\boxed{i}\boxed{j}\rangle + \frac{1}{\sqrt{2}} \left| \begin{smallmatrix} \boxed{i} \\ \boxed{j} \end{smallmatrix} \right\rangle, \quad \text{and} \quad |j, i\rangle = \frac{1}{\sqrt{2}} |\boxed{i}\boxed{j}\rangle - \frac{1}{\sqrt{2}} \left| \begin{smallmatrix} \boxed{i} \\ \boxed{j} \end{smallmatrix} \right\rangle, \quad \text{for } i < j. \quad (10)$$

From this fact, we can derive several useful consequences.

1. Because these vectors form an orthonormal basis for all of $(\mathbb{C}^d)^{\otimes 2}$, they also form orthonormal bases for their respective subspaces. Thus, we can conclude that $\dim(V_{\square\square}) = d(d+1)/2$ and $\dim(V_{\square}) = d(d-1)/2$.
2. Similarly, this means that the symmetric and antisymmetric subspaces together span all of $(\mathbb{C}^d)^{\otimes 2}$. In other words, $(\mathbb{C}^d)^{\otimes 2} = V_{\square\square} \oplus V_{\square}$.
3. Finally, since these vectors form an orthonormal basis for all of $(\mathbb{C}^d)^{\otimes 2}$, they also form a complete eigenbasis for SWAP, with eigenvalues $+1$ or -1 . More succinctly, $\text{SWAP} = \Pi_{\square\square} - \Pi_{\square}$.

We can now use these vectors to compute our desired expectation. In particular, when $i > 1$,

$$\begin{aligned} \int_U U^{\otimes 2} |1, i\rangle \langle 1, i| U^{\dagger, \otimes 2} \cdot dU &= \frac{1}{2} \int_U U^{\otimes 2} |\boxed{1}\boxed{i}\rangle \langle \boxed{1}\boxed{i}| U^{\dagger, \otimes 2} \cdot dU + \frac{1}{2} \int_U U^{\otimes 2} \left| \begin{smallmatrix} \boxed{1} \\ \boxed{i} \end{smallmatrix} \right\rangle \left\langle \begin{smallmatrix} \boxed{1} \\ \boxed{i} \end{smallmatrix} \right| U^{\dagger, \otimes 2} \cdot dU \\ &\quad + \frac{1}{2} \int_U U^{\otimes 2} |\boxed{1}\boxed{i}\rangle \left\langle \begin{smallmatrix} \boxed{1} \\ \boxed{i} \end{smallmatrix} \right| U^{\dagger, \otimes 2} \cdot dU + \frac{1}{2} \int_U U^{\otimes 2} \left| \begin{smallmatrix} \boxed{1} \\ \boxed{i} \end{smallmatrix} \right\rangle \langle \boxed{1}\boxed{i}| U^{\dagger, \otimes 2} \cdot dU. \end{aligned}$$

This we can analyze using our three integration formulas Equations (7) to (9), which tell us that for $i > 1$,

$$\int_U U^{\otimes 2} |1, i\rangle \langle 1, i| U^{\dagger, \otimes 2} \cdot dU = \frac{1}{2} \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} + \frac{1}{2} \cdot \frac{\Pi_{\square}}{\dim(V_{\square})}.$$

Computing M_{avg} . Using our integration formulas, we can compute M_{avg} as

$$\begin{aligned}
M_{\text{avg}} &= d \cdot \sum_{i=1}^d \lambda_i^\dagger \cdot \int_U U^{\otimes 2} |1, i\rangle\langle 1, i| U^{\dagger, \otimes 2} \cdot dU \\
&= d \cdot \left(d \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} + (d-1) \cdot (-1) \cdot \left(\frac{1}{2} \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} + \frac{1}{2} \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} \right) \right) \\
&= \frac{d(d+1)}{2} \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} - \frac{d(d-1)}{2} \cdot \frac{\Pi_{\square\square}}{\dim(V_{\square\square})} \\
&= \Pi_{\square\square} - \Pi_{\square\square} \\
&= \text{SWAP}.
\end{aligned}$$

Now that we have computed M_{avg} , we can conclude our calculation of $\hat{\rho}$'s first moment as follows:

$$\mathbf{E}[\hat{\rho}] = \text{tr}_1(M_{\text{avg}} \cdot \rho \otimes I) = \text{tr}_1(\text{SWAP} \cdot \rho \otimes I) = \rho,$$

where we have here used the fact that $\text{tr}_1(\text{SWAP} \cdot A \otimes I) = A$ for any matrix $A \in \mathbb{C}^{d \times d}$. This shows that $\hat{\rho}$ is an unbiased estimator for ρ and therefore completes the proof.

1.7.2 The two copy case

Exploiting permutation symmetry. Next, let us illustrate how these ideas generalize to the case where our input is two copies $\rho \otimes \rho$. Here, there is a new feature not present in the single copy case, in that the input has “permutation symmetry”, meaning that it is invariant under swapping the two registers. Mathematically, this can be expressed through the equation

$$\text{SWAP} \cdot \rho \otimes \rho \cdot \text{SWAP} = \rho \otimes \rho.$$

One consequence of this is that $\rho \otimes \rho$ commutes with SWAP, i.e. that $\rho \otimes \rho \cdot \text{SWAP} = \text{SWAP} \cdot \rho \otimes \rho$, and so $\rho \otimes \rho$ and SWAP share a common eigenbasis. As we have seen, SWAP has two eigenspaces, a $+1$ eigenspace corresponding to the subspace $V_{\square\square}$ and a -1 eigenspace corresponding to $V_{\square\square}$, and so $\rho \otimes \rho$'s eigenvectors also split into those that live in $V_{\square\square}$ and those that live in $V_{\square\square}$. This means that $\rho \otimes \rho$ can be written as $\rho \otimes \rho = \rho_{\square\square} + \rho_{\square\square}$, where $\rho_{\square\square}$ is a subnormalized mixed state entirely supported on $V_{\square\square}$ and $\rho_{\square\square}$ is a subnormalized mixed state entirely supported on $V_{\square\square}$.

Inside the symmetric subspace. In the case of two copies of ρ , the debiased Keyl's algorithm begins by measuring $\rho \otimes \rho$ with the projective measurement $\{\Pi_{\square\square}, \Pi_{\square\square}\}$. Let us first consider the case when the outcome λ is equal to $\lambda = \square\square$. This case occurs with probability $\text{tr}(\Pi_{\square\square} \cdot \rho^{\otimes 2}) = \text{tr}(\rho_{\square\square})$, and when it occurs the state collapses to $\rho_{\square\square} / \text{tr}(\rho_{\square\square})$. Recalling that $\lambda^\dagger = (d+2, -1, \dots, -1)$, the debiased Keyl's algorithm will output an estimator $\hat{\rho}$ with eigenvalues $\lambda^\dagger/2$. But what are its eigenvectors? To learn these, we must perform a further measurement on the collapsed state $\rho_{\square\square} / \text{tr}(\rho_{\square\square})$, which lies within $V_{\square\square}$. The debiased Keyl's algorithm uses the POVM

$$\{\dim(V_{\square\square}) \cdot U^{\otimes 2} | \square\square \rangle\langle \square\square | U^{\dagger, \otimes 2} \cdot dU\}.$$

This is indeed a measurement on the space $V_{\square\square}$, as

$$\dim(V_{\square\square}) \cdot \int_U U^{\otimes 2} | \square\square \rangle\langle \square\square | U^{\dagger, \otimes 2} \cdot dU = \Pi_{\square\square}$$

due to Equation (7). Letting $\mathbf{U}$ be the result of this measurement, the debiased Keyl's algorithm outputs the estimator

$$\hat{\rho} = \sum_{i=1}^d \lambda_i^\dagger \cdot |u_i\rangle\langle u_i| = \frac{(d+2)}{2} \cdot |u_1\rangle\langle u_1| - \frac{1}{2} \cdot |u_2\rangle\langle u_2| - \dots - |u_d\rangle\langle u_d|,$$

where $|\mathbf{u}_i\rangle = \mathbf{U} \cdot |i\rangle$ for all $1 \leq i \leq d$. As in the single copy case, $|\mathbf{u}_1\rangle$ should be a rough approximation of the largest eigenvector of ρ , or should at least have decent overlap with ρ 's largest eigenvectors, because the probability density function of $\mathbf{U}$ is given by

$$\begin{aligned} \mathbf{Pr}[\mathbf{U} = U \mid \rho, \square\square] &= \text{tr} \left(\left(\dim(V_{\square\square}) \cdot U^{\otimes 2} |\square\square\rangle\langle\square\square| U^{\dagger, \otimes 2} \cdot dU \right) \cdot \frac{\rho_{\square\square}}{\text{tr}(\rho_{\square\square})} \right) \\ &= \frac{\dim(V_{\square\square})}{\text{tr}(\rho_{\square\square})} \cdot \langle \square\square | U^{\dagger, \otimes 2} \rho_{\square\square} U^{\otimes 2} | \square\square \rangle \cdot dU \\ &= \frac{\dim(V_{\square\square})}{\text{tr}(\rho_{\square\square})} \cdot (\langle 1 | U^\dagger \rho U | 1 \rangle)^2 \cdot dU. \end{aligned}$$

In fact, this PDF scales not just with $\langle 1 | U^\dagger \rho U | 1 \rangle$ but with its *square*, meaning that $|\mathbf{u}_1\rangle$ should be even better concentrated around ρ 's highest eigenvectors than it was in the single copy case.

Highest weight vectors. If, instead, the measurement outcome is $\lambda = \square$, then the state collapses to $\rho_{\square}/\text{tr}(\rho_{\square})$. The debiased Keyl's algorithm then performs within the $V_{\square}$ space the POVM

$$\left\{ \dim(V_{\square}) \cdot U^{\otimes 2} \left| \begin{smallmatrix} 1 \\ 2 \end{smallmatrix} \right\rangle\left\langle \begin{smallmatrix} 1 \\ 2 \end{smallmatrix} \right| U^{\dagger, \otimes 2} \cdot dU \right\},$$

which is indeed a measurement due to [Equation \(8\)](#). As in the single copy case and the $\lambda = \square\square$ case, this measurement is biased towards U 's for which $U \cdot |1\rangle$ and $U \cdot |2\rangle$ have good overlap with ρ 's largest eigenvectors. In both the one- and two-copy cases, our measurements are chosen fixing a particular vector,

$$\text{either } |1\rangle, \left| \begin{smallmatrix} 1 \\ 1 \end{smallmatrix} \right\rangle, \text{ or } \left| \begin{smallmatrix} 1 \\ 2 \end{smallmatrix} \right\rangle,$$

and considering all possible rotations of this vector by a unitary matrix U . (Note that one can unify the notation by observing that in the $n = 1$ case, the input space $\mathbb{C}^d$ is also known as the irreducible representation $V_{\square}$, and within this representation $|1\rangle$ is written as the vector $\left| \begin{smallmatrix} 1 \\ 1 \end{smallmatrix} \right\rangle$.) These vectors are chosen to have the “most 1’s” possible within their subspace, then the “most 2’s” possible, and so forth, in order to bias the U 's towards ρ 's largest eigenvectors, as described above. Vectors of this form are known in representation theory as *highest weight vectors*; as a result, this measurement, which was first introduced by Keyl [\[Key06\]](#) in his tomography algorithm, is sometimes referred to as the “rotated highest weight measurement”.

Computing M_{avg} . Computing M_{avg} in the two copy case splits into two expressions, one for the symmetric subspace and another for the antisymmetric subspace. Writing $\lambda = (2, 0, \dots, 0)$, the expression for the symmetric subspace is

$$\begin{aligned} &\int_U (\dim(V_{\square\square}) \cdot U^{\otimes 2} |\square\square\rangle\langle\square\square| U^{\dagger, \otimes 2} \cdot dU) \otimes \sum_{i=1}^d \lambda_i^\uparrow \cdot U |i\rangle\langle i| U^\dagger \\ &= \dim(V_{\square\square}) \cdot \sum_{i=1}^d \lambda_i^\uparrow \cdot \int_U U^{\otimes 3} |\square\square\rangle\langle\square\square| \otimes |i\rangle\langle i| U^{\dagger, \otimes 3} \cdot dU. \end{aligned}$$

Computing this involves computing the integrals

$$\int_U U^{\otimes 3} |\square\square\rangle\langle\square\square| \otimes |i\rangle\langle i| U^{\dagger, \otimes 3} \cdot dU.$$

As in the $n = 1$ case, this not easy to directly compute because $|\square\square\rangle \otimes |i\rangle$ is not contained inside an irreducible representation of the unitary group over $(\mathbb{C}^d)^{\otimes 3}$. Instead, we must first split this vector into a component for each irreducible representation. As it turns out, the irreducible representations of the unitary group on $(\mathbb{C}^d)^{\otimes 3}$ are labeled by Young diagrams λ with 3 boxes, and in particular

$$|\square\square\rangle \otimes |1\rangle = |\square\square\square\rangle, \quad \text{and} \quad |\square\square\rangle \otimes |i\rangle = \sqrt{\frac{1}{3}} \cdot |\square\square\square_i\rangle + \sqrt{\frac{2}{3}} \cdot \left| \begin{smallmatrix} \square & \square \\ i \end{smallmatrix} \right\rangle, \quad \text{for } i > 1, \quad (11)$$

where the vectors in the decomposition are elements of the irreducible representations corresponding to $\square\square\square$ and $\square\square$, respectively. This process, in which we take a vector in an irreducible representation of the unitary group corresponding, tensor on a $|i\rangle$ vector, and then decompose that into vectors which live in separate irreducible representations, is known as a *Clebsch-Gordan transform*, and the coefficients that arise in this process (for example, the coefficients in Equation (11), as well as the coefficients in Equation (10)) are known as *Clebsch-Gordan coefficients*. Computing this integral, and therefore computing M_{avg} , requires being able to understand and manipulate these Clebsch-Gordan coefficients and evaluate large expressions involving them. Eventually, we aim to show that $M_{\text{avg}} = \frac{1}{2} \cdot \text{SWAP}_{1,3} + \frac{1}{2} \cdot \text{SWAP}_{2,3}$. If this is true, then

$$\begin{aligned} \mathbf{E}[\hat{\rho}] &= \text{tr}_{[2]}(M_{\text{avg}} \cdot \rho \otimes \rho \otimes I) = \text{tr}_{[2]} \left(\left(\frac{1}{2} \cdot \text{SWAP}_{1,3} + \frac{1}{2} \cdot \text{SWAP}_{2,3} \right) \cdot \rho \otimes \rho \otimes I \right) \\ &= \frac{1}{2} \cdot \text{tr}_{[2]}(\text{SWAP}_{1,3} \cdot \rho \otimes \rho \otimes I) + \frac{1}{2} \cdot \text{tr}_{[2]}(\text{SWAP}_{2,3} \cdot \rho \otimes \rho \otimes I) \\ &= \frac{1}{2} \cdot \text{tr}_{[1]}(\text{SWAP} \cdot \rho \otimes I) + \frac{1}{2} \cdot \text{tr}_{[1]}(\text{SWAP} \cdot \rho \otimes I) \\ &= \rho, \end{aligned}$$

proving that $\hat{\rho}$ is an unbiased estimator for ρ .

1.7.3 The general copy case

The case of general n closely resembles the case of $n = 2$ copies. First, our goal is to show that

$$M_{\text{avg}} = \frac{1}{n} \cdot (\text{SWAP}_{1,n+1} + \text{SWAP}_{2,n+1} + \cdots + \text{SWAP}_{n,n+1}), \quad (12)$$

which implies that the debiased Keyl's algorithm is indeed an unbiased estimator. (As an aside, we note that the right-hand side of this expression is a well-known element of the group algebra $\mathbb{C}[S_{n+1}]$ called the *Jucys-Murphy element*.) Computing M_{avg} then breaks into a separate sub-expression for each irreducible representation of the unitary group corresponding to a Young diagram λ with n boxes. For each such λ , computing this sub-expression involves taking the highest weight vector T^λ of λ and applying the Clebsch-Gordan transform to vectors of the form $|T^\lambda\rangle \otimes |i\rangle$ in order to understand how these vectors decompose into irreducible representations corresponding to Young diagrams with $n + 1$ boxes. Having done this, proving Equation (12) then involves deriving exact formulas for certain complicated expressions involving Clebsch-Gordan coefficients, which is made all the more difficult since for general n even individual Clebsch-Gordan coefficients can be quite complicated themselves. We also then have to compute the *second* moment of our estimator in order to prove Theorem 1.4, and this involves computing a formula for the matrix

$$M_{\text{avg}}^{(2)} = \sum_{\hat{\rho}} M_{\hat{\rho}} \otimes \hat{\rho} \otimes \hat{\rho}.$$

This is the most technically challenging part of our work, as it involves understanding the effect that two back-to-back Clebsch-Gordan transforms have on a highest weight vector $|T^\lambda\rangle$.

In recent years, Clebsch-Gordan coefficients and the Clebsch-Gordan transform have become increasingly central to the study of quantum computing and information. Perhaps their first appearance in the quantum computing literature dates back to the efficient algorithm for the quantum Schur transform by Bacon, Chuang, and Harrow from 2005 [BCH05], but more recently they have been used to compute the dual and mixed Schur transforms [Ngu24, GBO24], perform optimal qudit purification [LFIC24], efficiently implement port-based teleportation [GBO24], take a quantum majority vote [BLM+23], and prove lower bounds on the number of queries needed to invert a unitary [CYZ25]. To our knowledge, our work is the first to apply the Clebsch-Gordan transform in the domain of quantum learning theory, though we expect that they will see further use in this domain in the future. Our primary technical contribution, then, is to develop a host of new tools for evaluating natural expressions that arise when studying Clebsch-Gordan coefficients, and we hope that our tools will find further applications in future works which make use of the Clebsch-Gordan transform.

1.8 Open problems

Our work leaves open a number of interesting follow-up questions, which we list below.

1. **(Learning with very high probability)** In this work, we have studied algorithms for quantum tomography which succeed with a high constant probability, say 99%. But what if we want to succeed with probability $1 - \delta$, where $\delta > 0$ is a parameter which is allowed to vary? As we show in [Section 2.2](#) below, an immediate corollary of our results is that $n = O(d^2/\varepsilon^2 \cdot \log(1/\delta))$ copies suffice to estimate a mixed state $\rho \in \mathbb{C}^{d \times d}$ to error ε with success probability $1 - \delta$ in both trace and Bures distances. (In fact, in the case of trace distance, this also follows from the work of O’Donnell and Wright [[OW16](#)].) However, it is natural to conjecture that a stronger bound of

$$n = O\left(\frac{d^2}{\varepsilon^2} + \frac{\log(1/\delta)}{\varepsilon^2}\right)$$

copies should hold for both distance measures. A similar bound is known to hold in the classical setting of learning discrete distributions: in particular, $n = O(d/\varepsilon^2 + \log(1/\delta)/\varepsilon^2)$ samples are sufficient to learn a distribution $p = (p_1, \dots, p_d)$ up to error ε with probability $1 - \delta$, in both trace and Hellinger distances [[Can20](#)]. We note that a bound of this form, with some additional logarithmic factors, can be derived from [[HHJ⁺16](#), Equation 14], and so the challenge here is to achieve this bound precisely, with no additional logarithms. One natural route for accomplishing this is to study higher moments of our estimator, generalizing [Theorem 1.4](#) beyond the second moment.

Let us note that this question also appears to be open in the case of tomography algorithms which use independent measurements, even for learning in trace distance. It follows from [[GKKT20](#), Theorem 2] that $n = O(d^3/\varepsilon^2 + d^2 \log(1/\delta)/\varepsilon^2)$ copies are sufficient to learn in trace distance with independent measurements (in fact, the uniform POVM tomography algorithm achieves this bound); on the other hand, it appears to follow from the argument of [[CHL⁺23](#)] that $n = \Omega(d^3/\varepsilon^2 + d \log(1/\delta)/\varepsilon^2)$ copies are necessary in this setting [[CL24](#)]. We believe that closing this gap, and determining the bound for Bures distance, is also an interesting open problem.

2. **(Spectrum estimation)** How many copies do we need to learn a mixed state ρ ’s eigenvalues $\alpha = (\alpha_1, \dots, \alpha_d)$? This problem is certainly no harder than full state tomography, which requires learning ρ ’s eigenvalues *and* eigenvectors, and so $n = O(d^2/\varepsilon^2)$ samples, which suffice for tomography, also suffice for spectrum estimation. But can we do better? Or does one first have to learn the eigenvectors in order to then learn the eigenvalues? Currently, the best known spectrum estimation algorithm is the EYD algorithm, and although it is known to only require $n = O(d^2/\varepsilon^2)$ copies [[OW16](#)], it is also known to fail at spectrum estimation when $n = o(d^2/\varepsilon^2)$ [[OW15b](#)]. On the flip side, in the classical setting, it is known that $n = O(d/\log(d))$ samples are necessary and sufficient to learn the set of probability values $\{p_1, \dots, p_d\}$ of a probability distribution $p = (p_1, \dots, p_d)$ [[VV11](#), [VV17](#), [HJW18](#)] (which is the classical analogue of a mixed state’s spectrum), which does outperform the $n = \Omega(d)$ samples needed to learn p itself.

Whether we can do better than full state tomography was recently resolved in the setting of algorithms which use unentangled measurements by Pelecanos, Tan, Tang, and Wright [[PTTW25](#)], who showed that $n = o(d^3)$ samples suffice for spectrum estimation, which beats the $n = \Omega(d^3)$ copies which are required for full state tomography with unentangled measurements [[CHL⁺23](#)]. Their argument crucially relied on the fact that we have guarantees for full state tomography with unentangled measurements in the “very high probability” regime. As a result, we believe that a resolution of [Item 1](#) above is likely to help design improved algorithms for spectrum estimation with entangled measurements.

3. **(Other distance measures)** For the first time, our work shows optimal bounds for learning a mixed state in Bures distance, which is a more stringent distance metric than trace distance. There are even more stringent divergences than Bures which are commonly studied in the quantum information literature, namely the quantum relative entropy and the Bures χ^2 -divergence. How many copies are needed to learn a state ρ for either of these two divergences? This question was studied by Flammia and O’Donnell [[FO24](#)], who were motivated by the application of testing whether a bipartite quantum state has zero quantum mutual information. In fact, combining our [Theorem 1.6](#) with their [Theorem 2.34](#)

produces a tomography algorithm on rank r states $\rho \in \mathbb{C}^{d \times d}$ which learns in quantum relative distance error ε using only $n = O(rd/\varepsilon) \cdot \log(d/\varepsilon)$ copies, improving on the bound in their Corollary 1.8 by a factor of $\log(d/\varepsilon)$ (see their Page 12 for a computation of their precise bound). But the best lower bound we know for this problem is $n = \Omega(rd/\varepsilon)$ (which follows from the Bures distance lower bound of [Yue23]), suggesting that $n = O(rd/\varepsilon)$ might be the optimal upper bound.

4. **(Efficient algorithms for Keyl’s measurement)** Although the algorithms we consider in this work are sample-efficient, they are not yet known to be computationally efficient. The one exception is when $n = 1$, as it is known that the uniform POVM tomography algorithm still works if one substitutes the Haar random unitary used to specify the measurement basis with a random unitary chosen from a unitary 2-design or 3-design, depending on one’s application [GKKT20, HKP20]. This yields an efficient algorithm, as there are examples of 2-designs and 3-designs, such as Clifford unitaries, which are computationally efficient to sample from. Can Keyl’s measurement be made efficient for general n ?
5. **(Optimality of the debiased Keyl’s algorithm)** Is the debiased Keyl’s algorithm *the* optimal entangled tomography algorithm? This is of course not a well-defined question, but perhaps one way of making it formal is asking whether the debiased Keyl’s algorithm is the minimum variance estimator among all unbiased estimators, where the variance is defined as $\mathbf{E} \|\hat{\rho} - \rho\|_2^2$. We show in [Theorem 3.9](#) below that the answer is yes, at least for a large family of estimators based on Keyl’s measurement, but this doesn’t rule out an even better unbiased estimator which uses a different measurement. A good starting point for proving this, if it is indeed true, is the $n = 1$ case and the general n pure state case.
6. **(Simpler proofs of our bounds)** Can our proofs be made simpler? As we mention in [Section 1.7.1](#), there is a substantially simpler proof than the one we use that the debiased Keyl’s algorithm is unbiased in the $n = 1$ case; could this proof be generalized to larger n ? More broadly, is there a more conceptual reason why our algorithm gives an unbiased estimator, beyond just what we learn from working out the Clebsch-Gordan coefficients? Possibly even one that doesn’t involve representation theory? There have been several instances in quantum learning in which results which were first derived using heavy representation theory were then given new proofs which are substantially cleaner and more conceptual (e.g. the new proof of the mixedness testing upper bound from O’Donnell and Wright [OW15b] due to Badescu, O’Donnell, and Wright [BOW19], and the new proof of the mixedness testing *lower* bound from [OW15b] due to O’Donnell and Wadhwa [OW25]). It would be interesting if the same could be done here.

1.9 Paper organization

This paper is organized into five parts.

- [Part I](#) contains this introduction, as well as preliminaries and a full description of the debiased Keyl’s algorithm. The preliminaries contain general background material on quantum learning theory, as well as an overview of all the representation theory needed to state the debiased Keyl’s algorithm.
- [Part II](#) contains all five of our applications. This part uses [Theorems 1.3](#) and [1.4](#), our main results concerning the first and second moments of the debiased Keyl’s algorithm, as a black box. With the exception of the tomography with limited entanglement application, which relies on some lemmas from the trace distance tomography section, the applications are meant to be self-contained and readable in any order.
- [Part III](#) contains an introduction to the Clebsch-Gordan transform and the Clebsch-Gordan coefficients.
- [Part IV](#) uses the Clebsch-Gordan technology developed in the previous part to prove [Theorems 1.3](#) and [1.4](#), our first and second moment bounds.
- [Part V](#) is an appendix which contains a technical result that we make use of in our first and second moment proofs. Informally, it shows that the algorithm for the Schur transform due to Bacon, Chuang, and Harrow [BCH05] also happens to compute a specific choice of basis for the irreducible representations of the symmetric group known as *Young’s orthogonal basis*. Previously, it was known only that it

computes a *Young-Yamanouchi basis*, a more general class of bases of which Young’s orthogonal basis is a specific and especially convenient example.

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.

2 Preliminaries

We use **boldface** to denote random variables, and define $[d] = \{1, \dots, d\}$. Additionally, we will use δ_{ij} for the Kronecker delta, that is, $\delta_{ij} := \mathbb{1}[i = j]$.

2.1 Quantum distances

We first review the quantum distances (and divergences) that we will use throughout this paper. The interested reader is encouraged to refer to [BOW19] for a systematic treatment of these quantum distances and their relationships.

Definition 2.1 (Schatten p -norm). Let $M \in \mathbb{C}^{d \times d}$ be a matrix with singular values $\lambda_1, \dots, \lambda_d$. The *Schatten p -norm*, for $p \geq 1$, is defined as

$$\|M\|_p = \left(\sum_{i=1}^d |\lambda_i|^p \right)^{1/p}.$$

Let $\rho, \sigma \in \mathbb{C}^{d \times d}$ be quantum states.

Definition 2.2 (Trace distance). The *trace distance* between ρ and σ is

$$D_{\text{tr}}(\rho, \sigma) = \frac{1}{2} \|\rho - \sigma\|_1.$$

Definition 2.3 (Fidelity). The *fidelity* of ρ and σ is

$$F(\rho, \sigma) = \|\sqrt{\rho}\sqrt{\sigma}\|_1 = \text{tr} \sqrt{\sqrt{\rho}\sigma\sqrt{\rho}}.$$

Our version for the fidelity of quantum states is sometimes referred to as the “square root fidelity”. Since fidelity is not a metric, we use the closely related Bures distance metric.

Definition 2.4 (Bures distance). The *Bures distance* between ρ and σ is

$$D_B(\rho, \sigma) = \sqrt{2(1 - F(\rho, \sigma))}.$$

Definition 2.5 (Bures χ^2 -divergence). The *Bures χ^2 -divergence* of ρ from σ is given by the following formula when σ is a diagonal full-rank state with eigenvalues $(\alpha_1, \dots, \alpha_d)$:

$$D_{\chi^2}(\rho \parallel \sigma) = \sum_{i,j=1}^d \frac{2}{\alpha_i + \alpha_j} \cdot |\rho_{ij} - \sigma_{ij}|^2.$$

In this paper, we will compute the trace distance of matrices that are not necessarily PSD, or even Hermitian, as well as the the fidelity, Bures distance, and Bures χ^2 -divergence of sub-normalized PSD matrices. The definitions given above can be extended to these cases using the same formulas. We now state some useful results about the quantum distance measures above. The first set of inequalities relates trace distance and fidelity [NC10, Section 9.2].

Lemma 2.6 (Fuchs-van de Graaf inequalities). $1 - F(\rho, \sigma) \leq D_{\text{tr}}(\rho, \sigma) \leq \sqrt{1 - F(\rho, \sigma)^2}$.

The Fuchs-van de Graaf inequalities can be used to show a pair of inequalities between trace distance and Bures distance.

Corollary 2.7. $\frac{1}{2}D_B(\rho, \sigma)^2 \leq D_{\text{tr}}(\rho, \sigma) \leq D_B(\rho, \sigma)$.

Bures distance is also related to the Bures χ^2 -divergence, via the following inequality.

Lemma 2.8 (Proposition 2.31, [FO24]). $D_B(\rho, \sigma) \leq \sqrt{D_{\chi^2}(\rho \parallel \sigma)}$.

Finally, the Gentle Measurement Lemma gives an expression for the fidelity between pre- and post-measurement versions of a state.

Lemma 2.9 (Gentle Measurement Lemma [Win02]). *Given a quantum state ρ and a projector Π , let $\rho|_{\Pi}$ be the post-measurement state if ρ is measured with the projective measurement $\{\Pi, I - \Pi\}$ and Π is obtained. Then*

$$\left(1 - \frac{1}{2}D_B(\rho, \rho|_{\Pi})^2\right)^2 = F(\rho, \rho|_{\Pi})^2 = \text{tr}(\Pi \cdot \rho).$$

2.2 Learning with high probability and amplification

Our goal throughout this paper is to obtain tomography algorithms that output an estimator which is ε -close to the true state, in some distance D , with low probability of failure. Such an algorithm is said to succeed “with high probability”. The exact threshold for success is not particularly important, but for concreteness, we will take it to mean with probability at least 0.99.

Sometimes our algorithms involve multiple (though, a small number of) steps, which each, individually, might only succeed with high probability themselves (possibly conditioned on the success of prior steps). We can use a union bound to lower bound the success probability of the overall algorithm. For example, an algorithm consisting of three steps, which each succeed with high probability, succeeds with probability at least 0.97. However, we have now dropped under our threshold. Fortunately, there is a standard procedure we may apply to amplify the success probability over 0.99, with only a constant factor increase in the number of samples, and a constant factor increase in error.

Let D be a metric (e.g. trace distance, or Bures distance). Consider a tomography algorithm $\mathcal{A}$ that takes as input n samples of a state ρ , and outputs an estimator $\hat{\rho}$ such that $D(\hat{\rho}, \rho) \leq \varepsilon$ with probability at least, say, 0.51. We define a new algorithm $\mathcal{A}'$ that uses $n' = O(n \cdot \log(1/\delta))$ samples, and outputs a state $\hat{\rho}'$ such that $D(\hat{\rho}', \rho) \leq 3\varepsilon$ with probability at least $1 - \delta$.

The new algorithm divides the samples into $k = O(\log(1/\delta))$ groups of n samples, and executes $\mathcal{A}$ on each group of samples to obtain estimates $\hat{\rho}_1, \dots, \hat{\rho}_k$. Then $\mathcal{A}'$ will output an estimate $\hat{\rho}_i$ maximizing the following quantity:

$$N_i := |\{j \in [k] \mid D(\hat{\rho}_j, \hat{\rho}_i) \leq 2\varepsilon\}|.$$

See, e.g. [HKOT23, Proposition 2.4] for details.

For our purposes, we will set $\delta = 0.01$, and thus $\mathcal{A}'$ succeeds with high probability with $n' = O(n)$ samples, with a constant factor degradation in error. If $\mathcal{A}$ scales, say, inverse-polynomially in ε , then $\mathcal{A}'$ can be used to output estimates to the same accuracy, with high probability, with the same overall sample complexity.

2.3 On low-rank tomography algorithms

Our algorithms will often take copies of a rank- r quantum state ρ , and output an estimator $\hat{\rho}$ that is close to ρ in some distance. However, it is not necessarily the case that $\hat{\rho}$ is, itself, a rank- r state. However, when we are dealing with tomography in some metric (e.g. the trace distance, or Bures distance) there exists a generic transformation to our algorithm that makes it output an estimator $\hat{\rho}'$ of rank r .

Remark 2.10. Let D be a metric. Also let $\mathcal{A}$ be a tomography algorithm that takes n samples of a rank- r state ρ and outputs an estimator $\hat{\rho}$, not necessarily of rank r , such that $D(\hat{\rho}, \rho) \leq \varepsilon$ with high probability. Then there exists a tomography algorithm $\mathcal{A}'$ that takes n samples of a rank- r state ρ and outputs a rank- r estimator $\hat{\rho}'$ such that $D(\hat{\rho}', \rho) \leq 2\varepsilon$ with high probability.

The new algorithm $\mathcal{A}'$ will set $\hat{\rho}'$ to be the rank- r quantum state that is closest to $\hat{\rho}$ in the metric D . Since ρ is also rank- r , this means that $\hat{\rho}'$ is at least as close to $\hat{\rho}$:

$$D(\hat{\rho}, \hat{\rho}') \leq D(\hat{\rho}, \rho).$$

Thus $D(\hat{\rho}, \hat{\rho}') \leq \varepsilon$ with high probability. From the triangle inequality, we conclude that with high probability,

$$D(\hat{\rho}', \rho) \leq D(\hat{\rho}', \hat{\rho}) + D(\hat{\rho}, \rho) \leq 2\varepsilon.$$

2.4 Representation theory

In this section, we review the representation theory needed to describe the debiased Keyl's algorithm. For a more detailed treatment of applying representation theory to quantum learning algorithms, see [Wri16, Chapter 2].

2.4.1 Basics

The *symmetric group* S_n is the group of permutations of the set $[n]$. The *unitary group* $U(d)$ is the group of $d \times d$ unitary matrices. More generally, for V a complex, finite-dimensional Hilbert space, $U(V)$ is the group of unitary operators on V .

Definition 2.11 (Representations). Let G be a group. A *complex, unitary, finite-dimensional representation* (henceforth, a *representation*) of G is a pair (μ, V) , where V is a finite-dimensional, complex Hilbert space, and $\mu : G \rightarrow U(V)$ is a group homomorphism. That is, $\mu(e) = I_V$, $\mu(g)$ is unitary for all $g \in G$, and $\mu(g \cdot h) = \mu(g) \cdot \mu(h)$ for all $g, h \in G$. The *dimension* of the representation, denoted $\dim(\mu)$, is the dimension of V .

Commonly, a representation of G , (μ, V) , is referred to either by μ or V , when the meaning is clear from context.

Definition 2.12 (Characters). Let G be a group, and let μ be a representation. The *character* of μ is the function $\chi_\mu : G \rightarrow \mathbb{C}$ given by $\chi_\mu(g) = \text{tr}(\mu(g))$.

Definition 2.13 (Intertwining operators). Let G be a group, and let (μ_1, V_1) and (μ_2, V_2) be representations of G . Then a map $T : V_1 \rightarrow V_2$ is an *intertwining operator*, or *intertwiner*, if $T \cdot \mu_1(g) = \mu_2(g) \cdot T$ for all $g \in G$. We say T *intertwines* μ_1 and μ_2 .

Definition 2.14 (Isomorphic representations). Let G be a group, and let (μ_1, V_1) and (μ_2, V_2) be representations of G . We say μ_1 and μ_2 are *isomorphic* representations if there exists an invertible map T that intertwines μ_1 and μ_2 . We write $\mu_1 \cong \mu_2$, or, if we want to emphasize the group, $\mu_1 \stackrel{G}{\cong} \mu_2$.

Definition 2.15 (Irreducible representations). Let G be a group, and let (μ, V) be a representation of G . A subspace $W \subseteq V$ is *invariant* if $\mu(g) \cdot W \subseteq W$ for all $g \in G$. An invariant subspace is *trivial* if $W = \{0\}$ or $W = V$. The representation μ is *irreducible* if V has no nontrivial invariant subspaces. An irreducible representation is also called an *irrep* for short. The set of all isomorphism classes of irreducible representations of G is denoted $\hat{G}$.

We will usually fix an irrep μ_i from each isomorphism class, and identify $\hat{G}$ with the set $\{\mu_i\}$.

Lemma 2.16 (Schur's lemma). Let G be a group, and let (μ_1, V_1) and (μ_2, V_2) be two irreducible representations of G . Let $T : V_1 \rightarrow V_2$ intertwine μ_1 and μ_2 . If μ_1 and μ_2 are non-isomorphic, then $T = 0$. If instead $\mu_1 = \mu_2$, then

$$T = \frac{\text{tr}(T)}{\dim(\mu_1)} \cdot I_{V_1}.$$

The following result, which we will make frequent use of, can be shown using Schur's lemma.

Lemma 2.17. Let G be a group, and let (μ_1, V_1) and (μ_2, V_2) be unitary, isomorphic representations of G . Then there exists a unitary $U : V_1 \rightarrow V_2$ intertwining μ_1 and μ_2 . That is, for all $g \in G$,

$$U \cdot \mu_1(g) \cdot U^\dagger = \mu_2(g).$$

Definition 2.18 (Direct sum of representations). Let G be a group, and let $(\mu_1, V_1), \dots, (\mu_k, V_k)$ be representations of G . Then the *direct sum* of $\mu_1, \dots, \mu_k$ is the representation (μ, V) , where $V := V_1 \oplus \dots \oplus V_k$ and $\mu(g) := \mu_1(g) \oplus \dots \oplus \mu_k(g)$, for all $g \in G$. Equivalently, we can write $\mu(g) = \sum_{i=1}^k |i\rangle\langle i| \otimes \mu_i(g)$.

Generally, each representation μ_i may appear multiple times. If μ_i appears m_i times, we also write $\mu = \bigoplus_{i=1}^k m_i \cdot \mu_i$, and refer to m_i as the *multiplicity* of μ_i .

Definition 2.19 (Complete reducibility). Let G be a group, and let (μ, V) be a representation. Then μ is *completely reducible* if

$$\mu \cong \bigoplus_{i=1}^k m_i \cdot \mu_i,$$

with each μ_i irreducible.

Finite-dimensional unitary representations are completely reducible.

Definition 2.20 (Restriction of representations). Let H be a subgroup of G , and let μ be a representation of G . The *restriction of μ to H* , a representation of H , is denoted $\mu \downarrow_H^G$, and is given by

$$\mu \downarrow_H^G(h) = \mu(h).$$

The restriction of an irreducible representation may itself be reducible. *Branching rules* describe how restrictions decompose into irreps.

Definition 2.21 (Branching rules). Let H be a subgroup of G , and let μ be a representation of G . A *branching rule* is a statement of the form:

$$\mu \downarrow_H^G \cong \bigoplus_{\mu_i \in \hat{H}} m_i \cdot \mu_i.$$

If $m_i \in \{0, 1\}$ for all i , then we say the branching rule is *multiplicity-free*.

2.4.2 Partitions and Young diagrams

Definition 2.22 (Partitions). Let n be a nonnegative integer. A *partition* of n is a finite list of nonnegative integers $\lambda = (\lambda_1, \dots, \lambda_m)$ such that $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_m \geq 0$ and $\lambda_1 + \dots + \lambda_m = n$. We write $\lambda \vdash n$. The *size* of λ is $|\lambda| = n$, and the *length* of λ , denoted $\ell(\lambda)$, is the largest index i with $\lambda_i > 0$.

Remark 2.23. Partitions which only differ by some number of trailing zeroes are considered equivalent. As we will see later, when considering the problem of learning a quantum state $\rho \in \mathbb{C}^d$, only partitions of length at most d will be relevant to us. In this case, we will often view a partition as a tuple of length exactly d .

Partitions have diagrammatic representations known as Young diagrams.

Definition 2.24 (Young diagrams). Let $\lambda \vdash n$. The *Young diagram of shape λ* is a diagram consisting of n boxes, organized into $\ell(\lambda)$ left-justified rows, such that the i -th row contains λ_i boxes. We will also refer to boxes as *cells*.

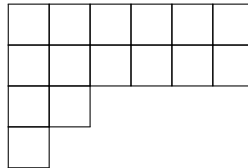


Figure 3: The Young diagram above corresponds to the partition $\lambda = (6, 6, 2, 1)$, with $|\lambda| = 15$ and $\ell(\lambda) = 4$.

See [Figure 3](#) for an example of a Young diagram. We use the following non-standard definition to refer to a consecutive group of rows with the same length, something that will be useful for describing our family of unbiased estimators.

Definition 2.25 (Blocks in Young diagrams). Let $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ be a Young diagram. Then $\{i, i+1, \dots, j\}$ is a *block* of λ if

$$\lambda_i = \lambda_{i+1} = \dots = \lambda_j,$$

and $\lambda_{i-1} > \lambda_i$ when $i > 1$, and $\lambda_j > \lambda_{j+1}$ when $j < d$.

In the Young diagram in [Figure 3](#), $\{1, 2\}$, $\{3\}$ and $\{4\}$ are blocks. If we regard the corresponding partition as a tuple of length d , and $d \geq 5$, then $\{5, \dots, d\}$ is also a block.

Definition 2.26 (Standard Young tableaux). Let $\lambda \vdash n$. A *standard Young tableau* (SYT) of shape λ is a filling of λ with the numbers $\{1, \dots, n\}$, such that the entries strictly increase both from left to right, and from top to bottom. The set of all SYTs of shape λ is denoted $\text{SYT}(\lambda)$.

Definition 2.27 (Semistandard Young tableaux). Let $\lambda \vdash n$, and let d be a positive integer. A *semistandard Young tableau* of shape λ and alphabet $[d]$ is a filling of λ with numbers from $[d]$, such that the entries weakly increase from left to right, and strictly increase from top to bottom. The set of all SSYT of shape λ and alphabet $[d]$ is denoted $\text{SSYT}(\lambda, d)$.

1	2	4	6
3	5		
7			

1	1	2	3
2	3		
3			

Figure 4: Let $\lambda = (4, 2, 1)$. On the left, a standard Young tableau of shape λ . On the right, a semistandard Young tableau of shape λ and alphabet $[3]$.

Notation 2.28. We use the following additional notation.

- We also write $\square$ to denote the Young diagram corresponding to $\lambda = (1)$.
- The box in the i -th row and j -th column is denoted by (i, j) .
- The *content* of (i, j) is $\text{cont}_\lambda(i, j) := j - i$.
- The *hook* of (i, j) is the set of boxes in λ , (k, ℓ) , such that $i = k$ and $j \leq \ell$, or $j = \ell$ and $i \leq k$. The *hook-length* of (i, j) , denoted $\text{hook}_\lambda(i, j)$, is the number of boxes in the hook of (i, j) .
- If S is an SYT of shape λ , and $m \in [n]$ is the label of box (i, j) in S , then we also write $\text{cont}_S(m) := \text{cont}_\lambda(i, j)$.
- Let λ and μ are Young diagrams such that $\ell(\mu) \geq \ell(\lambda)$, and suppose for each $i \in \ell(\lambda)$ we have $\lambda_i \leq \mu_i$. We can view λ as a subset of the boxes of μ , and write $\lambda \subseteq \mu$. We also write $\mu \setminus \lambda$ for the subset of boxes of μ not in λ .
- If $\lambda \subseteq \mu$ and $|\mu \setminus \lambda| = 1$, then we write $\lambda \nearrow \mu$. The Young diagram of μ can be obtained by adding a single box to some row of λ . If this is the i -th row, we also write $\mu = \lambda + e_i$, or $\lambda = \mu - e_i$.
- If $\lambda \subseteq \mu$, and μ can be obtained from λ by adding some number of boxes to λ , such that none of these additional boxes are in the same column, we write $\lambda \lesssim \mu$. Equivalently, the row-lengths of λ and μ *interlace*, meaning $\mu_{i+1} \leq \lambda_i \leq \mu_i$ for all i , so that $\lambda_d \leq \mu_d \leq \lambda_{d-1} \leq \dots \leq \mu_2 \leq \lambda_1 \leq \mu_1$.

2.4.3 The irreducible representations of S_n

We now describe the irreducible representations of the symmetric group. For more, see [\[Sag01\]](#) for general theory, [\[Har05, Chapter 7\]](#) or [\[VO05\]](#) for subgroup adapted bases, and [\[VO05\]](#) for Young's orthogonal representation and the Jucys-Murphy element.

Theorem 2.29. *The irreducible representations of S_n are indexed by partitions $\lambda \vdash n$.*

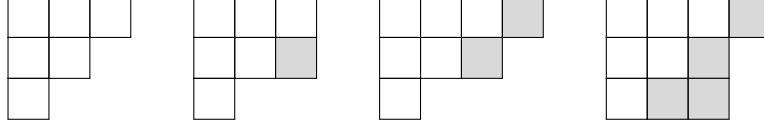


Figure 5: Four Young diagrams. On the left is $\lambda = (3, 2, 1)$. The remaining Young diagrams are μ_i , for $i = 1, 2, 3$. Each μ_i satisfies $\lambda \subseteq \mu_i$, with the boxes of $\mu \setminus \lambda$ shaded. We have $\lambda \nearrow \mu_1$, with $\mu_1 = \lambda + e_2$, since the additional box is in the second row. The shaded box in μ_1 is $(2, 3)$, and the content of this cell is $\text{cont}(2, 3) = 1$. Both μ_1 and μ_2 satisfy $\lambda \leq \mu_i$, since there is at most one added box per column. However, $\lambda \not\leq \mu_3$, since there are two additional boxes in the third column.

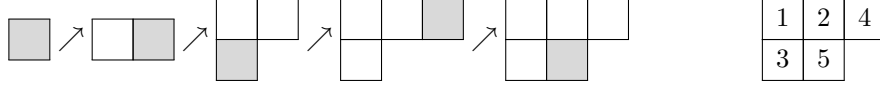


Figure 6: On the left, a sequence of Young diagrams $(\lambda^{(i)})_{i \in [5]}$, each obtained from the last by adding a single box. The new box is shaded in each diagram. On the right, the SYT of shape $\lambda^{(5)}$ corresponding to the chain. The box added at the i -th step is labeled i .

Definition 2.30. The irrep corresponding to λ is $(\kappa_\lambda, \text{Sp}_\lambda)$, where Sp_λ is called the *Specht module*. We will abbreviate $\dim(\kappa_\lambda)$ by $\dim(\lambda)$.

The symmetric group has the following branching rule.

Theorem 2.31. Let $\lambda \vdash n$ be a partition. Consider S_{n-1} as a subgroup of S_n via the embedding that maps $\pi \in S_{n-1}$ to the permutation on $[n]$ that fixes n and permutes everything else according to π . The branching rule for κ_λ , restricted to S_{n-1} , is

$$(\kappa_\lambda) \downarrow_{S_{n-1}}^{S_n} \cong \bigoplus_{\mu \nearrow \lambda} \kappa_\mu.$$

Note that the branching rule is multiplicity-free.

A multiplicity-free branching rule can be used to construct an explicit orthonormal basis for each irreducible representation. Such a basis is generally known as a *subgroup adapted basis*. We now outline this construction for the symmetric group. From [Theorem 2.31](#), we have the following decomposition of vector spaces.

$$\text{Sp}_\lambda \cong \bigoplus_{\mu \nearrow \lambda} \text{Sp}_\mu. \quad (13)$$

Taking the summands to be orthogonal reduces the problem to finding orthonormal bases for Sp_μ . We now recursively apply the branching rule to $S_{n-1}, S_{n-2}, \dots, S_1$ to obtain

$$\text{Sp}_\lambda \cong \bigoplus_{\lambda^{(1)} \nearrow \lambda^{(2)} \nearrow \dots \nearrow \lambda^{(n-1)} \nearrow \lambda^n} \text{Sp}_{\lambda^{(1)}}. \quad (14)$$

However, since S_1 is the trivial group, $\text{Sp}_{\lambda^{(1)}}$ is a one-dimensional subspace, and contains a unique vector (up to phase). As a result, Sp_λ decomposes into the direct sum of one-dimensional subspaces indexed by chains of the form $\square = \lambda^{(1)} \nearrow \lambda^{(2)} \nearrow \dots \nearrow \lambda^{(n-1)} \nearrow \lambda^n = \lambda$. Moreover, such chains are in bijection with standard Young tableaux of shape λ , where the number $i \in [n]$ is placed in the box in the unique cell in $\lambda^{(i)} \setminus \lambda^{(i-1)}$ (taking $\lambda^{(0)} = \emptyset$). See [Figure 6](#) for an illustration.

The above construction gives us a canonical decomposition into one-dimensional subspaces. Taking these subspaces to be orthogonal, and choosing a unit-norm vector in each subspace gives us a *Young-Yamanouchi basis*.

Definition 2.32 (Young-Yamanouchi basis). Given $\lambda \vdash n$, a *Young-Yamanouchi basis* for Sp_λ has a unit vector $|S\rangle$ for each standard Young tableau S of shape λ , associated to the corresponding one-dimensional subspace in [Equation \(14\)](#).

Corollary 2.33. $\dim(\lambda) = |\text{SYT}(\lambda)|$.

Young-Yamanouchi bases are canonical up to a choice of phase in each basis vector. *Young's orthogonal basis* is a particularly nice basis, in which each permutation is represented by an orthogonal matrix.

Definition 2.34 (Young's orthogonal basis). Let $\lambda \vdash n$. *Young's orthogonal basis* is a particular Young-Yamanouchi basis, which can be specified by the action of the transpositions $(i, i+1)$, for $i \in [n-1]$. If i and $i+1$ occur in the same row of S , then

$$\kappa_\lambda(i, i+1) \cdot |S\rangle = +|S\rangle.$$

If, instead, i and $i+1$ occur in the same column of S , then

$$\kappa_\lambda(i, i+1) \cdot |S\rangle = -|S\rangle.$$

Finally, if i and $i+1$ occur in neither the same row nor the same column of S , then swapping the locations of i and $i+1$ in S gives us another valid SYT S' , and

$$\kappa_\lambda(i, i+1) \cdot |S\rangle = \frac{1}{\Delta_S(i)} |S\rangle + \sqrt{1 - \frac{1}{\Delta_S(i)^2}} |S'\rangle,$$

where $\Delta_S(i) := \text{cont}_S(i+1) - \text{cont}_S(i)$.

When written in this basis, the representation κ_λ is also known as *Young's orthogonal form*. Henceforth, we will take κ_λ to mean Young's orthogonal form, unless otherwise mentioned.

For a group G , the group algebra $\mathbb{C}[G]$ is the associative algebra whose elements are formal linear combinations of group elements with complex coefficients. The multiplication of basis elements is inherited from the group's operation, and can be linearly extended to multiplication of generic elements. Representations can be extended to the group algebra, as $\rho(\sum_i c_i g_i) = \sum_i c_i \rho(g_i)$.

Jucys-Murphy elements play a prominent role in our analysis.

Definition 2.35 (Jucys-Murphy element). For $i \in [n]$, the k -th *Jucys-Murphy element* is the element of the group algebra $\mathbb{C}[S_n]$ given by

$$X_i := \sum_{j < i} (j, i).$$

Theorem 2.36. Let $\lambda \vdash n$. Consider $\kappa_\lambda(X_k)$ as an operator on Sp_λ . Then *Young's orthogonal basis* is an eigenbasis of $\kappa_\lambda(X_k)$, with

$$\kappa_\lambda(X_i) \cdot |S\rangle = \text{cont}_S(i) \cdot |S\rangle$$

for all standard Young tableaux S of shape λ .

We will sometimes identify X_k with its action on a representation space, and use X_k in place of $\mathcal{P}(X_k)$ or $\kappa_\lambda(X_k)$ whenever the representation is clear from context.

2.4.4 The irreducible representations of $U(d)$

In this section, we describe the irreducible representations of the unitary group. For more, see [GW09].

Definition 2.37 (Polynomial representations). A finite-dimensional representation (μ, V) of a matrix group G is *polynomial*, if, expressed in some basis of V , every matrix element $\mu(g)_{ij}$ is a polynomial in the entries of $g \in G$.

A representation being polynomial is a basis-independent notion, since if the entries of $\mu(g)$ are polynomials in one basis, then they are also polynomials in any other basis.

Theorem 2.38. The polynomial irreps of $U(d)$ are indexed by partitions λ of length at most d . The irrep corresponding to λ is written $(\nu_\lambda, V_\lambda^d)$, where V_λ^d is called the Schur(-Weyl) module.

Since ν_λ is polynomial, we can extend the domain to be all matrices $M \in \mathbb{C}^{d \times d}$. Also, note that here we have an irrep corresponding to the empty partition, $\lambda = \emptyset$.

We now turn to the characters of these irreps.

Definition 2.39. Let $x_1, \dots, x_d$ be indeterminates. For $\lambda \vdash n$, the *Schur polynomial* $s_\lambda(x_1, \dots, x_d)$ is the degree- n homogeneous polynomial defined as $s_\lambda(x_1, \dots, x_d) := \sum_T x^T$, where the sum is over all SSYT's T of shape λ and alphabet $[d]$, and

$$x^T = \prod_{i=1}^d x_i^{w_T(i)}.$$

Here, $w_T(i)$ is the number of boxes of T labeled by i .

Theorem 2.40. Let U be a unitary with eigenvalues $\alpha_1, \dots, \alpha_d$. Then $\chi_{\nu_\lambda}(U) = s_\lambda(\alpha_1, \dots, \alpha_d)$.

We can extend the domain of s_λ to all $M \in \mathbb{C}^{d \times d}$ by $s_\lambda(M) := \text{tr}(\nu_\lambda(M))$.

Next, we construct an explicit orthonormal basis for each polynomial irrep of the unitary group using the following result.

Theorem 2.41. Let λ be a partition with $\ell(\lambda) \leq d$, and fix a basis $\{|i\rangle\}_{i \in [d]}$. For $d \geq 2$, consider the embedding of $U(d-1)$ into $U(d)$ via the following map:

$$U \mapsto \begin{bmatrix} U & 0 \\ 0 & 1 \end{bmatrix}.$$

Under this mapping, we can identify $U(d-1)$ with its image, a subgroup of $U(d)$. When restricted to $U(d-1)$, we have the multiplicity-free branching rule

$$(\nu_\lambda)_{\downarrow U(d-1)}^{U(d)}(U) \cong \bigoplus_{\substack{\mu \preceq \lambda \\ \ell(\mu) \leq d-1}} \nu_\mu(U).$$

We can use this branching rule to decompose V_λ^d as

$$V_\lambda^d \cong \bigoplus_{\substack{\mu \preceq \lambda \\ \ell(\mu) \leq d-1}} V_\mu^{d-1}.$$

As in the construction of Young's orthogonal basis in [Section 2.4.3](#), we can recursively restrict, eventually decomposing

$$V_\lambda^d \cong \bigoplus_{\substack{\lambda^{(1)} \preceq \dots \preceq \lambda^{(n)} \\ \ell(\lambda^{(i)}) \leq i}} V_{\lambda^{(1)}}^1,$$

where $\lambda^{(n)} = \lambda$. However, since $U(1)$ is Abelian, and since $V_{\lambda^{(1)}}^1$ an irrep of $U(1)$, $V_{\lambda^{(1)}}^1$ is a one-dimensional representation, and contains a unique vector, up to phase. Hence, V_λ^d is isomorphic to a direct sum of one-dimensional subspaces labeled by chains

$$\lambda^{(1)} \preceq \dots \preceq \lambda^{(n)} = \lambda,$$

where $\ell(\lambda^{(i)}) \leq i$. Every such chain can be associated bijectively to a semistandard Young tableau of shape λ and alphabet $[d]$, by filling the boxes of $\lambda^{(i)} \setminus \lambda^{(i-1)}$ with the number i , for each $i \in [d]$. See [Figure 7](#) for an illustration.

Choosing a vector in each one-dimensional subspace gives a *Gelfand-Tsetlin basis*.

Definition 2.42 (Gelfand-Tsetlin basis). Given λ with $\ell(\lambda) \leq d$, the *Gelfand-Tsetlin (GT) basis* for V_λ^d has a unit vector $|T\rangle$ for each semistandard Young tableau T of shape λ and alphabet $[d]$.

Corollary 2.43. $\dim(V_\lambda^d) = |\text{SSYT}(\lambda, d)|$.

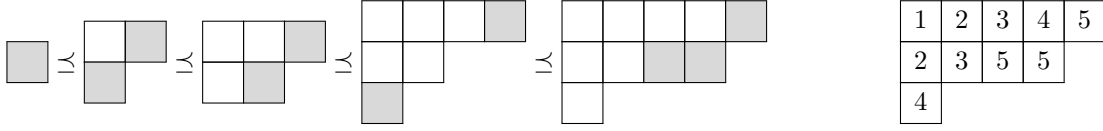


Figure 7: On the left, a sequence of Young diagrams $(\lambda^{(i)})_{i \in [5]}$, such that $\lambda^{(i)} \preceq \lambda^{(i+1)}$ and $\ell(\lambda^{(i)}) \leq i$. The new boxes added at each step are shaded in each diagram. On the right, the SSYT of shape $\lambda^{(5)}$ corresponding to the chain. The boxes added at the i -th step are labeled i .

The dimension of V_λ^d can be computed using the Weyl dimension formula [GW09, Equation (7.18)]:

$$\dim(V_\lambda^d) = \prod_{1 \leq i < j \leq d} \frac{(\lambda_i - \lambda_j) + (j - i)}{j - i}. \quad (15)$$

It is also given by Stanley's hook-content formula [Sta99, Equation (7.21.2)]:

$$\dim(V_\lambda^d) = \prod_{(i,j) \in \lambda} \frac{d + \text{cont}_\lambda(i,j)}{\text{hook}_\lambda(i,j)}. \quad (16)$$

Certain GT basis vectors, called the *highest weight vectors*, play a prominent role in Keyl's algorithm and our debiased version.

Definition 2.44 (Highest weight SSYTs, highest weight vectors). The *highest weight SSYT*, denoted T^λ , is the SSYT of shape λ for which every box in the i -th row is labelled i . The corresponding GT basis vector, $|T^\lambda\rangle$, is called the *highest weight vector*.

1	1	1	1	1
2	2	2	2	
3				

Figure 8: The highest weight SSYT of shape $\lambda = (5, 4, 1)$.

We conclude this section with an example which will be particularly important to us later on: the irrep corresponding to $\lambda = (1)$.

Example 2.45. The irrep $V_{(1)}^d$ turns out to be isomorphic to the *defining representation*, in which which U acts on $\mathbb{C}^d$ as itself. The Gelfand-Tsetlin basis is the computational basis $\{|i\rangle\}_{i \in [d]}$ (up to phase).

2.4.5 Schur-Weyl duality

In state tomography, we are given as input n copies of some unknown state ρ , i.e. we are given the state $\rho^{\otimes n}$. In this setting, there are two particularly natural representations of the groups S_n and $U(d)$, acting on $(\mathbb{C}^d)^{\otimes n}$.

Definition 2.46. The groups S_n and $U(d)$ have the following representations acting on $(\mathbb{C}^d)^{\otimes n}$. The representation $\mathcal{P}^{(n)}(\pi)$ acts by permuting the n tensor factors according to π , i.e. $\mathcal{P}^{(n)}(\pi)$ acts on a standard basis element $|i_1\rangle \otimes \cdots \otimes |i_n\rangle$ as

$$\mathcal{P}^{(n)}(\pi) |i_1\rangle \otimes \cdots \otimes |i_n\rangle = |i_{\pi^{-1}(1)}\rangle \otimes \cdots \otimes |i_{\pi^{-1}(n)}\rangle.$$

The representation $\mathcal{Q}^{(n,d)}(U)$ acts on each tensor factor as U , i.e. $\mathcal{Q}^{(n,d)}(U)$ acts on a standard basis element as

$$\mathcal{Q}^{(n,d)}(U) |i_1\rangle \otimes \cdots \otimes |i_n\rangle = U |i_1\rangle \otimes \cdots \otimes U |i_n\rangle.$$

We will often just write $\mathcal{P}(\pi)$ and $\mathcal{Q}(U)$, when n and d are clear from context.

Since $\mathcal{P}(\pi)$ and $\mathcal{Q}(U)$ commute for any choices of π and U , the product $\mathcal{P} \cdot \mathcal{Q}$ forms a representation of the product group $S_n \times U(d)$. Schur-Weyl duality provides us with a nice decomposition of this representation into products of irreps of S_n and $U(d)$.

Theorem 2.47 (Schur-Weyl duality). *Consider the representation $\mathcal{P} \cdot \mathcal{Q}$ of $S_n \times U(d)$. The action of this representation on $(\mathbb{C}^d)^{\otimes n}$ decomposes as*

$$(\mathbb{C}^d)^{\otimes n} \cong \bigoplus_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} \text{Sp}_\lambda \otimes V_\lambda^d.$$

Therefore, there exists a fixed unitary $\mathcal{U}_{\text{SW}}^{(n)}$ such that for all $\pi \in S_n$ and $U \in U(d)$, we have

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \mathcal{P}(\pi) \mathcal{Q}(U) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\pi) \otimes \nu_\lambda(U). \quad (17)$$

In [Section 11](#) we will construct the unitary $\mathcal{U}_{\text{SW}}^{(n)}$, though for now, it suffices to know that such a unitary exists. After changing basis with $\mathcal{U}_{\text{SW}}^{(n)}$, we will refer to the second register as the *permutation register*, and the third register as the *unitary register*, since S_n and $U(d)$ act on these registers only. The collection of vectors of the form $|\lambda\rangle \otimes |S\rangle \otimes |T\rangle$, with $\lambda \vdash n$, $S \in \text{SYT}(\lambda)$, and $T \in \text{SSYT}(\lambda, d)$ forms the *Schur basis*. The λ -subspace is the subspace spanned by Schur basis vectors with fixed λ . We write Π_λ for the projector defined by

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \Pi_\lambda \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes I_{\dim(V_\lambda^d)}. \quad (18)$$

Recalling that the domain of $\mathcal{Q}$ can be extended to all matrices $M \in \mathbb{C}^{d \times d}$, we can apply Schur-Weyl duality to the state $\rho^{\otimes n}$, regarded as $\mathcal{P}(e) \cdot \mathcal{Q}(\rho)$, with e the identity permutation. This gives us the following result.

Corollary 2.48 (Schur-Weyl duality for quantum states). *There is a fixed unitary change of basis $\mathcal{U}_{\text{SW}}^{(n)}$, and therefore an allowed quantum mechanical transformation, that puts an arbitrary input state $\rho^{\otimes n}$ into the following form:*

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \rho^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(\rho).$$

In other words, there is a fixed change-of-basis unitary which simultaneously block-diagonalizes all possible input states.

As a further example, which will be useful to us later, we have

$$\begin{aligned} \mathcal{U}_{\text{SW}}^{(n)} \cdot \mathcal{P}(X_k) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} &= \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(X_k) \otimes I_{\dim(V_\lambda^d)} \\ &= \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} \sum_{S \in \text{SYT}(\lambda)} \text{cont}_S(k) \cdot |\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes I_{\dim(V_\lambda^d)}, \end{aligned} \quad (19)$$

using [Theorems 2.36](#) and [2.47](#). Thus, in the Schur basis, the Jucys-Murphy element acts diagonally.

2.4.6 Weak Schur sampling

Definition 2.49 (Weak Schur sampling). *Weak Schur sampling* (WSS) refers to performing the projective measurement $\{\Pi_\lambda\}_{\lambda \vdash n, \ell(\lambda) \leq d}$ on $\rho^{\otimes n}$, and induces a probability distribution on partitions $\lambda \vdash n$, denoted $\text{WSS}_n(\rho)$.

If ρ has spectrum $\alpha = (\alpha_1, \dots, \alpha_d)$, weak Schur sampling yields outcome λ with probability

$$\Pr_{\lambda \sim \text{WSS}_n(\rho)}[\lambda = \lambda] = \text{tr}(\Pi_\lambda \rho^{\otimes n}) = \dim(\lambda) \cdot \text{tr}(\nu_\lambda(\rho)) = \dim(\lambda) \cdot s_\lambda(\alpha). \quad (20)$$

Remark 2.50. Suppose ρ has rank r . Recall that the boxes of an SSYT are strictly increasing as we move down a column. Therefore, any SSYT with more than r rows necessarily has a box containing a number larger than r . Therefore, each term of $s_\lambda(\alpha)$ contains at least one α_i for $i > r$ (see [Theorem 2.40](#)). So if $\ell(\lambda) > r$, then $s_\lambda(\alpha) = 0$. That is, we always receive a Young diagram λ with $\ell(\lambda) \leq r$.

If we perform WSS on $\rho^{\otimes n}$, and obtain outcome λ , the resulting post-measurement state is

$$\frac{\Pi_\lambda \cdot \rho^{\otimes n} \cdot \Pi_\lambda}{\text{tr}(\Pi_\lambda \cdot \rho^{\otimes n})} = \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes \frac{I_{\dim(\lambda)}}{\dim(\lambda)} \otimes \frac{\nu_\lambda(\rho)}{s_\lambda(\alpha)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n)}. \quad (21)$$

Note that, in the Schur basis, the first two registers contain no quantum information, and may be discarded as they are classically simulable.

3 The debiased Keyl's algorithm

Here, we formally introduce our unbiased entangled tomography algorithm. In [Section 3.1](#), we describe Keyl's algorithm. In [Section 3.2](#), we formally state our debiased version of Keyl's algorithm. Lastly, in [Section 3.3](#) we show that in addition to our debiased Keyl's algorithm, there is an entire family of unbiased estimators for quantum state tomography. Some of these estimators have simpler descriptions than our original estimator and will be useful in some intermediate steps of our proofs. We also show that our original debiased Keyl's algorithm is the family member with the minimal variance.

3.1 Keyl's algorithm

Keyl's algorithm leverages tools from representation theory to perform quantum state tomography [\[Key06\]](#). It starts by performing weak Schur sampling, and uses the Young diagram $\lambda \vdash n$ it receives to form an estimate for the spectrum of ρ , as described in [Section 2.4.6](#). It continues by measuring in the irrep V_λ^d to learn an estimate for the change of basis unitary that rotates the computational basis onto the eigenbasis of ρ , using the highest weight vector introduced in [Theorem 2.42](#). Finally, it combines these ingredients to produce an estimate for the state ρ itself.

Definition 3.1 (Keyl's algorithm). Given n copies of a quantum state $\rho \in \mathbb{C}^d$:

1. Perform WSS on $\rho^{\otimes n}$ to obtain a Young diagram $\lambda \vdash n$, and the corresponding state ρ_λ , as in [Equation \(21\)](#). Discard the first two registers, leaving only $\nu_\lambda(\rho)/s_\lambda(\alpha)$. Set $\hat{\alpha}_0 = \lambda/n$.
2. Measure in V_λ^d using the POVM M_λ , where $M_\lambda = \{M_{\lambda,U}\}_U$ is a POVM whose elements are labeled by $U \in U(d)$ and defined as follows:

$$M_{\lambda,U} = \dim(V_\lambda^d) \cdot \nu_\lambda(U) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(U)^\dagger \cdot dU,$$

where dU is the Haar measure on $U(d)$. Let U be the measurement outcome. We write $U \sim K_\lambda(\rho)$.

3. Output $\hat{\rho} := U \cdot \hat{\alpha}_0 \cdot U^\dagger$.

Equivalently, Keyl's algorithm can be described by a single measurement, with outcomes labeled by the pair (λ, U) , and

$$M_{\lambda,U} = \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(\dim(V_\lambda^d) \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(U) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(U)^\dagger \cdot dU \right) \cdot \mathcal{U}_{\text{SW}}^{(n)}. \quad (22)$$

To see that the POVM described in the second step of the algorithm is a valid POVM, let us define

$$N_\lambda := \int_U M_{\lambda,U} = \dim(V_\lambda^d) \int_U \nu_\lambda(U) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(U)^\dagger \cdot dU.$$

For M_λ to form a POVM, we need $N_\lambda = I_{\dim(V_\lambda^d)}$. We note that for any $V \in U(d)$,

$$\begin{aligned} \nu_\lambda(V) N_\lambda \nu_\lambda(V)^\dagger &= \dim(V_\lambda^d) \int_U \nu_\lambda(VU) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(VU)^\dagger \cdot dU \\ &= \dim(V_\lambda^d) \int_W \nu_\lambda(W) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(W)^\dagger \cdot dW = N_\lambda, \end{aligned}$$

so that N_λ is an intertwining operator. Here we have defined $W = VU$, and used the left-invariance of the Haar measure. Then, since ν_λ is an irreducible representation, Schur's lemma ([Theorem 2.16](#)) implies that $N_\lambda = \text{tr}(N_\lambda) \cdot I_{\dim(V_\lambda^d)} / \dim(V_\lambda^d)$. However,

$$\text{tr}(N_\lambda) = \dim(V_\lambda^d) \int_U \text{tr}(\nu_\lambda(U) |T_\lambda\rangle\langle T_\lambda| \nu_\lambda(U)^\dagger) \cdot dU = \dim(V_\lambda^d) \int_U dU = \dim(V_\lambda^d),$$

so that $N_\lambda = I_{\dim(V_\lambda^d)}$.

We next provide some intuition for why the second step is able to learn a unitary that approximately diagonalizes ρ . We have

$$\text{tr}(M_{\lambda,U} \cdot \nu_\lambda(\rho)) = \dim(V_\lambda^d) \cdot \langle T_\lambda | \nu_\lambda(U^\dagger \rho U) | T_\lambda \rangle. \quad (23)$$

As noted by Keyl [[Key06](#), Equation 141], it turns out that the matrix elements $\langle T_\lambda | \nu_\lambda(U^\dagger \rho U) | T_\lambda \rangle$ can be re-expressed as $\Delta_\lambda(U^\dagger \rho U)$. This quantity is defined for a matrix $Z \in \mathbb{C}^{d \times d}$ as

$$\Delta_\lambda(Z) := \prod_{i=1}^d \text{pm}_i(Z)^{\lambda_i - \lambda_{i+1}},$$

where pm_i is the i -th principal minor of Z , and we take $\lambda_{d+1} = 0$. Since $U^\dagger \rho U$ has eigenvalues $\alpha_1 \geq \dots \geq \alpha_d$, and $\text{pm}_i(U^\dagger \rho U) \leq \prod_{j=1}^i \alpha_j$, we have

$$\langle T_\lambda | \nu_\lambda(U^\dagger \rho U) | T_\lambda \rangle = \prod_{i=1}^d \text{pm}_i(U^\dagger \rho U)^{\lambda_i - \lambda_{i+1}} \leq \prod_{i=1}^d \prod_{j=1}^i \alpha_i^{\lambda_i - \lambda_{i+1}} = \prod_{i=1}^d \alpha_i^{\lambda_i}.$$

This maximum is attained if and only if $U^\dagger \rho U$ is diagonal in the computational basis, with decreasingly sorted eigenvalues. In other words, this maximum is attained if $\rho = U \cdot \text{diag}(\alpha_1, \dots, \alpha_d) \cdot U^\dagger$, so that U is the change of basis unitary rotating the computational basis onto the eigenbasis of ρ . Therefore, the probability density of obtaining U is maximal at the correct unitary.

3.2 The debiased Keyl's algorithm

Keyl's algorithm measures a Young diagram $\lambda \vdash n$ and a unitary matrix $U \in U(d)$, and outputs $U \cdot (\lambda/n) \cdot U^\dagger$. In order to debias Keyl's algorithm, it turns out that we can perform a simple, deterministic transformation on λ , called *donation*, and output $U \cdot (\text{donate}(\lambda)/n) \cdot U^\dagger$. Donation and the debiased Keyl's algorithm were both defined in the introduction, but we restate their definitions here for convenience.

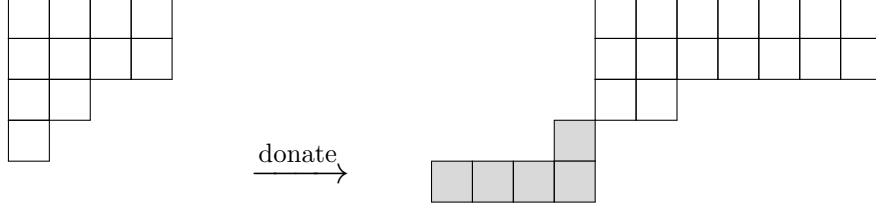


Figure 9: Take $d = 5$. On the left is the Young diagram corresponding to $\lambda = (4, 4, 2, 1, 0)$, with $d = 5$. On the right is $\text{donate}(\lambda) = (7, 7, 2, -1, -4)$. The gray-colored boxes correspond to rows with a negative length. Note that $\text{donate}(\lambda)$ is not necessarily a Young diagram, since it may have negative entries, but $\text{donate}(\lambda)$ is a weakly decreasing list of integers with sum $|\lambda|$, and $\text{donate}(\lambda)$ is constant on the blocks of λ .

Definition 3.2 (Young diagram box donation; [Theorem 1.1](#), restated). Let $\lambda = (\lambda_1, \dots, \lambda_d)$ be a Young diagram. Then $\text{donate}(\lambda) \in \mathbb{Z}^d$ is the vector defined as

$$\text{donate}(\lambda)_i = \lambda_i - \sum_{j=1}^{i-1} 1[\lambda_j > \lambda_i] + \sum_{j=i+1}^d 1[\lambda_i > \lambda_j].$$

Intuitively, each row of the Young diagram λ “donates” a box to each row which starts off longer than it and receives a box from each row which starts off shorter than it. Note that $\text{donate}(\lambda)$ need not be a Young diagram, even if λ is, as it might have negative entries. We will also use the notation $\lambda^\uparrow$ for $\text{donate}(\lambda)$.

See [Figure 9](#) for an example of donation.

Definition 3.3 (Debiased Keyl’s algorithm; [Theorem 1.2](#), restated). Given n copies of a quantum state $\rho \in \mathbb{C}^{d \times d}$, the *debiased Keyl’s algorithm* works as follows.

1. Measure $\rho^{\otimes n}$ as in Keyl’s algorithm ([Theorem 3.1](#)), producing a Young diagram $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ and a unitary matrix $U \in U(d)$.
2. Set $\hat{\alpha} = \lambda^\uparrow / n$. Output $\hat{\rho} = U \cdot \hat{\alpha} \cdot U^\dagger$.

In [Section 13](#), we prove that the debiased Keyl’s algorithm is, in fact, an unbiased estimator of ρ .

3.3 A family of unbiased estimators

Note that our estimator is equivalent to performing Keyl’s algorithm, and then replacing the empirical Young diagram λ/n with $\hat{\alpha}$. It turns out that in addition to the algorithm given in [Theorem 3.3](#), there is a family of related unbiased estimators that use slightly different Young diagram transformations than $\lambda^\uparrow$. In this section, we describe this family of unbiased estimators, while highlighting a few particular members that we will use in our proofs. We end this section by showing that the original estimator $\hat{\alpha}$ is the family member with the minimum variance.

We start by defining the set of legal Young diagram transformations, which generalizes the box donation from [Theorem 1.1](#).

Definition 3.4 (Legal Young diagram transformation). A map f from Young diagrams to $\mathbb{R}^d$ is a *legal Young diagram transformation* if for all Young diagrams $\lambda = (\lambda_1, \dots, \lambda_d)$, $\mu = f(\lambda) \in \mathbb{R}^d$ is a vector that satisfies

$$\sum_{i \in B} \mu_i = \sum_{i \in B} \lambda_i^\uparrow$$

for every block of rows B in λ .

In words, any Young diagram transformation f of $\lambda \vdash n$ is legal if the total number of “boxes” in a block B after transformation f (allowing for the possibility of negative or, more generally, real-valued amounts of boxes) equals the number of boxes in that block after the box donation transformation of [Theorem 1.1](#). An

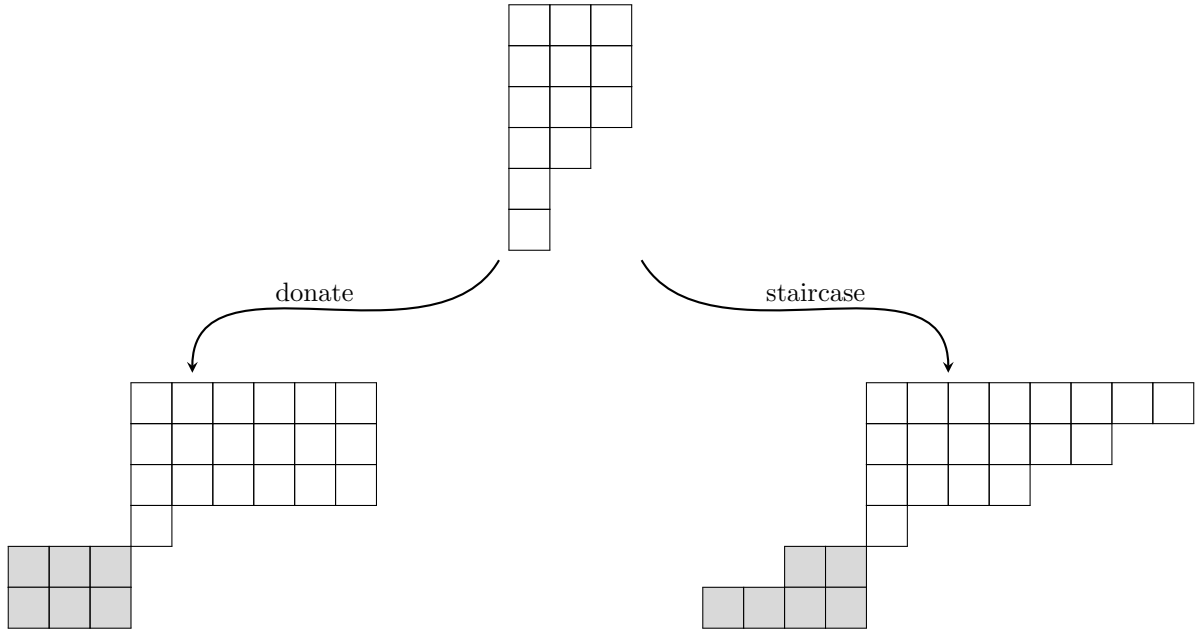


Figure 10: Two examples of legal Young diagram transformations for $\lambda = (3, 3, 3, 2, 1, 1)$. The Young diagram λ has 3 blocks that consist of the first three rows, the fourth row, and the last two rows, respectively. The diagram on the left is the one produced from the original box donation algorithm: $\lambda^\uparrow = (6, 6, 6, 1, -3, -3)$. The one on the right can be obtained by moving boxes within its blocks, and it coincides with the staircase box donation from [Theorem 3.7](#): $\lambda^{\uparrow\uparrow} = (8, 6, 4, 1, -2, -4)$. The gray-colored boxes correspond to rows with negative lengths.

equivalent condition is that any legal transformation f can be obtained from $\lambda^\uparrow$ by arbitrarily redistributing boxes among the rows in the same block B . In addition to having real-valued lengths which may be negative or non-integral, the resulting shape is not necessarily a Young diagram, also due to the possibility of non-monotone row lengths, as illustrated in [Figure 12](#). Given any legal Young diagram transformation, we can define an unbiased estimator for entangled state tomography.

Definition 3.5 (Family of unbiased estimators). Suppose we are given n copies of a mixed state $\rho \in \mathbb{C}^{d \times d}$. For any legal Young diagram transformation f , we define an estimator as follows.

1. Measure $\rho^{\otimes n}$ as in Keyl’s algorithm, producing a Young diagram $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ and a unitary matrix $U \in U(d)$.
2. Set $\lambda' = f(\lambda)$ and $\hat{\alpha}' = \lambda'/n$. Output $\hat{\rho}' = U \cdot \hat{\alpha}' \cdot U^\dagger$.

We will often use the term *legal estimator* to refer to such an estimator.

We will prove in [Section 13](#) that all legal Young diagram transformations produce unbiased estimators.

Theorem 3.6 (All legal estimators are unbiased). Let $\hat{\rho}'$ be the output of a legal estimator when run on $\rho^{\otimes n}$. Then $\mathbb{E}[\hat{\rho}'] = \rho$.

Below, we highlight a few legal estimators that will be useful for proving the results in this paper. We start by defining the *staircase estimator*, which is perhaps the simplest one, as its defining Young diagram transformation can be described using a closed-form expression.

Definition 3.7 (Young diagram staircase box donation). Let $\lambda = (\lambda_1, \dots, \lambda_d)$ be a Young diagram. Then $\text{staircase}(\lambda) \in \mathbb{Z}^d$ is the vector defined

$$\text{staircase}(\lambda)_i = \lambda_i - \sum_{j=1}^{i-1} 1 + \sum_{j=i+1}^d 1 = \lambda_i - (i-1) + (d-i) = \lambda_i + d - 2i + 1.$$

Compared to the original box donation from [Theorem 1.1](#), in the staircase donation transformation, each row of the Young diagram λ donates a box to each row above it, regardless of whether they started with the same row length. We can also model it as a two-step process, where each row first donates boxes according to the original box donation algorithm, and then the rows further donate one box to each row above them in the same block. Since the second step only involves moving boxes within the same block, it does not change the total number of boxes in each block. This confirms that the staircase donation transformation is a legal donation transformation. We include a diagram of $\text{staircase}(\lambda)$ in [Figure 10](#). Similar to $\text{donate}(\lambda)$, $\text{staircase}(\lambda)$ need not be a Young diagram, due to potentially negative entries.

Using this Young diagram transformation, we define the *staircase estimator* and use the notation $\lambda^{\uparrow\uparrow} := \text{staircase}(\lambda)$ and $\hat{\alpha}^{\uparrow\uparrow} := \lambda^{\uparrow\uparrow}/n$. The final state estimate from this algorithm is denoted by $\hat{\rho}^{\uparrow\uparrow} = U \cdot \hat{\alpha}^{\uparrow\uparrow} \cdot U^\dagger$. We chose the double arrow superscript notation for this specific estimator, since we view the staircase box donation as a more “extreme” version of the original box donation.

The advantage of the staircase estimator is the simple closed-form expression of the transformation $\lambda^{\uparrow\uparrow}$. For this reason, the staircase estimator plays a central role in proving our expressions for the first and second moments. In particular, we first prove in [Section 13](#) that $\hat{\rho}^{\uparrow\uparrow}$ is an unbiased estimator, and then extend this result to all legal estimators, including our original estimator $\hat{\rho}$ as in [Theorem 1.3](#). Moreover, it turns out that the staircase estimator achieves the optimal sample complexity for tomography problems over general states, when the number of copies required is $n = \Omega(d^2)$. However, it can be suboptimal in the special case when ρ is promised to be of low rank. One example is the pure state case, where weak Schur sampling always produces the tableau $\lambda = (n, 0, \dots, 0)$. The two estimators will have spectra

$$\hat{\alpha} = \left(\frac{n+d-1}{n}, -\frac{1}{n}, \dots, -\frac{1}{n} \right), \quad \hat{\alpha}^{\uparrow\uparrow} = \left(\frac{n+d-1}{n}, \frac{d-3}{n}, \frac{d-5}{n}, \dots, \frac{-d+1}{n} \right).$$

Observe that $\hat{\rho}^{\uparrow\uparrow}$ is far from any pure state because the absolute sum of all of its eigenvalues except the first one is going to be $O(d^2/n)$. This implies that $\hat{\rho}^{\uparrow\uparrow}$ is only guaranteed to be close to a pure state when

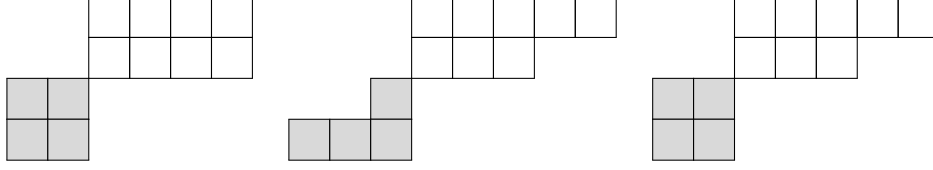


Figure 11: The original, staircase, and rank-2 Young diagram transformations for $\lambda = (2, 2, 0, 0)$.

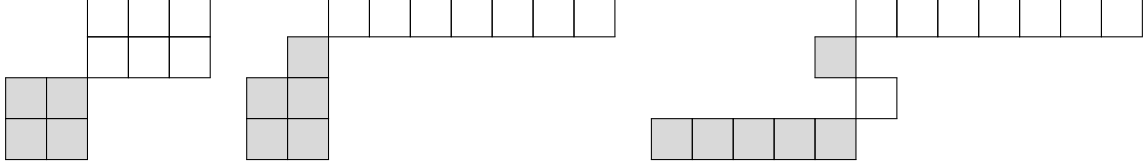


Figure 12: The transformations of $\lambda = (1, 1, 0, 0)$ represented above are also legal transformations, even though some of them seem weird. Indeed, the transformation on the right gives a shape whose row lengths are not monotonically decreasing.

the number of samples is $n \gg d^2$, which is much more than the optimal bound of $O(d/\varepsilon^2)$. The issue with the staircase transformation is that the large negative eigenvalues increase the variance of our estimator, and thus, we require more samples to approximate its true mean. As a result, when dealing with low-rank states, it will not suffice to use the staircase estimator if we want optimal bounds. However, there is a natural modification to the staircase estimator for the case of low-rank states, which reduces its variance while keeping its simplicity.

Definition 3.8 (Young diagram rank- r box donation). Let $\lambda = (\lambda_1, \dots, \lambda_d)$ be a Young diagram, with at most r nonzero rows. Then $\text{staircase}_r(\lambda) \in \mathbb{Z}^d$ is the vector defined as follows.

$$\text{staircase}_r(\lambda)_i = \begin{cases} \lambda_i + d - 2i + 1 & i \leq r \\ -r & i \in \{r+1, \dots, d\}. \end{cases}$$

We first argue that this Young diagram transformation is a legal one. We model it as a two-step process: First, each row donates boxes according to the staircase algorithm. Since λ has at most r nonzero rows, there is a block of rows of λ that includes the last $d - r$ rows. The total number of boxes in the last $d - r$ rows of the block after the staircase donation is $n - \sum_{i=1}^r (\lambda_i + d - 2i + 1) = -(d - r)r$, which is a multiple of $d - r$. The second phase of the box donation will have the last $d - r$ rows donate to each other until they have the same number of boxes, for a total $-r$ each. The last step only involves donations within the same block, so the total number of boxes in each block does not change from the first phase. As the staircase transformation is legal, we conclude that the rank- r donation procedure is also legal.

First, observe that the transformation $\text{staircase}_r(\lambda)$ also has a simple closed-form expression. Using this Young diagram transformation, the final estimator will be $\hat{\rho}' = U \cdot \hat{\alpha}' \cdot U^\dagger$, where $\hat{\alpha}' = \text{staircase}_r(\lambda)/n$. Moreover, when ρ is a rank- r state, the absolute sum of the $d - r$ smallest eigenvalues of $\hat{\rho}'$ will be $O(rd/n)$, compared to $O(d^2/n)$ for $\hat{\rho}^{\uparrow\uparrow}$. Looking forward, we will prove that this “rank- r ” modification to the staircase estimator suffices to achieve the optimal bounds for rank- r states, where the number of copies required is $n = \Omega(rd)$.

Finally, we prove that the original estimator is the one with the minimum variance among all unbiased estimators from [Theorem 3.5](#).

Lemma 3.9. *The estimator that corresponds to the Young diagram box donation of [Theorem 1.1](#) has the minimal variance among all estimators in the family of estimators from [Theorem 3.5](#).*

Proof. Let $\hat{\rho} = U \cdot \hat{\alpha} \cdot U^\dagger$ for $\hat{\alpha} = \lambda^\dagger/n$ be the original unbiased estimator, and let $\hat{\rho}' = U \cdot \hat{\alpha}' \cdot U^\dagger$ be the unbiased estimator that corresponds to a legal Young diagram transformation f . Since f is a legal transformation, it holds that for any λ obtained from weak Schur sampling, $\hat{\alpha}'$ has the same total weight as

$\hat{\alpha}$ on each block of λ , but not necessarily uniformly distributed among each row. The expected distance of $\hat{\rho}'$ from ρ in the Schatten 2-norm is

$$\begin{aligned}
\mathbf{E}\|\hat{\rho}' - \rho\|_2^2 &= \mathbf{E}\|(\hat{\rho}' - \hat{\rho}) + (\hat{\rho} - \rho)\|_2^2 \\
&= \mathbf{E}\|\hat{\rho}' - \hat{\rho}\|_2^2 + 2 \cdot \mathbf{E} \operatorname{tr}((\hat{\rho}' - \hat{\rho})(\hat{\rho} - \rho)) + \mathbf{E}\|\hat{\rho} - \rho\|_2^2 \\
&= \mathbf{E}\|\hat{\rho}' - \hat{\rho}\|_2^2 + 2 \cdot \mathbf{E} \operatorname{tr}((\hat{\rho}' - \hat{\rho})\hat{\rho}) - 2 \cdot \mathbf{E} \operatorname{tr}((\hat{\rho}' - \hat{\rho})\rho) + \mathbf{E}\|\hat{\rho} - \rho\|_2^2 \\
&= \mathbf{E}\|\hat{\rho}' - \hat{\rho}\|_2^2 + 2 \cdot \mathbf{E} \operatorname{tr}((\hat{\rho}' - \hat{\rho})\hat{\rho}) + \mathbf{E}\|\hat{\rho} - \rho\|_2^2 \quad (\mathbf{E}\hat{\rho} = \mathbf{E}\hat{\rho}' = \rho) \\
&= \mathbf{E}\|\hat{\alpha}' - \hat{\alpha}\|_2^2 + 2 \cdot \mathbf{E} \operatorname{tr}((\hat{\alpha}' - \hat{\alpha})\hat{\alpha}) + \mathbf{E}\|\hat{\rho} - \rho\|_2^2. \tag{24}
\end{aligned}$$

In the last line, we have used the fact that $\hat{\rho}'$ and $\hat{\rho}$ are simultaneously diagonalized by U . Now fix a tableau λ and consider a block of indices with λ_i constant. We will analyze the contribution of the first two terms on this block. Without loss of generality, let this block of indices be $\{1, \dots, m\}$, and let $\{\hat{\mu}'_i\}_{i=1}^m$ and $\{\hat{\mu}_i\}_{i=1}^m$ be the entries of $f(\lambda)$ and $\lambda^\dagger$ respectively on the rows in this block. Note that $\hat{\mu}_i = \hat{\mu} := \frac{1}{m} \sum_{j=1}^m \hat{\mu}'_j$ for all $i \in \{1, \dots, m\}$. Then we have

$$\begin{aligned}
\sum_{i=1}^m (\hat{\mu}'_i - \hat{\mu}_i)^2 + 2 \sum_{i=1}^m (\hat{\mu}'_i - \hat{\mu}_i) \hat{\mu}_i &= \sum_{i=1}^m (\hat{\mu}'_i - \hat{\mu})^2 + 2 \sum_{i=1}^m (\hat{\mu}'_i - \hat{\mu}) \hat{\mu} \\
&= \sum_{i=1}^m \hat{\mu}'_i{}^2 - 2\hat{\mu} \sum_{i=1}^m \hat{\mu}'_i + m\hat{\mu}^2 + 2\hat{\mu} \sum_{i=1}^m \hat{\mu}'_i - 2m\hat{\mu}^2 \\
&= \sum_{i=1}^m \hat{\mu}'_i{}^2 - m\hat{\mu}^2 \\
&= m \left(\frac{1}{m} \sum_{i=1}^m \hat{\mu}'_i{}^2 - \left(\frac{1}{m} \sum_{i=1}^m \hat{\mu}'_i \right)^2 \right) \geq 0,
\end{aligned}$$

where the last inequality follows from Cauchy-Schwarz. Since $\hat{\alpha}'$ and $\hat{\alpha}$ are diagonal matrices with entries $f(\lambda)/n$ and $\lambda^\dagger/n$ respectively, we conclude that for every tableau λ produced by weak Schur sampling, the first two terms of Equation (24) are positive, and thus

$$\mathbf{Var}(\hat{\rho}') = \mathbf{E}\|\hat{\rho}' - \rho\|_2^2 \geq \mathbf{E}\|\hat{\rho} - \rho\|_2^2 = \mathbf{Var}(\hat{\rho}). \quad \square$$

Remark 3.10 (Unitary invariance of Keyl measurement). Throughout this paper, we will consider distances or divergences that are unitarily invariant, that is, distances D that satisfy $D(\rho, \sigma) = D(U\rho U^\dagger, U\sigma U^\dagger)$. Any algorithm from our family of unbiased estimators, when run on a state $\rho = V \cdot \alpha \cdot V^\dagger$ with spectrum α , has the form:

$$\lambda \sim \text{SW}^n(\alpha), \quad U \sim K_\lambda(\rho), \quad \hat{\alpha} = \operatorname{diag}(f(\lambda)/n), \quad \hat{\rho} = U \cdot \hat{\alpha} \cdot U^\dagger.$$

Here, f is some legal Young diagram transformation. It holds that

$$D(\hat{\rho}, \rho) = D(U \cdot \hat{\alpha} \cdot U^\dagger, V \cdot \alpha \cdot V^\dagger) = D(V^\dagger U \cdot \hat{\alpha} \cdot U^\dagger V, \alpha).$$

Moreover, from Equation (23) we see that

$$\begin{aligned}
\operatorname{tr}(M_{\lambda, V^\dagger U} \cdot \nu_\lambda(\alpha)) &= \dim(V_\lambda^d) \cdot \langle T_\lambda | \nu_\lambda(U^\dagger V \alpha V^\dagger U) | T_\lambda \rangle \\
&= \dim(V_\lambda^d) \cdot \langle T_\lambda | \nu_\lambda(U^\dagger \rho U) | T_\lambda \rangle = \operatorname{tr}(M_{\lambda, U} \cdot \nu_\lambda(\rho)).
\end{aligned}$$

In words, the probability density of U under $K_\lambda(\rho)$ is equal to the probability density of $V^\dagger U$ under $K_\lambda(\alpha)$. Therefore, when analyzing a unitarily invariant distance, we may assume, without loss of generality, that $\rho = \alpha$ is diagonal, and our algorithm operates as follows:

$$\lambda \sim \text{SW}^n(\alpha), \quad U \sim K_\lambda(\alpha), \quad \hat{\alpha} = \operatorname{diag}(f(\lambda)/n), \quad \hat{\rho} = U \cdot \hat{\alpha} \cdot U^\dagger.$$

Part II

Applications

4 Full state tomography in trace distance

In this section, we formally prove [Theorem 1.5](#), that our debiased Keyl's estimator can learn a rank- r quantum state to trace distance ε , using $O(rd/\varepsilon^2)$ copies of ρ . Our bound matches the upper bound from Corollary 1.4 of [\[OW16\]](#), and the lower bound from Theorem 1.1 of [\[SSW25\]](#).

Theorem 4.1 ([Theorem 1.5](#) restated). *Let ρ be a rank- r quantum state, and $\hat{\rho}$ be the output of the debiased Keyl's algorithm when run on $\rho^{\otimes n}$. Then $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability when $n = O(rd/\varepsilon^2)$.*

Before we prove the main theorem of this section, we prove a bound on the variance of our estimator $\hat{\rho}$.

Lemma 4.2 (Variance of $\hat{\rho}$). *Let ρ be a quantum state, and $\hat{\rho}$ be the output of the debiased Keyl's algorithm when run on $\rho^{\otimes n}$. Then*

$$\text{Var}[\hat{\rho}] = \mathbf{E}\|\hat{\rho} - \rho\|_2^2 \leq \frac{2d}{n} + \frac{d^2 \cdot \mathbf{E} \ell(\boldsymbol{\lambda})}{n^2}.$$

Proof. Since $\hat{\rho}$ is unbiased, we have

$$\mathbf{E}\|\hat{\rho} - \rho\|_2^2 = \mathbf{E} \text{tr}((\hat{\rho} - \rho)^2) = \mathbf{E} \text{tr}(\hat{\rho}^2) - \mathbf{E} \text{tr}(\rho^2) = \mathbf{E} \text{tr}(\hat{\rho}^2) - \text{tr}(\rho^2). \quad (25)$$

We use our expression for the second moment from [Theorem 1.4](#) to write

$$\begin{aligned} \mathbf{E} \text{tr}(\hat{\rho}^2) &= \mathbf{E} \text{tr}(\hat{\rho} \otimes \hat{\rho} \cdot \text{SWAP}) \\ &= \text{tr}(\mathbf{E}[\hat{\rho} \otimes \hat{\rho}] \cdot \text{SWAP}) \\ &= \frac{n-1}{n} \text{tr}(\rho \otimes \rho \cdot \text{SWAP}) + \frac{1}{n} \text{tr}(\rho \otimes I + I \otimes \rho) + \frac{\mathbf{E} \ell(\boldsymbol{\lambda})}{n^2} \cdot \text{tr}(I \otimes I) - \text{tr}(\text{Lower}_\rho \cdot \text{SWAP}) \\ &\leq \text{tr}(\rho^2) + \frac{2d}{n} + \frac{d^2 \cdot \mathbf{E} \ell(\boldsymbol{\lambda})}{n^2}. \end{aligned}$$

In the last step, we are able to drop $\text{tr}(\text{Lower}_\rho \cdot \text{SWAP})$, since Lower_ρ is a positive linear combination of terms of the form $(P \otimes P) \cdot \text{SWAP}$, and since $\text{tr}((P \otimes P) \cdot \text{SWAP} \cdot \text{SWAP}) = \text{tr}(P \otimes P) = \text{tr}(P)^2$. Plugging this back into [Equation \(25\)](#) completes the proof. $\square$

We are now ready to prove the main result of this section.

Proof of [Theorem 4.1](#). Since ρ has rank r , the Young tableau $\boldsymbol{\lambda}$ that we obtain from weak Schur sampling will have at most r nonempty rows (see [Theorem 2.50](#)). Let us decompose $\hat{\rho}$ as

$$\hat{\rho} = U \cdot \hat{\alpha}_{>0} \cdot U^\dagger + U \cdot \hat{\alpha}_{=0} \cdot U^\dagger = \hat{\rho}_{>0} + \hat{\rho}_{=0},$$

where

$$\hat{\alpha}_{>0} := \sum_{i=1}^{\ell(\boldsymbol{\lambda})} \hat{\alpha}_i |i\rangle\langle i| = \sum_{i, \lambda_i > 0} \hat{\alpha}_i |i\rangle\langle i|, \quad \hat{\alpha}_{=0} := \sum_{i=\ell(\boldsymbol{\lambda})+1}^d \hat{\alpha}_i |i\rangle\langle i| = \sum_{i, \lambda_i = 0} \hat{\alpha}_i |i\rangle\langle i|.$$

In words, we split $\hat{\rho}$ into the state we obtain from the top $\ell(\boldsymbol{\lambda})$ eigenvalues (that correspond to nonempty rows of the Young diagram $\boldsymbol{\lambda}$), plus the state from the bottom $d - \ell(\boldsymbol{\lambda})$ eigenvalues (that correspond to the empty rows of $\boldsymbol{\lambda}$). Observe that $\hat{\alpha}_i = -\frac{\ell(\boldsymbol{\lambda})}{n}$ for all $i > \ell(\boldsymbol{\lambda})$, and thus

$$\hat{\rho}_{=0} = U \cdot \left(-\frac{\ell(\boldsymbol{\lambda})}{n} \sum_{i=\ell(\boldsymbol{\lambda})+1}^d |i\rangle\langle i| \right) \cdot U^\dagger = -\frac{\ell(\boldsymbol{\lambda})}{n} \cdot \Pi_{=0},$$

where $\Pi_{=0}$ is some projector of rank $d - \ell(\boldsymbol{\lambda})$. By the triangle inequality,

$$\mathbf{E}\|\hat{\rho} - \rho\|_1 \leq \mathbf{E}\|\hat{\rho}_{>0} - \rho\|_1 + \mathbf{E}\|\hat{\rho}_{=0}\|_1. \quad (26)$$

First, since $\hat{\rho}_{=0} = -\frac{\ell(\lambda)}{n} \cdot \Pi_{=0}$ and $\ell(\lambda) \leq r$:

$$\|\hat{\rho}_{=0}\|_1 = \frac{\ell(\lambda)}{n} \cdot (d - \ell(\lambda)) \leq \frac{rd}{n}. \quad (27)$$

Next, by Cauchy-Schwarz and Jensen's inequality:

$$\|\hat{\rho}_{>0} - \rho\|_1 \leq \sqrt{2r} \cdot \mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2 \leq \sqrt{2r} \cdot \sqrt{\mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2^2}.$$

Then, we have

$$\mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2^2 = \mathbf{E}\|(\hat{\rho} - \rho) - \hat{\rho}_{=0}\|_2^2 = \mathbf{E}\|\hat{\rho} - \rho\|_2^2 - 2\mathbf{E}\text{tr}((\hat{\rho} - \rho)\hat{\rho}_{=0}) + \mathbf{E}\|\hat{\rho}_{=0}\|_2^2.$$

We expand the second term as follows

$$\begin{aligned} 2\mathbf{E}\text{tr}((\hat{\rho} - \rho)\hat{\rho}_{=0}) &= 2\mathbf{E}\text{tr}(\hat{\rho}_{>0} + \hat{\rho}_{=0} - \rho)\hat{\rho}_{=0} \\ &= 2\mathbf{E}\text{tr}(\hat{\rho}_{>0} \cdot \hat{\rho}_{=0}) + 2\mathbf{E}\|\hat{\rho}_{=0}\|_2^2 - 2\mathbf{E}\text{tr}(\rho \cdot \hat{\rho}_{=0}) \\ &= 2\mathbf{E}\|\hat{\rho}_{=0}\|_2^2 - 2\mathbf{E}\text{tr}(\rho \cdot \hat{\rho}_{=0}), \end{aligned}$$

where we have used the fact that $\text{tr}(\hat{\rho}_{>0} \cdot \hat{\rho}_{=0}) = \text{tr}(\hat{\alpha}_{>0} \cdot \hat{\alpha}_{=0}) = 0$, as $\hat{\alpha}_{>0}$ and $\hat{\alpha}_{=0}$ are supported on complementary subsets of computational basis vectors. Thus $\mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2^2$ becomes

$$\mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2^2 = \mathbf{E}\|\hat{\rho} - \rho\|_2^2 + 2\mathbf{E}\text{tr}(\rho \cdot \hat{\rho}_{=0}) - \mathbf{E}\|\hat{\rho}_{=0}\|_2^2 \leq \mathbf{E}\|\hat{\rho} - \rho\|_2^2 + 2\mathbf{E}\text{tr}(\rho \cdot \hat{\rho}_{=0}).$$

Then, since $\ell(\lambda)$ is always at most r , [Theorem 4.2](#) implies that

$$\mathbf{E}\|\hat{\rho} - \rho\|_2^2 \leq \frac{2d}{n} + \frac{rd^2}{n^2}.$$

Moreover,

$$\text{tr}(\rho \cdot \hat{\rho}_{=0}) = -\frac{\ell(\lambda)}{n} \cdot \text{tr}(\rho \cdot \Pi_{=0}) \leq 0.$$

Thus,

$$\mathbf{E}\|\hat{\rho}_{>0} - \rho\|_2^2 \leq \frac{2d}{n} + \frac{rd^2}{n^2}. \quad (28)$$

Plugging [Equations \(27\)](#) and [\(28\)](#) into [Equation \(26\)](#), we get

$$\mathbf{E}\|\hat{\rho} - \rho\|_1 \leq \sqrt{\frac{4rd}{n} + \frac{2r^2d^2}{n^2}} + \frac{rd}{n},$$

which is $O(\varepsilon)$ when $n = O(rd/\varepsilon^2)$. $\square$

Alternative proof using Young diagram statistics. There is another way of proving [Theorem 4.1](#) without relying on our second moment calculation. Let $\hat{\rho}' = U \cdot \hat{\alpha}' \cdot U^\dagger$ be the member of the family of unbiased estimators that uses the rank- r box donation transformation from [Theorem 3.8](#). The spectrum of $\hat{\rho}'$ satisfies

$$\hat{\alpha}'_i = \begin{cases} \frac{\lambda_i + d - 2i + 1}{n} & i \leq r \\ -\frac{r}{n} & i \in \{r+1, \dots, d\}. \end{cases}$$

We will obtain a slightly weaker bound for the variance of $\hat{\rho}$ as follows: $\text{Var}[\hat{\rho}] \leq \text{Var}[\hat{\rho}'] \leq \frac{2d}{n} + \frac{rd^2}{n^2}$. We proved the first inequality in [Theorem 3.9](#), and for the second inequality, we write

$$\begin{aligned} \text{Var}[\hat{\rho}'] &= \mathbf{E}\text{tr}(\hat{\rho}'^2) - p_2(\alpha) && \text{(Equation (25))} \\ &= \mathbf{E}\left[\sum_{i=1}^d \hat{\alpha}'_i{}^2\right] - p_2(\alpha) \\ &= \mathbf{E}\left[\sum_{i=1}^r \frac{(\lambda_i + (d - 2i + 1))^2}{n^2}\right] + \sum_{i=r+1}^d \frac{(-r)^2}{n^2} - p_2(\alpha) \end{aligned} \quad (29)$$

The first term is related to the second moment of the row lengths of the Young diagram random variable λ . Fortunately, these moments can be computed using *Kerov's algebra of observables of diagrams* [Wri16]. In particular, the following member of this algebra captures the second moment of a Young diagram λ :

$$p_2^*(\lambda) = \sum_{i=1}^d \left(\left(\lambda_i - i + \frac{1}{2} \right)^2 - \left(-i + \frac{1}{2} \right)^2 \right) = \sum_{i=1}^d (\lambda_i^2 - \lambda_i(2i-1)).$$

Proposition 2.34 of [OW15b] computes the expectation of this polynomial to be

$$\mathbf{E}[p_2^*(\lambda)] = n(n-1)p_2(\alpha). \quad (30)$$

We conclude that

$$\begin{aligned} (29) &\leq \frac{1}{n^2} \mathbf{E} \left[\sum_{i=1}^d (\lambda_i^2 - \lambda_i(2i-1)) \right] + \frac{1}{n^2} \mathbf{E} \left[\sum_{i=1}^r \lambda_i(2d-2i+1) \right] + \frac{1}{n^2} \sum_{i=1}^r (d-2i+1)^2 + \frac{r^2 d}{n^2} - p_2(\alpha) \\ &\leq \frac{1}{n^2} \cdot \mathbf{E}[p_2^*(\lambda)] + \frac{2d}{n^2} \mathbf{E} \left[\sum_{i=1}^r \lambda_i \right] + \frac{rd^2}{n^2} + \frac{r^2 d}{n^2} - p_2(\alpha) \\ &\leq \frac{2d}{n} + \frac{rd^2}{n^2}. \end{aligned} \quad (\text{Equation (30)}, \sum_{i=1}^r \lambda_i = n)$$

Given the variance bound, we can prove the statement of Theorem 4.1 using the same proof.

5 Full state tomography in Bures distance

In this section, we formally prove Theorem 1.6, that our debiased Keyl's estimator can learn a rank- r quantum state to Bures distance ε , using $O(rd/\varepsilon^2)$ copies of ρ . Equivalently, $O(rd/\delta)$ samples suffice to learn to infidelity error δ . This matches the lower bound proven in [Yue23].

This section is structured as follows.

1. As a warm-up, we present in Section 5.1 a simpler argument that shows how to learn a full-rank quantum state ρ to Bures distance ε using $O(d^2/\varepsilon^2)$ copies.
2. The complete argument for rank- r states appears in Section 5.2.
3. As a further application of our Bures distance tomography result, we present in Section 5.3 an algorithm for learning a bipartite pure state $|\psi\rangle_{AB}$ with Schmidt rank r using $O(r(d_A + d_B)/\varepsilon^2)$ copies of $|\psi\rangle_{AB}$.

Let us now make a couple of remarks that will be useful in multiple subsections below. First, we will use the Bures χ^2 -divergence to bound the Bures distance, both of which are unitarily invariant. Following Theorem 3.10, we will therefore assume throughout this section, without loss of generality, that the state we are trying to learn ρ is diagonal with sorted eigenvalues $\alpha = (\alpha_1, \dots, \alpha_d)$. Second, we will use the following lemma to upper bound the expectation of the square of the entries of $\hat{\rho}$.

Lemma 5.1. *Given a rank- r diagonal mixed state $\rho \in \mathbb{C}^{d \times d}$ with eigenvalues $\alpha_1, \dots, \alpha_d$, let $\hat{\rho}$ be the output of the debiased Keyl's algorithm on $\rho^{\otimes n}$. Then*

$$\mathbf{E}[|\hat{\rho}_{ij}|^2] \leq |\rho_{ij}|^2 + \frac{\alpha_i + \alpha_j}{n} + \frac{r}{n^2}.$$

Proof. Using the fact that $\hat{\rho}$ is Hermitian, we write

$$\mathbf{E}[|\hat{\rho}_{ij}|^2] = \mathbf{E}[\hat{\rho}_{ij} \cdot \hat{\rho}_{ji}] = \mathbf{E}[\langle i | \hat{\rho} | j \rangle \cdot \langle j | \hat{\rho} | i \rangle] = \langle ij | \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}] \cdot |ji\rangle.$$

We use our second moment result from [Theorem 1.4](#) to deduce that $\mathbf{E}|\widehat{\rho}_{ij}|^2$ is equal to

$$\begin{aligned}
&= \frac{n-1}{n} \langle ij | \rho \otimes \rho | ji \rangle + \frac{1}{n} \langle ij | (\rho \otimes I + I \otimes \rho) \cdot \text{SWAP} | ji \rangle + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \cdot \langle ij | \text{SWAP} | ji \rangle - \langle ij | \text{Lower}_\rho | ji \rangle \\
&= \frac{n-1}{n} \cdot |\rho_{ij}|^2 + \frac{1}{n} \cdot \langle ij | (\rho \otimes I + I \otimes \rho) | ij \rangle + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} - \langle ij | \text{Lower}_\rho | ji \rangle \\
&\leq |\rho_{ij}|^2 + \frac{1}{n} (\rho_{ii} + \rho_{jj}) + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \\
&\leq |\rho_{ij}|^2 + \frac{\alpha_i + \alpha_j}{n} + \frac{r}{n^2}. \quad (\text{Because } \ell(\boldsymbol{\lambda}) \leq r \text{ when } \rho \text{ is rank } r)
\end{aligned}$$

The first inequality follows from our characterization of Lower_ρ : we can expand $\langle ij | \cdot \text{Lower}_\rho | ji \rangle$ as a positive linear combination of terms of the form

$$\langle ij | \cdot (P \otimes P) \cdot \text{SWAP} | ji \rangle = \langle ij | \cdot (P \otimes P) \cdot | ij \rangle = P_{ii} \cdot P_{jj} \geq 0,$$

where the last step uses the fact that P is positive semidefinite. $\square$

In contrast to learning in trace distance, the Bures distance introduces additional challenges that we have to overcome in the following subsections:

1. The first one is that our debiased Keyl's algorithm can return a matrix $\widehat{\rho}$ with negative eigenvalues. Since the definition of the Bures distance involves the square root of matrices, it is not well-defined for an estimator that is not positive semidefinite.
2. The second challenge arises because we use [Theorem 2.8](#) to upper bound the Bures distance by the Bures χ^2 -divergence, which is sensitive to the small eigenvalues of ρ . Thus, we would like to make sure that the eigenvalues of ρ are large enough to obtain a useful bound on the Bures χ^2 -divergence.

To deal with these challenges that are specific to the Bures distance, we will use two modified estimators in [Sections 5.1](#) and [5.2](#) that go beyond our debiased Keyl's estimator.

5.1 Full-rank tomography

The main result of this section is an algorithm for learning a mixed state ρ to Bures distance ε using $O(d^2/\varepsilon^2)$ copies. We deal with the second challenge of learning in Bures distance, that is, when ρ has some eigenvalues which are too small, by first applying the depolarizing channel to each copy of ρ with noise rate $C \cdot \frac{d^2}{n}$, where C is a constant that we will choose later. That is, we obtain copies of the regularized state

$$\rho' := \left(1 - C \cdot \frac{d^2}{n}\right) \cdot \rho + C \cdot \frac{d^2}{n} \cdot \frac{I}{d}.$$

Note that we will ultimately take $n = \Theta(d^2/\varepsilon^2)$ large enough so that the above procedure is well-defined (i.e., the noise rate does not exceed 1). The depolarizing noise implies that each eigenvalue of ρ' is at least $C \cdot \frac{d^2}{n} \cdot \frac{1}{d} = C \cdot \frac{d}{n}$, and as we will see, our Bures χ^2 -divergence argument does not have to deal with eigenvalues that are too small.

We will show that when the number of samples is $n = \Theta(d^2/\varepsilon^2)$, the original ρ and depolarized state ρ' are $O(\varepsilon)$ -close in Bures distance. Thus, we will use the debiased Keyl's algorithm to produce an estimate $\widehat{\rho}$ that is $O(\varepsilon)$ -close to ρ' in Bures distance, and it will also be a good estimate for the original state ρ .

In addition, we will show that the regularized state also deals with the first challenge of learning in the Bures distance, which is the negative eigenvalues of the estimator. In particular, we show in [Theorem 5.2](#) that there exists a constant C such that the eigenvalues of the estimate $\widehat{\rho}$ are all non-negative with high probability.

Lemma 5.2. *Let ρ' be a diagonal mixed state with sorted eigenvalues $\alpha'_1, \dots, \alpha'_d$, and $\widehat{\rho}$ be the output of the debiased Keyl's algorithm when run on $(\rho')^{\otimes n}$. There exists a constant C such that, if the eigenvalues of ρ' are all at least $C \cdot d/n$, then $\widehat{\rho}$ is a positive semidefinite matrix with high probability.*

Proof. Recall that $\hat{\rho} = \mathbf{U} \cdot \hat{\alpha} \cdot \mathbf{U}^\dagger$, where $\hat{\alpha}$ is obtained by performing the Young diagram box donation procedure on λ and normalizing by n . Then the smallest eigenvalue of $\hat{\rho}$ is lower bounded by

$$\hat{\alpha}_d \geq \frac{\lambda_d - d}{n}. \quad (31)$$

Our goal is to show that $\hat{\alpha}_d$ is non-negative. For this, we will need to understand the distribution of λ when drawn from $\text{SW}^n(\alpha')$. Theorems 1.4 and 1.5 of [OW17] give bounds on the first and second moments of λ_k respectively. In particular, for $k \in [d]$ and $\nu_k := \min\{1, \alpha'_k d\}$, the first moment of λ_k satisfies

$$\alpha'_k n - 2\sqrt{\nu_k n} \leq \mathbf{E}_{\lambda \sim \text{SW}^n(\alpha')}[\lambda_k] \leq \alpha'_k n + 2\sqrt{\nu_k n}, \quad (32)$$

and the second moment of λ_k satisfies

$$\mathbf{E}_{\lambda \sim \text{SW}^n(\alpha')}[(\lambda_k - \alpha'_k n)^2] \leq O(\nu_k n). \quad (33)$$

We can bound the variance of λ_d using [Equations \(32\) and \(33\)](#):

$$\begin{aligned} \text{Var}[\lambda_d] &= \mathbf{E}_{\lambda}[(\lambda_d - \mathbf{E}[\lambda_d])^2] = \mathbf{E}_{\lambda}[(\lambda_d - \alpha'_d n + (\alpha'_d n - \mathbf{E}[\lambda_d]))^2] \\ &\leq \mathbf{E}_{\lambda}[2(\lambda_d - \alpha'_d n)^2 + 2(\alpha'_d n - \mathbf{E}[\lambda_d])^2] \leq O(\nu_d n). \end{aligned} \quad (\text{Cauchy-Schwarz})$$

Now we will use the fact that $\nu_d \leq \alpha'_d d$ to deduce that

$$\mathbf{E}_{\lambda \sim \text{SW}^n(\alpha')}[\lambda_d] \geq \alpha'_d n - 2\sqrt{\alpha'_d d n}, \quad \text{Var}[\lambda_d] \leq O(\alpha'_d d n).$$

So, Chebyshev's inequality says that with high probability, λ_d is within $O(\sqrt{\text{Var}[\lambda_d]}) \leq O(\sqrt{\alpha'_d d n})$ of its expectation. In particular, there exists a constant C' such that, with high probability,

$$\lambda_d \geq \alpha'_d n - C'\sqrt{\alpha'_d d n}.$$

Therefore, from [Equation \(31\)](#), it suffices to show that we can pick our C large enough so that

$$\alpha'_d n - C'\sqrt{\alpha'_d d n} - d \geq 0.$$

From the AM-GM inequality, $C'\sqrt{\alpha'_d d n} \leq \frac{1}{2}\alpha'_d n + \frac{1}{2}(C')^2 d$, and thus it suffices to show that

$$\frac{1}{2}\alpha'_d n - \left(\frac{1}{2}(C')^2 + 1\right)d \geq 0.$$

Setting C to be a constant at least $(C')^2 + 2$, and recalling that $\alpha'_d \geq C \cdot \frac{d}{n}$, we conclude that

$$\frac{1}{2}\alpha'_d n - \left(\frac{1}{2}(C')^2 + 1\right)d \geq \left(\frac{1}{2}C - \frac{1}{2}(C')^2 - 1\right)d \geq 0. \quad \square$$

We are now ready to prove the main result of this subsection: that $\hat{\rho}$ is a good estimator for ρ in Bures distance.

Theorem 5.3 (Bures distance tomography). *Given a mixed state $\rho \in \mathbb{C}^{d \times d}$, $n = O(d^2/\varepsilon^2)$ samples suffice to output an estimator $\hat{\rho}$ such that $D_B(\rho, \hat{\rho}) \leq \varepsilon$ with high probability.*

Proof. Let C be the constant from the statement of [Theorem 5.2](#), and let ρ' be the state we obtain when we apply the depolarizing channel with noise rate $C \cdot d^2/n$ to ρ :

$$\rho' := \left(1 - C \cdot \frac{d^2}{n}\right) \cdot \rho + C \cdot \frac{d^2}{n} \cdot \frac{I}{d}.$$

We will take $n \geq C \cdot d^2/\varepsilon^2$, and thus the noise rate does not exceed 1. We will denote the eigenvalues of ρ' by $\alpha'_1, \dots, \alpha'_d$, and observe that the smallest eigenvalue of ρ' is at least $C \cdot d^2/n \cdot 1/d = C \cdot d/n$. On the other hand, this regularized state ρ' is still $O(\varepsilon)$ -close to ρ in Bures distance, since

$$\frac{1}{2}D_B^2(\rho, \rho') \leq D_{\text{tr}}(\rho, \rho') = \frac{1}{2}\|\rho - \rho'\|_1 \leq \frac{1}{2}\|C \cdot d^2/n \cdot \rho\|_1 + \frac{1}{2}\|C \cdot d^2/n \cdot I/d\|_1 = O(d^2/n) = O(\varepsilon^2).$$

As a result, it suffices to learn ρ' up to $O(\varepsilon)$ error instead of ρ . From [Theorem 5.2](#), we know that running the debiased Keyl's algorithm on $(\rho')^{\otimes n}$ will produce an estimator $\hat{\rho}$ that is positive semidefinite with high probability. Hence the Bures distance between $\hat{\rho}$ and ρ' is well defined, and can be upper bounded by

$$\mathbf{E}[D_B(\hat{\rho}, \rho')] \leq \mathbf{E}\left[\sqrt{D_{\chi^2}(\hat{\rho} \parallel \rho')}\right] \leq \sqrt{\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho')]}. \quad (34)$$

The right-hand side of the above equation becomes:

$$\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho')] = \sum_{i,j=1}^d \frac{2}{\alpha'_i + \alpha'_j} \cdot \mathbf{E}[|\hat{\rho}_{ij} - \rho'_{ij}|^2] = \sum_{i,j=1}^d \frac{2}{\alpha'_i + \alpha'_j} \cdot (\mathbf{E}[|\hat{\rho}_{ij}|^2] - |\rho'_{ij}|^2). \quad (35)$$

Let us consider the term corresponding to a fixed pair (i, j) :

$$\begin{aligned} \frac{2}{\alpha'_i + \alpha'_j} \cdot (\mathbf{E}[|\hat{\rho}_{ij}|^2] - |\rho'_{ij}|^2) &\leq \frac{2}{\alpha'_i + \alpha'_j} \cdot \left(\frac{\alpha'_i + \alpha'_j}{n} + \frac{d}{n^2} \right) && \text{(Theorem 5.1)} \\ &\leq \frac{2}{n} + O\left(\frac{1}{n}\right) = O\left(\frac{1}{n}\right). && (\alpha'_i, \alpha'_j \geq \Omega(d/n) \implies \alpha'_i + \alpha'_j \geq \Omega(d/n)) \end{aligned}$$

[Equation \(35\)](#) has a total of d^2 terms, each of which is at most $O(\frac{1}{n})$ in expectation. We conclude that

$$\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho')] = \sum_{i,j=1}^d \frac{2}{\alpha'_i + \alpha'_j} \cdot (\mathbf{E}[|\hat{\rho}_{ij}|^2] - |\rho'_{ij}|^2) \leq O\left(\frac{d^2}{n}\right).$$

From our choice of $n \geq C \cdot d^2/\varepsilon^2 = \Theta(d^2/\varepsilon^2)$ sufficiently large, it holds that $C \cdot d^2/n = O(\varepsilon^2)$, and $\mathbf{E}[D_{\chi^2}(\hat{\rho} \parallel \rho')] \leq O(\varepsilon^2)$. We conclude from [Equation \(34\)](#) that with high probability:

$$\mathbf{E}[D_B(\rho, \hat{\rho})] \leq D_B(\rho, \rho') + \mathbf{E}[D_B(\rho', \hat{\rho})] \leq O(\varepsilon).$$

The theorem statement follows by reducing ε by a constant factor, and applying Markov's inequality to deduce that $D_B(\rho, \hat{\rho}) \leq \varepsilon$ with high probability. $\square$

5.2 Low-rank tomography

The main result of this section is an algorithm for learning a mixed state of rank r to Bures distance ε using $O(rd/\varepsilon^2)$ copies.

Theorem 5.4. *Given a rank r mixed state $\rho \in \mathbb{C}^{d \times d}$, $n = O(rd/\varepsilon^2)$ samples suffice to output an estimator $\tilde{\rho}$ such that $D_B(\rho, \tilde{\rho}) \leq \varepsilon$ with high probability.*

Recall the first challenge of learning in the Bures distance, which is the fact that $\hat{\rho}$ is not necessarily a positive semidefinite estimator. In the full rank case of the previous subsection, we dealt with this challenge by applying a particular depolarizing channel to each copy of our input state ρ to obtain copies of the regularized state ρ' . We then showed that the debiased Keyl's algorithm on $(\rho')^{\otimes n}$ will output a positive semidefinite estimator $\hat{\rho}$ with high probability, as long as the eigenvalues of ρ' are sufficiently large.

In the low rank case that we consider now, the number of samples does not suffice for this argument to succeed. Instead, we will execute our debiased Keyl's algorithm on the copies of the original state $\rho^{\otimes n}$ to produce $\hat{\rho}$, shift its eigenvalues by r/n , and renormalize:

$$\tilde{\rho} := \frac{n}{n+rd} \left(\hat{\rho} + \frac{r}{n} \cdot I \right).$$

Since the smallest eigenvalue of $\hat{\rho}$ is at least $-\ell(\lambda)/n \geq -r/n$, the resulting state $\tilde{\rho}$ is positive semidefinite and will serve as our low-rank estimator for the Bures distance.

We deal with the second challenge of learning in the Bures distance, which is sensitivity of the Bures χ^2 -divergence on the small eigenvalues of ρ , by following a strategy similar to the one Flammia and O'Donnell use in [FO24, Corollary 3.20] to analyze their Bures distance estimator. In particular, we will bound the Bures χ^2 -divergence between $\tilde{\rho}$ and ρ when restricted to the large eigenvalues of ρ . We then show that the small eigenvalues of ρ do not contribute much more to the Bures distance.

To define the restriction of a state to a subset of eigenvalues, we introduce the following notation, which appears in [FO24, Notation 2.1].

Notation 5.5. Let $\sigma \in \mathbb{C}^{d \times d}$ be a matrix, and $R \subseteq [d]$. We define $\sigma[R] \in \mathbb{C}^{|R| \times |R|}$ to be the submatrix of σ formed by the rows and columns in R .

Note that when σ is a quantum state, $\sigma[R]$ may have trace less than one, but by Cauchy's interlacing theorem below, it still holds that every eigenvalue of $\sigma[R]$ is in $[0, 1]$.

Theorem 5.6 (Cauchy's interlacing theorem, [Fis05]). *Let $A \in \mathbb{C}^{d \times d}$ be a Hermitian matrix with eigenvalues $\lambda_1 \leq \dots \leq \lambda_d$, and R be a subset of $[d]$ of size $d-1$. Then the Hermitian matrix $B = A[R] \in \mathbb{C}^{(d-1) \times (d-1)}$ has eigenvalues $\mu_1 \leq \dots \leq \mu_{d-1}$ that satisfy*

$$\lambda_1 \leq \mu_1 \leq \lambda_2 \leq \dots \leq \lambda_{d-1} \leq \mu_{d-1} \leq \lambda_d.$$

Roughly speaking, we will define $R \subseteq [d]$ to be the subset of indices corresponding to the large eigenvalues of ρ , and thus the expression

$$D_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R]) = \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot |\tilde{\rho}_{ij} - \rho_{ij}|^2$$

will capture the contribution of the large eigenvalues to the Bures χ^2 -divergence. Note that this expression is well-defined as long as $\alpha_i > 0$ for each $i \in R$.

In addition, we use $\sigma^{\text{norm}}[R]$ to denote the normalized version of $\sigma[R]$.

Notation 5.7. Let $\sigma \in \mathbb{C}^{d \times d}$ be a quantum state, and $R \subseteq [d]$. Then $\sigma^{\text{norm}}[R] \in \mathbb{C}^{|R| \times |R|}$ is the $|R|$ -dimensional quantum state given by

$$\sigma^{\text{norm}}[R] := \frac{\sigma[R]}{\text{tr}(\sigma[R])}.$$

Notation 5.8. For a mixed state σ and projector Π , we use

$$\sigma|_{\Pi} := \frac{\Pi \cdot \sigma \cdot \Pi}{\text{tr}(\Pi \cdot \sigma)}$$

to denote the post-measurement state if σ is measured with the projective measurement $\{\Pi, I - \Pi\}$ and the outcome Π is obtained.

Observe that if Π is the projector onto the space defined by the subset R , then the matrices $\sigma^{\text{norm}}[R]$ and $\sigma|_{\Pi}$ correspond to the same quantum state, with the difference that the latter matrix corresponds to the state that lives in the larger ambient space $\mathbb{C}^{d \times d}$.

Proof of Theorem 5.4. To make sure that our constants work out, we will always set $n \geq 400 \cdot rd/\varepsilon^2$. Recall that we can assume without loss of generality that ρ is diagonal with eigenvalues $\alpha_1, \dots, \alpha_r, 0, \dots, 0$. We define R to be the subset of indices i such that $\alpha_i \geq \frac{r}{n}$, and $L := [d] \setminus R$ to be the remaining indices that correspond to eigenvalues less than $\frac{r}{n}$. Given these two subsets, we define the projector Π to project onto the eigenvectors of ρ that lie in R . That is,

$$\Pi = \sum_{i \in R} |i\rangle\langle i|.$$

Since L is the complement of R , the matrix $I - \Pi$ projects onto the eigenvalues that correspond to the subset L .

We will prove the theorem by using the triangle inequality to write

$$D_B(\tilde{\rho}, \rho) \leq D_B(\tilde{\rho}, \tilde{\rho}|_\Pi) + D_B(\tilde{\rho}|_\Pi, \rho|_\Pi) + D_B(\rho|_\Pi, \rho), \quad (36)$$

and bounding each of the three terms by $O(\varepsilon)$ with high probability. We bound the first and third terms using the Gentle Measurement Lemma (Theorem 2.9) on ρ and $\tilde{\rho}$. Recall that the Gentle Measurement Lemma relates the Bures distance between a state σ and the post-measurement state $\sigma|_\Pi$ with the following trace:

$$\left(1 - \frac{1}{2} D_B(\sigma, \sigma|_\Pi)^2\right)^2 = \text{tr}(\Pi \cdot \sigma).$$

For the third term, we note that L is the subset of the eigenvalues of ρ smaller than $\frac{r}{n}$, and since there are at most r nonzero eigenvalues in total, we have

$$\text{tr}(\Pi \cdot \rho) = 1 - \sum_{i \notin R} \alpha_i \geq 1 - \frac{r^2}{n} \geq 1 - \frac{1}{400} \varepsilon^2.$$

For the first term, we bound $\text{tr}(\Pi \cdot \tilde{\rho})$ by computing the expectation of the diagonal entries of $\tilde{\rho}$:

$$\mathbf{E}[\tilde{\rho}_{ii}] = \frac{n}{n+rd} \left(\mathbf{E}[\hat{\rho}_{ii}] + \frac{r}{n} \right) = \frac{n}{n+rd} \left(\alpha_i + \frac{r}{n} \right).$$

This means that

$$\sum_{i \notin R} \mathbf{E}[\tilde{\rho}_{ii}] = \sum_{i \notin R} \frac{n}{n+rd} \left(\alpha_i + \frac{r}{n} \right) \leq \sum_{i \notin R} \frac{n}{n+rd} \cdot \frac{2r}{n} \leq \frac{n}{n+rd} \cdot \frac{2rd}{n} = \frac{2rd}{n+rd} \leq \frac{2rd}{n} \leq \frac{1}{200} \cdot \varepsilon^2.$$

Since $\tilde{\rho}$ is positive semidefinite, the random variable $\text{tr}((I - \Pi) \cdot \tilde{\rho}) = \sum_{i \notin R} \tilde{\rho}_{ii}$ is nonnegative, and Markov's inequality implies that $\text{tr}((I - \Pi) \cdot \tilde{\rho}) \leq \frac{1}{2} \cdot \varepsilon^2$ with probability at least 0.99. This implies that $\text{tr}(\Pi \cdot \tilde{\rho}) \geq 1 - \frac{1}{2} \cdot \varepsilon^2$ with high probability. Therefore, for the rest of this proof, we condition on the event that these inequalities hold with high probability. A straightforward application of the Gentle Measurement Lemma implies that

$$\left(1 - \frac{1}{2} D_B^2(\rho|_\Pi, \rho)\right)^2 = 1 - O(\varepsilon^2) \implies D_B^2(\rho|_\Pi, \rho) \leq O(\varepsilon^2), \quad (37)$$

and

$$\left(1 - \frac{1}{2} D_B^2(\tilde{\rho}, \tilde{\rho}|_\Pi)\right)^2 = 1 - O(\varepsilon^2) \implies D_B^2(\tilde{\rho}, \tilde{\rho}|_\Pi) \leq O(\varepsilon^2). \quad (38)$$

Our argument for bounding the second term of Equation (36) will contain several steps. First, we use the second moment expression of the debiased Keyl's estimator to bound the Bures χ^2 -divergence between the unnormalized states $\tilde{\rho}[R]$ and $\rho[R]$. We then use this to deduce a bound on the Bures χ^2 -divergence between the normalized states $\tilde{\rho}^{\text{norm}}[R]$ and $\rho^{\text{norm}}[R]$, which also implies a bound on their Bures distance, by Theorem 2.8. Finally, we observe that the same bound holds for the Bures distance between $\tilde{\rho}|_\Pi$ and $\rho|_\Pi$, since $(\tilde{\rho}^{\text{norm}}[R], \rho^{\text{norm}}[R])$ and $(\tilde{\rho}|_\Pi, \rho|_\Pi)$ correspond to the same pair of states, with the difference that the first pair lives in the space spanned by R , whereas the second pair lives in the entire ambient space $\mathbb{C}^{d \times d}$.

For the first step, we write

$$\mathbf{E}[D_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] = \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot \mathbf{E}[|\hat{\rho}_{ij} - \rho_{ij}|^2] = \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot (\mathbf{E}[|\hat{\rho}_{ij}|^2] - |\rho_{ij}|^2). \quad (39)$$

Let us consider the term from Equation (39) that corresponds to a fixed pair (i, j) :

$$\begin{aligned} \frac{2}{\alpha_i + \alpha_j} \cdot (\mathbf{E}[|\hat{\rho}_{ij}|^2] - |\rho_{ij}|^2) &\leq \frac{2}{\alpha_i + \alpha_j} \cdot \left(\frac{\alpha_i + \alpha_j}{n} + \frac{r}{n^2} \right) && \text{(Theorem 5.1)} \\ &\leq \frac{2}{n} + \frac{1}{n} = \frac{3}{n}. && (i, j \in R \implies \alpha_i + \alpha_j \geq \frac{2r}{n}) \end{aligned}$$

Equation (39) has a total of $|R|^2 \leq r^2$ terms, each of which is at most $\frac{3}{n}$ in expectation. We deduce that

$$\mathbf{E} [\mathcal{D}_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] = \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot \left(\mathbf{E} [|\hat{\rho}_{ij}|^2] - |\rho_{ij}|^2 \right) \leq \frac{3r^2}{n} \leq O(\varepsilon^2). \quad (40)$$

We now turn our attention to bounding $\mathbf{E} [\mathcal{D}_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R])]$. From our definition of $\tilde{\rho}$, it holds that

$$\tilde{\rho}_{ij} = \begin{cases} \frac{n}{n+rd}(\hat{\rho}_{ij} + \frac{r}{n}) & i = j, \\ \frac{n}{n+rd} \cdot \hat{\rho}_{ij} & i \neq j. \end{cases}$$

Then we write

$$\begin{aligned} \mathbf{E} [\mathcal{D}_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R])] &= \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot \mathbf{E} [|\tilde{\rho}_{ij} - \rho_{ij}|^2] \\ &= \sum_{\substack{i,j \in R \\ i \neq j}} \frac{2}{\alpha_i + \alpha_j} \cdot \mathbf{E} [|\tilde{\rho}_{ij}|^2] + \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\tilde{\rho}_{ii} - \alpha_i|^2] \\ &= \sum_{\substack{i,j \in R \\ i \neq j}} \frac{2}{\alpha_i + \alpha_j} \mathbf{E} \left[\left| \frac{n}{n+rd} \cdot \hat{\rho}_{ij} \right|^2 \right] + \sum_{i \in R} \frac{1}{\alpha_i} \mathbf{E} \left[\left| \frac{n}{n+rd} \left(\hat{\rho}_{ii} + \frac{r}{n} \right) - \alpha_i \right|^2 \right]. \end{aligned} \quad (41)$$

The first term can be bounded using $\mathbf{E} \left[\left| \frac{n}{n+rd} \cdot \hat{\rho}_{ij} \right|^2 \right] \leq \mathbf{E} [|\hat{\rho}_{ij}|^2]$. We bound the second term as follows

$$\begin{aligned} \mathbf{E} \left[\left| \frac{n}{n+rd} \left(\hat{\rho}_{ii} + \frac{r}{n} \right) - \alpha_i \right|^2 \right] &= \mathbf{E} \left[\left| \frac{n}{n+rd} \cdot \hat{\rho}_{ii} + \frac{r}{n+rd} - \alpha_i \right|^2 \right] \\ &= \mathbf{E} \left[\left| \hat{\rho}_{ii} - \alpha_i + \frac{r}{n+rd} - \frac{rd}{n+rd} \cdot \hat{\rho}_{ii} \right|^2 \right] \\ &\leq 3 \mathbf{E} [|\hat{\rho}_{ii} - \alpha_i|^2] + 3 \left(\frac{r}{n+rd} \right)^2 + 3 \mathbf{E} \left[\left| \frac{rd}{n+rd} \cdot \hat{\rho}_{ii} \right|^2 \right] \quad (\text{Cauchy-Schwarz}) \\ &\leq 3 \mathbf{E} [|\hat{\rho}_{ii} - \alpha_i|^2] + 3 \left(\frac{r}{n} \right)^2 + 3 \left(\frac{rd}{n} \right)^2 \cdot \mathbf{E} [|\hat{\rho}_{ii}|^2]. \end{aligned}$$

We plug these bounds back into Equation (41) to obtain

$$\begin{aligned} &\mathbf{E} [\mathcal{D}_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R])] \\ &\leq \sum_{\substack{i,j \in R \\ i \neq j}} \frac{2}{\alpha_i + \alpha_j} \cdot \mathbf{E} [|\hat{\rho}_{ij}|^2] + 3 \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii} - \alpha_i|^2] + 3 \left(\frac{r}{n} \right)^2 \sum_{i \in R} \frac{1}{\alpha_i} + 3 \left(\frac{rd}{n} \right)^2 \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii}|^2] \\ &\leq 3 \mathbf{E} [\mathcal{D}_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + \frac{3r^2}{n} + 3 \left(\frac{rd}{n} \right)^2 \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii}|^2]. \end{aligned}$$

Here, the second inequality follows because all α_i in R are at least $\frac{r}{n}$ and there are at most r of them, so that $\sum_{i \in R} \frac{1}{\alpha_i} \leq n$. Now we bound the last term as follows:

$$\begin{aligned} \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii}|^2] &= \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii} - \alpha_i + \alpha_i|^2] \\ &\leq 2 \sum_{i \in R} \frac{1}{\alpha_i} \cdot \mathbf{E} [|\hat{\rho}_{ii} - \alpha_i|^2] + 2 \sum_{i \in R} \frac{1}{\alpha_i} \cdot \alpha_i^2 \quad (\text{Cauchy-Schwarz}) \\ &\leq 2 \mathbf{E} [\mathcal{D}_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + 2 \sum_{i \in R} \alpha_i \\ &\leq 2 \mathbf{E} [\mathcal{D}_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + 2. \end{aligned}$$

Hence

$$\begin{aligned} \mathbf{E} [D_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R])] &\leq 3 \mathbf{E} [D_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + \frac{3r^2}{n} + 3 \left(\frac{rd}{n} \right)^2 (2 \mathbf{E} [D_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + 2) \\ &\leq O(1) \cdot \mathbf{E} [D_{\chi^2}(\hat{\rho}[R] \parallel \rho[R])] + \frac{3r^2}{n} + \frac{6r^2 d^2}{n^2} \leq O(\varepsilon^2). \end{aligned} \quad (\text{Equation (40)})$$

Therefore, Markov's inequality implies that $D_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R]) \leq O(\varepsilon^2)$ with high probability. For the rest of this proof, we condition on the event that the above inequality holds.

Now, we would like to translate our restricted Bures χ^2 -divergence bound on the (potentially) unnormalized states $\tilde{\rho}[R]$ and $\rho[R]$ to the restricted Bures χ^2 -divergence between their normalized versions $\tilde{\rho}^{\text{norm}}[R]$ and $\rho^{\text{norm}}[R]$. Recall that $\tilde{\rho}^{\text{norm}}[R], \rho^{\text{norm}}[R] \in \mathbb{C}^{|R| \times |R|}$. In particular, for $i, j \in R$, we will need to rescale the entries of the two states as follows

$$(\tilde{\rho}^{\text{norm}}[R])_{ij} = \frac{1}{\text{tr}(\tilde{\rho}[R])} \cdot \tilde{\rho}_{ij} = (1 + c_1) \cdot \tilde{\rho}_{ij}, \quad (\rho^{\text{norm}}[R])_{ij} = \frac{1}{\text{tr}(\rho[R])} \cdot \rho_{ij} = (1 + c_2) \cdot \rho_{ij},$$

where $\text{tr}(\tilde{\rho}[R]) = \text{tr}(\Pi \cdot \tilde{\rho}) \geq 1 - \frac{1}{2} \cdot \varepsilon^2$ implies that $c_1 = O(\varepsilon^2)$, and similarly $c_2 = O(\varepsilon^2)$. Therefore,

$$D_{\chi^2}(\tilde{\rho}^{\text{norm}}[R] \parallel \rho^{\text{norm}}[R]) = \sum_{\substack{i,j \in R \\ i \neq j}} \frac{2}{\alpha_i + \alpha_j} \cdot |(\tilde{\rho}^{\text{norm}}[R])_{ij}|^2 + \sum_{i \in R} \frac{1}{\alpha_i} \cdot |(\tilde{\rho}^{\text{norm}}[R])_{ii} - (\rho^{\text{norm}}[R])_{ii}|^2. \quad (42)$$

We bound the first term using $|(\tilde{\rho}^{\text{norm}}[R])_{ij}|^2 \leq (1 + c_1)^2 \cdot |\tilde{\rho}_{ij}|^2 \leq O(1) \cdot |\tilde{\rho}_{ij}|^2$, and the second term using

$$\begin{aligned} &|(\tilde{\rho}^{\text{norm}}[R])_{ii} - (\rho^{\text{norm}}[R])_{ii}|^2 \\ &\leq 3 |(\tilde{\rho}^{\text{norm}}[R])_{ii} - \tilde{\rho}_{ii}|^2 + 3 |\tilde{\rho}_{ii} - \rho_{ii}|^2 + 3 |\rho_{ii} - (\rho^{\text{norm}}[R])_{ii}|^2 \quad (\text{Cauchy-Schwarz}) \\ &= 3c_1^2 |\tilde{\rho}_{ii}|^2 + 3 |\tilde{\rho}_{ii} - \rho_{ii}|^2 + 3c_2^2 |\rho_{ii}|^2 \\ &\leq 6c_1^2 |\tilde{\rho}_{ii} - \rho_{ii}|^2 + 6c_1^2 |\rho_{ii}|^2 + 3 |\tilde{\rho}_{ii} - \rho_{ii}|^2 + 3c_2^2 |\rho_{ii}|^2 \quad (\text{Cauchy-Schwarz}) \\ &= O(1) \cdot |\tilde{\rho}_{ii} - \rho_{ii}|^2 + O(\varepsilon^4) \cdot |\rho_{ii}|^2 \\ &= O(1) \cdot |\tilde{\rho}_{ii} - \rho_{ii}|^2 + O(\varepsilon^4) \cdot \alpha_i^2. \end{aligned}$$

Thus $D_{\chi^2}(\tilde{\rho}^{\text{norm}}[R] \parallel \rho^{\text{norm}}[R])$ is at most:

$$\begin{aligned} D_{\chi^2}(\tilde{\rho}^{\text{norm}}[R] \parallel \rho^{\text{norm}}[R]) &\leq O(1) \cdot \sum_{\substack{i,j \in R \\ i \neq j}} \frac{2}{\alpha_i + \alpha_j} \cdot |\tilde{\rho}_{ij}|^2 + \sum_{i \in R} \frac{1}{\alpha_i} \cdot (O(1) \cdot |\tilde{\rho}_{ii} - \rho_{ii}|^2 + O(\varepsilon^4) \cdot \alpha_i^2) \\ &= O(1) \cdot \sum_{i,j \in R} \frac{2}{\alpha_i + \alpha_j} \cdot |\tilde{\rho}_{ij} - \rho_{ij}|^2 + O(\varepsilon^4) \cdot \sum_{i \in R} \frac{1}{\alpha_i} \cdot \alpha_i^2 \\ &= O(1) \cdot D_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R]) + O(\varepsilon^4) \cdot \sum_{i \in R} \alpha_i \\ &\leq O(1) \cdot D_{\chi^2}(\tilde{\rho}[R] \parallel \rho[R]) + O(\varepsilon^4) \\ &\leq O(\varepsilon^2). \end{aligned}$$

The inequality between the Bures distance and the Bures χ^2 -divergence implies that

$$D_B(\tilde{\rho}^{\text{norm}}[R], \rho^{\text{norm}}[R]) \leq \sqrt{D_{\chi^2}(\tilde{\rho}^{\text{norm}}[R] \parallel \rho^{\text{norm}}[R])} \leq O(\varepsilon).$$

Moreover, we remark that since Π is the projector onto the eigenvectors that lie in R , $\tilde{\rho}|_{\Pi}$ and $\tilde{\rho}^{\text{norm}}[R]$ correspond to the same state, embedded in two ambient spaces of dimension d and $|R|$ respectively. The same holds for $\rho|_{\Pi}$ and $\rho^{\text{norm}}[R]$. Then

$$F(\tilde{\rho}|_{\Pi}, \rho|_{\Pi}) = F(\tilde{\rho}^{\text{norm}}[R], \rho^{\text{norm}}[R]) \implies D_B(\tilde{\rho}|_{\Pi}, \rho|_{\Pi}) = D_B(\tilde{\rho}^{\text{norm}}[R], \rho^{\text{norm}}[R]) \leq O(\varepsilon).$$

Recall that our goal is to use the triangle inequality $D_B(\tilde{\rho}, \rho) \leq D_B(\tilde{\rho}, \tilde{\rho}|_\Pi) + D_B(\tilde{\rho}|_\Pi, \rho|_\Pi) + D_B(\rho|_\Pi, \rho)$. We have showed that each of the terms on the right-hand side is $O(\varepsilon)$ with high probability. Therefore, we conclude that with high probability,

$$D_B(\tilde{\rho}, \rho) \leq O(\varepsilon).$$

The theorem statement follows by reducing ε by a small constant factor. $\square$

5.3 Learning bipartite pure states

As an application of our tomography algorithm for Bures distance, we consider the problem of learning bipartite pure states.

Suppose $|\psi\rangle_{AB}$ is a bipartite state in a Hilbert space $\mathcal{H}_A \otimes \mathcal{H}_B$ with $\dim(\mathcal{H}_A) = d_A$ and $\dim(\mathcal{H}_B) = d_B$, and suppose $|\psi\rangle_{AB}$ has Schmidt rank r , i.e. there exist vectors $\{|u_i\rangle_A\}_{i \in [r]}$ in $\mathcal{H}_A$ and $\{|v_i\rangle_B\}_{i \in [r]}$ in $\mathcal{H}_B$, and nonnegative coefficients $\{\lambda_i\}_{i \in [r]}$ such that

$$|\psi\rangle_{AB} = \sum_{i=1}^r \lambda_i \cdot |u_i\rangle_A \otimes |v_i\rangle_B.$$

What is the sample complexity of learning $|\psi\rangle_{AB}$ to ε error in Bures distance? One strategy is to use the rank-1 special case of the Bures learning algorithm directly, [Theorem 5.4](#), which uses $O(d_A d_B / \varepsilon^2)$ samples. In this subsection, we give an algorithm with better sample complexity when $r \ll d_A, d_B$. This algorithm was inspired by the reduction from rank- r state tomography to pure state tomography from [\[Yue23\]](#).

Definition 5.9 (Algorithm for learning bipartite pure states). On input $|\psi\rangle_{AB}^{\otimes n}$, with $n = O(r(d_A + d_B)/\varepsilon^2)$:

1. Trace out the B register on $O(rd_A/\varepsilon^2)$ copies of $|\psi\rangle_{AB}$, obtaining copies of the rank- r mixed state ψ_A .
2. Perform rank- r tomography on the copies of ψ_A , using [Theorem 5.4](#), producing an estimate $\hat{\psi}_A$ such that $D_B(\hat{\psi}_A, \psi_A) \leq \varepsilon$ with high probability. Write Π_A for the projector onto the support of $\hat{\psi}_A$. By [Theorem 2.10](#), we can assume $\hat{\psi}_A$ is a mixed state of rank at most r .
3. Measure $O(rd_B/\varepsilon^2)$ copies of $|\psi\rangle_{AB}$ with $\{\Pi_A \otimes I_B, \bar{\Pi}_A \otimes I_B\}$. Write

$$|\pi\rangle_{AB} = \frac{\Pi_A \otimes I_B |\psi\rangle_{AB}}{\|(\Pi_A \otimes I_B) |\psi\rangle_{AB}\|}$$

for the pure state collapsed to upon obtaining the measurement outcome $\Pi_A \otimes I_B$.

4. Since $|\pi\rangle_{AB}$ is a pure state in $\text{supp}(\Pi_A) \otimes \mathcal{H}_B$, which has dimension at most rd_B , use the copies of $|\pi\rangle_{AB}$ as input to the rank-1 special case of [Theorem 5.4](#). We obtain a state $\hat{\pi}_{AB}$ such that $D_B(\hat{\pi}_{AB}, \pi_{AB}) \leq \varepsilon$ with high probability.
5. Output $\hat{\pi}_{AB}$.

Remark 5.10. The algorithm in [Theorem 5.9](#) does not necessarily output a pure state. However, by [Theorem 2.10](#), there exists an algorithm that learns a pure state $|\hat{\psi}\rangle_{AB}$ approximating $|\psi\rangle_{AB}$, using the same number of samples. If the base algorithm learns to Bures distance ε , then the modified algorithm learns to Bures distance 2ε .

Proposition 5.11. *The algorithm given in [Theorem 5.9](#) learns $|\psi\rangle_{AB}$ to Bures distance $O(\varepsilon)$ using $n = O(r(d_A + d_B)/\varepsilon^2)$ copies. As a consequence, if we modify the algorithm to output pure states as in [Theorem 5.10](#), then it outputs a pure state $|\hat{\psi}\rangle_{AB}$ with overlap $|\langle \psi | \hat{\psi} \rangle_{AB}| \geq 1 - O(\varepsilon^2)$.*

Proof. First we prove that if all of the steps succeed, then $D_B(\hat{\pi}_{AB}, \psi_{AB}) \leq 2\varepsilon$. Since $D_B(\hat{\psi}_A, \psi_A) \leq \varepsilon$, we have $F(\hat{\psi}_A, \psi_A) \geq 1 - \varepsilon^2/2$. Thus, by [\[Wat18, Proposition 3.12, \(4\) and \(6\)\]](#),

$$\begin{aligned} 1 - \frac{1}{2}\varepsilon^2 &\leq F(\hat{\psi}_A, \psi_A) = F(\hat{\psi}_A, \Pi_A \psi_A \Pi_A) \leq \sqrt{\text{tr}(\hat{\psi}_A) \cdot \text{tr}(\Pi_A \psi_A \Pi_A)} \\ &= \sqrt{\text{tr}(\Pi_A \psi_A)} = \sqrt{\text{tr}((\Pi_A \otimes I_B) |\psi\rangle\langle\psi|_{AB})} = \|(\Pi_A \otimes I_B) |\psi\rangle_{AB}\|. \end{aligned}$$

Therefore,

$$|\langle \pi | \psi \rangle_{AB}| = \frac{\langle \psi | (\Pi_A \otimes I_B) | \psi \rangle_{AB}}{\|(\Pi_A \otimes I_B) | \psi \rangle_{AB}\|} = \|(\Pi_A \otimes I_B) | \psi \rangle_{AB}\| \geq 1 - \frac{1}{2}\varepsilon^2.$$

This implies

$$D_B(\pi_{AB}, \psi_{AB}) = \sqrt{2(1 - |\langle \pi | \psi \rangle_{AB}|)} \leq \varepsilon.$$

So, by the triangle inequality, we have

$$D_B(\hat{\pi}_{AB}, \psi_{AB}) \leq D_B(\hat{\pi}_{AB}, \pi_{AB}) + D_B(\pi_{AB}, \psi_{AB}) \leq \varepsilon + \varepsilon = 2\varepsilon.$$

Now we consider the success probabilities. First note that the second and fourth steps succeed with high probability provided they are given the number of samples specified, by [Theorem 5.4](#). The only other step to consider is the third step. In this stage of the algorithm, conditioned on the success of the second step, we obtain outcome $\Pi_A \otimes I_B$ with probability

$$\|(\Pi_A \otimes I_B) | \psi \rangle_{AB}\|^2 \geq 1 - \frac{1}{2}\varepsilon^2.$$

We can assume $\varepsilon \leq 1$, since otherwise, we can learn to this higher, still constant, accuracy instead. The probability $\|(\Pi_A \otimes I_B) | \psi \rangle_{AB}\| \geq 1 - \frac{1}{2}\varepsilon^2$ is at least $\frac{1}{2}$. Thus, with high probability, $O(rd_B/\varepsilon^2)$ copies of $|\psi\rangle_{AB}$ suffice to produce the necessary $O(rd_B/\varepsilon^2)$ samples of $|\pi\rangle_{AB}$ which are required for the fourth step. Altogether, since each step succeeds with high probability, the overall algorithm does too. $\square$

6 Tomography with limited entanglement

In this section, we use our unbiased estimator to devise a tomography algorithm when we are limited to measurements spanning up to k copies of the state ρ .

Theorem 6.1. *Let $n = n' \cdot k$. Let $\hat{\rho}_1, \dots, \hat{\rho}_{n'}$ be the outputs of the debiased Keyl's algorithm when run in n' independent copies of $\rho^{\otimes k}$. Set $\hat{\rho} = (\hat{\rho}_1 + \dots + \hat{\rho}_{n'})/n'$. Then $D_{\text{tr}}(\rho, \hat{\rho}) \leq \varepsilon$ with high probability when*

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

Proof. From linearity of expectation and the unbiasedness of our tomography algorithm, we know that

$$\mathbf{E} \hat{\rho} = \rho.$$

We proceed to bound the variance of our estimator. Since $\hat{\rho}_1, \dots, \hat{\rho}_{n'}$ are independent and identically distributed,

$$\text{Var}[\hat{\rho}] = \frac{1}{(n')^2} \cdot \text{Var}[\hat{\rho}_1 + \dots + \hat{\rho}_{n'}] = \frac{1}{(n')^2} \cdot (\text{Var}[\hat{\rho}_1] + \dots + \text{Var}[\hat{\rho}_1]) = \frac{1}{n'} \cdot \text{Var}[\hat{\rho}_1].$$

Recall from [Theorem 4.2](#) that

$$\text{Var}[\hat{\rho}_1] \leq \frac{2d}{k} + \frac{d^2 \cdot \mathbf{E} \ell(\lambda)}{k^2}.$$

We show in [Theorem 6.2](#) that $\mathbf{E} \ell(\lambda) \leq 2\sqrt{k}$, which means that

$$\text{Var}[\hat{\rho}] = \frac{1}{n'} \cdot \text{Var}[\hat{\rho}_1] \leq \frac{2d}{n'k} + \frac{2d^2}{n'k\sqrt{k}} = \frac{2d}{n} + \frac{2d^2}{n\sqrt{k}}.$$

We upper bound the trace distance using

$$\mathbf{E} \|\hat{\rho} - \rho\|_1 \leq \sqrt{d} \sqrt{\mathbf{E} \|\hat{\rho} - \rho\|_2^2} \leq \sqrt{\frac{2d^2}{n} + \frac{2d^3}{n\sqrt{k}}}.$$

Therefore, taking

$$n \geq \frac{4d^2}{\varepsilon^2} + \frac{4d^3}{\sqrt{k}\varepsilon^2} = O\left(\max\left(\frac{d^3}{\sqrt{k}\varepsilon^2}, \frac{d^2}{\varepsilon^2}\right)\right)$$

many samples suffices to produce $\hat{\rho}$ with $\mathbf{E} \|\hat{\rho} - \rho\|_1 \leq \varepsilon$. Markov's inequality then implies that $\|\hat{\rho} - \rho\|_1 \leq O(\varepsilon)$ with high probability. The theorem statement follows by reducing ε by a small constant factor. $\square$

Alternative proof using Young diagram statistics. Similar to [Theorem 4.1](#), there is an alternative way of proving [Theorem 6.1](#) without relying on our second moment calculation. Let $\hat{\rho}' = \mathbf{U} \cdot \hat{\alpha}' \cdot \mathbf{U}^\dagger$ be the member of the family of unbiased estimators that uses the rank- $\ell(\lambda)$ box donation transformation, where $\ell(\lambda) := \ell(\lambda)$ is the height of the Young tableau we obtain from weak Schur sampling. The spectrum of $\hat{\rho}'$ satisfies

$$\hat{\alpha}'_i = \begin{cases} \frac{\lambda_i + d - 2i + 1}{n} & i \leq \ell(\lambda) \\ -\frac{\ell(\lambda)}{n} & i \in \{\ell(\lambda) + 1, \dots, d\}. \end{cases}$$

We will obtain a slightly weaker bound for the variance of $\hat{\rho}$ as follows: $\text{Var}[\hat{\rho}_1] \leq \text{Var}[\hat{\rho}'_1] \leq \frac{23d}{k} + \frac{2d^2}{k\sqrt{k}}$. We proved the first inequality in [Theorem 3.9](#), and for the second inequality, we write

$$\begin{aligned} \text{Var}[\hat{\rho}'] &= \mathbf{E} \text{tr}(\hat{\rho}'^2) - p_2(\alpha) & (\text{Equation (25)}) \\ &= \mathbf{E} \left[\sum_{i=1}^d \hat{\alpha}'_i{}^2 \right] - p_2(\alpha) \\ &= \mathbf{E} \left[\sum_{i=1}^{\ell(\lambda)} \frac{(\lambda_i + (d - 2i + 1))^2}{k^2} \right] + \mathbf{E} \sum_{i=\ell(\lambda)+1}^d \frac{(-\ell(\lambda))^2}{k^2} - p_2(\alpha). \end{aligned} \quad (43)$$

We will bound the first term again using Kerov's algebra of observables of diagrams, and in particular, the expectation of $p_2^*(\lambda)$ from [Equation \(30\)](#).

$$\begin{aligned} (43) &\leq \frac{1}{k^2} \mathbf{E} \left[\sum_{i=1}^d (\lambda_i^2 - \lambda_i(2i - 1)) + \sum_{i=1}^{\ell(\lambda)} \lambda_i(2d - 2i + 1) + \sum_{i=1}^{\ell(\lambda)} (d - 2i + 1)^2 \right] + \frac{d \cdot \mathbf{E} \ell(\lambda)^2}{k^2} - p_2(\alpha) \\ &\leq \frac{1}{k^2} \mathbf{E}[p_2^*(\lambda)] + \frac{2d}{k^2} \mathbf{E} \left[\sum_{i=1}^{\ell(\lambda)} \lambda_i \right] + \frac{d^2 \cdot \mathbf{E} \ell(\lambda)}{k^2} + \frac{d \cdot \mathbf{E} \ell(\lambda)^2}{k^2} - p_2(\alpha) \\ &\leq \frac{2d}{k} + \frac{d^2 \cdot \mathbf{E} \ell(\lambda)}{k^2} + \frac{d \cdot \mathbf{E} \ell(\lambda)^2}{k^2}. \end{aligned} \quad (\text{Equation (30)})$$

We conclude using [Theorem 6.2](#) that $\text{Var}[\hat{\rho}'_1] \leq \frac{23d}{k} + \frac{2d^2}{k\sqrt{k}}$. Given this variance bound, the statement of [Theorem 6.1](#) as in the original proof.

Lemma 6.2. *Let β be a distribution over $[d]$, and let λ be the Young diagram produced from the RSK algorithm with k symbols drawn from the distribution β . Then the first and second moment of the height of the tableau λ satisfy*

$$\mathbf{E}_{\lambda \sim \text{SW}_k^d(\beta)} \ell(\lambda) \leq 2\sqrt{k}, \quad \mathbf{E}_{\lambda \sim \text{SW}_k^d(\beta)} \ell(\lambda)^2 \leq 20k.$$

Proof. Theorem 1.5.4 of [\[Wri16\]](#) states that if β, v are two distributions such that $\beta \succ v$, then there is a coupling (λ, μ) of $\text{SW}(\beta)$ and $\text{SW}(v)$ such that $\lambda \supseteq \mu$. Here, $\beta \succ v$ means that $\sum_{i=1}^k \beta_i \geq \sum_{i=1}^k v_i$ for all $k \in [d]$, and for two partitions λ and μ , we have $\lambda \supseteq \mu$ means, similarly, that $\sum_{i=1}^k \lambda_i \geq \sum_{i=1}^k \mu_i$ for all $k \in [d]$.

It is not hard to see that if $\lambda \supseteq \mu$, then μ necessarily has at least as many rows as λ . In the opposite case, $\sum_{i=1}^{\ell(\mu)} \mu_i > \sum_{i=1}^{\ell(\mu)} \lambda_i$, thus contradicting the monotonicity. We will choose $v = (1/d, \dots, 1/d)$ and deduce that

$$\mathbf{E}_{\lambda \sim \text{SW}_k^d(\beta)} \ell(\lambda)^t \leq \mathbf{E}_{\mu \sim \text{SW}_k^d(v)} \ell(\mu)^t$$

for all moments $t \in \mathbb{N}$.

What is nice is that the distribution of $\text{SW}_k^d(v)$ is the limit as $d \rightarrow \infty$ is well-studied and also known as the Plancherel distribution Planch_k ([\[Wri16\]](#), Corollary 3.3.4). To make use of this limiting behavior, define β_D for $D \geq d$ as the inclusion of β as a probability distribution on $[D]$, i.e. $\beta_D = (\beta_1, \dots, \beta_d, 0, \dots, 0)$. Moreover, let v_D be the uniform distribution on $[D]$.

Then for all $D \geq d$ we have

$$\mathbf{E}_{\lambda \sim \text{SW}_k^d(\beta)} \ell(\lambda)^t = \mathbf{E}_{\lambda \sim \text{SW}_k^D(\beta_D)} \ell(\lambda)^t \leq \mathbf{E}_{\mu \sim \text{SW}_k^D(v_D)} \ell(\mu)^t,$$

where we've used that $\text{SW}_k^d(\beta)$ is identical to $\text{SW}_k^D(\beta_D)$. And because this holds for arbitrary $D \geq d$, we can take the limit as $D \rightarrow \infty$ to get

$$\mathbf{E}_{\lambda \sim \text{SW}_k^d(\beta)} \ell(\lambda)^t \leq \mathbf{E}_{\lambda \sim \text{Planch}_k} \ell(\lambda)^t.$$

For the Plancherel distribution, it holds that $\ell(\lambda)$ and λ_1 are identically distributed. This is because drawing λ from Planch_k is equivalent to sampling a uniformly random permutation π of $[k]$ and performing the RSK algorithm to produce a diagram. The reverse permutation π_{rev} will produce the transpose diagram λ' when we execute the RSK algorithm, which satisfies $\ell(\lambda') = \lambda_1$. Since the distribution of π and π_{rev} are equivalent, we conclude that for all moments $t \in \mathbb{N}$

$$\mathbf{E}_{\lambda \sim \text{Planch}_k} \ell(\lambda)^t = \mathbf{E}_{\lambda \sim \text{Planch}_k} \lambda_1^t.$$

We are now ready to bound the RHS for the first and second moments. When $t = 1$, [VK85] and [Pil90] showed that (see [Wri16, Theorem 3.5.4] for an exposition of the proof)

$$\mathbf{E}_{\lambda \sim \text{Planch}_k} \lambda_1 \leq 2\sqrt{k}. \quad (44)$$

When $t = 2$, we write

$$\mathbf{E}_{\lambda \sim \text{Planch}_k} \lambda_1^2 = \left(\mathbf{E}_{\lambda \sim \text{Planch}_k} \lambda_1 \right)^2 + \text{Var}[\lambda_1] \leq 4k + 16k = 20k,$$

where the inequality follows from Equation (44) and [OW17, Proposition 4.8]. $\square$

7 Shadow tomography

In this section, we prove the following theorem.

Theorem 7.1 (Theorem 1.9, restated). *The classical shadows task for states of rank r and observables $O_1, \dots, O_m$ with $\text{tr}(O_i^2) \leq F$ can be solved using*

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

copies of ρ .

We use a “plug-in estimator” approach based on the debiased Keyl’s algorithm, and show that it achieves the above sample complexity.

Definition 7.2 (Algorithm for classical shadows using the debiased Keyl’s algorithm). Let n' and k be parameters, and set $n = n' \cdot k$. On input $\rho^{\otimes n}$, with $n = n' \cdot k$:

1. Use n' samples of ρ to obtain an estimate $\hat{\rho}_1$ for ρ using the debiased Keyl’s algorithm.
2. Set $\hat{o}_i^{(1)} := \text{tr}(O_i \cdot \hat{\rho}_1)$, for each $i \in \{1, \dots, m\}$.
3. Repeat the first two steps $(k - 1)$ times to obtain estimates $\hat{o}_i^{(j)}$, for $j \in \{2, \dots, k\}$ and $i \in \{1, \dots, m\}$.
4. For each i , output $\hat{o}_i := \text{median}(\hat{o}_i^{(1)}, \dots, \hat{o}_i^{(k)})$.

Specifically, we will show below that this algorithm achieves the sample complexity in Theorem 7.1 with $k = O(\log m)$ and $n' = O(\min\{\frac{\sqrt{rF}}{\varepsilon}, \frac{F^{2/3}}{\varepsilon^{4/3}}\} + \frac{1}{\varepsilon^2})$. We begin our analysis with a lemma.

Lemma 7.3. *Let $O \in \mathbb{C}^{d \times d}$ be Hermitian. If $\hat{\rho}$ is the output of the debiased Keyl's algorithm on input $\rho^{\otimes n}$, with ρ rank- r , then*

$$\mathbf{Var}[\mathrm{tr}(O \cdot \hat{\rho})] \leq \frac{2}{n} \|O\|_\infty^2 + \frac{\min\{r, 2\sqrt{n}\}}{n^2} \mathrm{tr}(O^2).$$

Proof. We have

$$\mathbf{Var}[\mathrm{tr}(O \cdot \hat{\rho})] = \mathbf{E} \mathrm{tr}(O \cdot \hat{\rho})^2 - (\mathbf{E} \mathrm{tr}(O \cdot \hat{\rho}))^2 = \mathrm{tr}(O \otimes O \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}]) - \mathrm{tr}(O \cdot \rho)^2. \quad (45)$$

By [Theorem 1.4](#),

$$\begin{aligned} \mathrm{tr}(O \otimes O \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}]) &= \frac{n-1}{n} \mathrm{tr}(O \otimes O \cdot \rho \otimes \rho) + \frac{1}{n} \mathrm{tr}(O \otimes O \cdot \rho \otimes I \cdot \mathrm{SWAP}) + \frac{1}{n} \mathrm{tr}(O \otimes O \cdot I \otimes \rho \cdot \mathrm{SWAP}) \\ &\quad + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \mathrm{tr}(O \otimes O \cdot \mathrm{SWAP}) - \mathrm{tr}(O \otimes O \cdot \mathrm{Lower}_\rho) \\ &= \frac{n-1}{n} \mathrm{tr}(O \cdot \rho)^2 + \frac{2}{n} \mathrm{tr}(O^2 \cdot \rho) + \frac{\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \mathrm{tr}(O^2) - \mathrm{tr}(O \otimes O \cdot \mathrm{Lower}_\rho). \end{aligned}$$

First, we note that the last term is non-negative from our characterization of Lower_ρ . In particular, we can expand $\mathrm{tr}(O \otimes O \cdot \mathrm{Lower}_\rho)$ as a positive linear combination of terms of the form

$$\mathrm{tr}(O \otimes O \cdot P \otimes P \cdot \mathrm{SWAP}) = \mathrm{tr}(O P O P) = \mathrm{tr}((\sqrt{P} O \sqrt{P})^2) \geq 0,$$

where the last equality follows because P is positive semidefinite. Next, because $O^2 \preceq \|O\|_\infty^2 \cdot I_d = \|O\|_\infty^2 \cdot I_d$ in the PSD order, we have $\mathrm{tr}(O^2 \cdot \rho) \leq \|O\|_\infty^2 \cdot \mathrm{tr}(\rho) = \|O\|_\infty^2$. Finally, we note two upper bounds on $\mathbf{E}[\ell(\boldsymbol{\lambda})]$: on the one hand, we always have $\ell(\boldsymbol{\lambda}) \leq r$, so that $\mathbf{E}[\ell(\boldsymbol{\lambda})] \leq r$; on the other hand, $\mathbf{E}[\ell(\boldsymbol{\lambda})] \leq 2\sqrt{n}$, by [Theorem 6.2](#). Combining all of these observations gets us:

$$\mathrm{tr}(O \otimes O \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho}]) \leq \mathrm{tr}(O \cdot \rho)^2 + \frac{2}{n} \|O\|_\infty^2 + \frac{\min\{r, 2\sqrt{n}\}}{n^2} \mathrm{tr}(O^2).$$

Substituting this back into [Equation \(45\)](#) gives us the desired bound. $\square$

Proof of Theorem 7.1. Fix any $i \in [m]$, and any $j \in [k]$. First note that the estimate $\hat{\mathbf{o}}_i^{(j)}$ is unbiased, since by [Theorem 1.3](#), we have

$$\mathbf{E}[\hat{\mathbf{o}}_i^{(j)}] = \mathrm{tr}(O_i \cdot \mathbf{E}[\hat{\rho}]) = \mathrm{tr}(O_i \cdot \rho).$$

Moreover, [Theorem 7.3](#) tells us

$$\mathbf{Var}[\hat{\mathbf{o}}_i^{(j)}] \leq \frac{2}{n'} + \frac{2}{(n')^2} \cdot \min\{r, \sqrt{n'}\} \cdot F. \quad (46)$$

Taking $n' = O(1/\varepsilon^2)$ suffices to make the first term on the right-hand side of the above equation at most $O(\varepsilon^2)$. Note that taking $n' \geq \sqrt{rF}/\varepsilon$ suffices to make $2rF/(n')^2 \leq O(\varepsilon^2)$, and taking $n' \geq F^{2/3}/\varepsilon^{4/3}$ suffices to make $2F/(n')^{3/2} \leq O(\varepsilon^2)$. Thus, taking $n' = O(\min\{\sqrt{rF}/\varepsilon, F^{2/3}/\varepsilon^{4/3}\})$ suffices to make at least one of the terms in the minimum $O(\varepsilon^2)$, the second term $O(\varepsilon^2)$.

So, we have that

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

suffices to ensure $\mathbf{Var}[\hat{\mathbf{o}}_i^{(j)}] \leq O(\varepsilon^2)$ for all $i \in [m]$. By Chebyshev's inequality, this sample complexity is also enough to ensure that $|\hat{\mathbf{o}}_i^{(j)} - \mathrm{tr}(O_i \rho)| \leq \varepsilon$ with high probability, for any i , say 99%, by taking C sufficiently small.

We complete our proof with a standard median argument. Assume k is odd. Fix i and define a set $\mathcal{S}_i$ by

$$\mathcal{S}_i := \left| \left\{ j : |\hat{\mathbf{o}}_i^{(j)} - \mathrm{tr}(O_i \rho)| > \varepsilon \right\} \right|.$$

Then the median of $\widehat{\mathbf{o}}_i^{(1)}, \dots, \widehat{\mathbf{o}}_i^{(k)}$ is more than ε away from $\text{tr}(O_i \rho)$ only when $\mathbf{S}_i > k/2$. However, $\mathbf{E}[\mathbf{S}_i] \leq k/100$, so that

$$\begin{aligned} \Pr \left[\left| \text{median}(\mathbf{o}_i^{(j)})_{j \in [k]} - \text{tr}(O_i \cdot \rho) \right| > \varepsilon \right] &\leq \Pr [\mathbf{S}_i > k/2] \\ &\leq \Pr [\mathbf{S}_i > \mathbf{E}[\mathbf{S}_i] + 49k/100] \leq \exp(-2k \cdot (49/100)^2). \end{aligned}$$

Here we have used an additive Chernoff bound, and we conclude that

$$\Pr \left[\left| \text{median}(\mathbf{o}_i^{(1)}, \dots, \mathbf{o}_i^{(k)}) - \text{tr}(O_i \cdot \rho) \right| > \varepsilon \right] \leq \exp(-O(k)).$$

Taking $k = O(\log m)$ is enough to make this probability less than $1/(100m)$. So, by a union bound, the probability that *all* m of the $\widehat{\mathbf{o}}_i$ are within ε of the correct value is at most $1/100$.

We conclude that

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

samples suffices to solve classical shadows using the algorithm described in [Theorem 7.2](#). $\square$

8 Quantum metrology

In this section, we construct a locally unbiased estimator achieving twice the QCRB asymptotically. Specifically, we show the following theorem, stated in the introduction.

Theorem 8.1 ([Theorem 1.10](#), restated). *Let ρ_θ be a parameterized state family with probe state*

$$\rho_{\theta^*} = \sum_{i=1}^d \alpha_i \cdot |v_i\rangle\langle v_i|,$$

where we assume without loss of generality that the eigenvalues are sorted so that $\alpha_1 \geq \dots \geq \alpha_r > 0$ and $\alpha_{r+1} = \dots = \alpha_d = 0$ (so that ρ_{θ^*} is rank r). Then for every n , there is a locally unbiased estimator which takes as input n copies of ρ_θ and whose MSEM V_n has the following property. When $n = 2r/(\varepsilon \cdot \alpha_r)$,

$$V_n \preceq (1 + \varepsilon) \cdot \frac{2}{n \cdot \mathcal{F}}.$$

Hence, as $n \rightarrow \infty$, we have $n \cdot V_n \rightarrow 2/\mathcal{F}$, and so this estimator achieves twice the QCRB asymptotically.

In [\[ZC25\]](#), Zhou and Chen design a locally unbiased estimator for θ , given an unbiased estimator for ρ_θ . Since we prove the above theorem by using the debiased Keyl's algorithm as the basis for such a locally unbiased estimator, we will now give an overview of their construction. We start by reviewing some standard material. Throughout, we will write ∂_i as shorthand for $\partial/\partial\theta_i$. We will also assume without loss of generality that ρ_{θ^*} is diagonalized in the computational basis, with eigenvalues sorted decreasingly, i.e. that $|v_i\rangle = |i\rangle$.

Recall that the symmetric logarithm derivative matrices $\{L_i\}_{i \in [m]}$ are defined implicitly as solutions to the equations

$$\partial_i \rho_\theta|_{\theta=\theta^*} = \frac{1}{2} L_i \rho_{\theta^*} + \frac{1}{2} \rho_{\theta^*} L_i. \quad (47)$$

These equations do not uniquely determine $\{L_i\}$, since we can add to L_i any matrix ΔL_i supported on the kernel of ρ_{θ^*} without affecting the right-hand side of [Equation \(47\)](#). However, this is the only source of non-uniqueness, since

$$(\partial_i \rho_\theta)_{k\ell}|_{\theta=\theta^*} = \langle k | \left(\frac{1}{2} L_i \rho_{\theta^*} + \frac{1}{2} \rho_{\theta^*} L_i \right) | \ell \rangle = \frac{1}{2} (\alpha_k + \alpha_\ell) (L_i)_{k\ell}$$

implies $(L_i)_{k\ell} = 2(\partial_i \rho_\theta)_{k\ell}|_{\theta=\theta^*}/(\alpha_k + \alpha_\ell)$ whenever $\alpha_k + \alpha_\ell > 0$, i.e. whenever $\alpha_k \neq 0$ or $\alpha_\ell \neq 0$. Taking L_i to be zero on the kernel of ρ_{θ^*} then gives rise to a canonical choice of the matrices $\{L_i\}$, with

$$L_i = 2 \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{(\partial_i \rho_\theta)_{k\ell}|_{\theta=\theta^*}}{\alpha_k + \alpha_\ell} \cdot |k\rangle\langle\ell|. \quad (48)$$

Note that this choice is Hermitian, since $\partial_i \rho_\theta|_{\theta=\theta^*}$ is Hermitian. Note also that for any choice of $\{L_i\}$,

$$\text{tr}(L_i \rho_{\theta^*}) = \frac{1}{2} \text{tr}(L_i \rho_{\theta^*}) + \frac{1}{2} \text{tr}(\rho_{\theta^*} L_i) = \text{tr}(\partial_i \rho_\theta|_{\theta=\theta^*}) = \partial_i \text{tr}(\rho_\theta)|_{\theta=\theta^*} = 0, \quad (49)$$

since $\text{tr}(\rho_\theta) = 1$ for any θ .

Recall that the QFI matrix is defined as

$$\mathcal{F}_{ij} := \frac{1}{2} \text{tr}(\rho_{\theta^*} L_i L_j) + \frac{1}{2} \text{tr}(\rho_{\theta^*} L_j L_i) = \text{tr}(L_i \cdot \partial_j \rho_\theta)|_{\theta=\theta^*}. \quad (50)$$

This definition is independent of the choice of $\{L_i\}$, since if $L'_i = L_i + \Delta L_i$, and ΔL_i is supported on the kernel of ρ_{θ^*} , then $\rho_{\theta^*} \Delta L_i = \Delta L_i \rho_{\theta^*} = 0$, so that $\text{tr}(\rho_{\theta^*} L_i L_j) = \text{tr}(\rho_{\theta^*} L'_i L_j) = \text{tr}(\rho_{\theta^*} L'_i L'_j)$, for all $i, j \in [m]$.

Note that $\mathcal{F}_{ij}$ is symmetric. Moreover, since we can choose Hermitian $\{L_i\}$, $\mathcal{F}_{ij}$ is also real, as

$$\mathcal{F}_{ij}^* = \frac{1}{2} \text{tr}(\rho_{\theta^*} L_j^\dagger L_i^\dagger) + \frac{1}{2} \text{tr}(\rho_{\theta^*} L_i^\dagger L_j^\dagger) = \mathcal{F}_{ij}.$$

We can also give an explicit expression for the entries of $\mathcal{F}$, using the matrices $\{L_i\}$ given in Equation (48), and the expression for $\mathcal{F}_{ij}$ on the right-hand side of Equation (50):

$$\mathcal{F}_{ij} = 2 \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{((\partial_i \rho_\theta)_{k\ell} \cdot (\partial_j \rho_\theta)_{\ell k})|_{\theta=\theta^*}}{\alpha_k + \alpha_\ell}. \quad (51)$$

It is assumed that $\mathcal{F}$ is invertible. We briefly explain why. If $\mathcal{F}$ is not invertible, it must have a real eigenvector which has eigenvalue zero. Let this vector be $|\psi\rangle = \sum_{i=1}^d z_i |i\rangle$, with $z_i \in \mathbb{R}$. Then $\langle\psi| \mathcal{F} |\psi\rangle = 0 = \text{tr}(\rho_{\theta^*} Q^2)$, where $Q = \sum_{i=1}^m z_i L_i$. Since Q is Hermitian, this implies that Q is supported on the kernel of ρ_{θ^*} . But since each L_i is zero on this kernel, this means that Q is also zero on the kernel, and hence $Q = 0$ as an operator. Thus, there exists a linear combination of derivatives which is zero at θ^* , by Equation (48), which means that ρ_θ is not well-parameterized at $\theta = \theta^*$. That is, the parameterization is redundant at this point. We assume this is not the case, and hence take $\mathcal{F}$ invertible.

We now develop material specific to Zhou and Chen's construction.

Definition 8.2 (Deviation observables). Let $\{X_i\}_{i \in [m]}$ be a set of Hermitian operators satisfying, for all $i, j \in [m]$:

$$(i) \text{tr}(X_i \cdot \rho_{\theta^*}) = 0, \quad \text{and} \quad (ii) \text{tr}(X_i \cdot \partial_j \rho_\theta)|_{\theta=\theta^*} = \delta_{ij}.$$

Then the X_i are called *deviation observables*.

Lemma 8.3. Let $\{L_i\}_{i \in [m]}$ be a set of Hermitian symmetric logarithm derivative matrices. For each $i \in [m]$, define $X_i := \sum_{j=1}^m (\mathcal{F}^{-1})_{ij} \cdot L_j$. Then $\{X_i\}_{i \in [m]}$ is a set of deviation observables.

Proof. We verify the two conditions in Theorem 8.2. First,

$$\text{tr}(X_i \cdot \rho_{\theta^*}) = \sum_{j=1}^m (\mathcal{F}^{-1})_{ij} \cdot \text{tr}(L_j \cdot \rho_{\theta^*}) = 0, \quad (52)$$

since $\text{tr}(L_j \cdot \rho_{\theta^*}) = 0$ using Equation (49). Second,

$$\text{tr}(X_i \cdot \partial_j \rho_\theta)|_{\theta=\theta^*} = \sum_{k=1}^m (\mathcal{F}^{-1})_{ik} \cdot \text{tr}(L_k \cdot \partial_j \rho_\theta)|_{\theta=\theta^*} = \sum_{k=1}^m (\mathcal{F}^{-1})_{ik} \cdot \mathcal{F}_{kj} = (\mathcal{F}^{-1} \cdot \mathcal{F})_{ij} = \delta_{ij},$$

since $\text{tr}(L_k \cdot \partial_j \rho_\theta)|_{\theta=\theta^*} = \mathcal{F}_{kj}$ by Equation (50). $\square$

Definition 8.4 (Local shadow estimators). Let $\{X_i\}_{i \in [m]}$ be a set of deviation observables, and let $\hat{\rho}$ be an estimator for ρ_θ . Then the corresponding *local shadow estimator* $\hat{\theta} \in \mathbb{R}^m$ is given by

$$\hat{\theta}_i := \theta_i^* + \text{tr}(X_i \cdot \hat{\rho}),$$

for all $i \in [m]$.

Recall that an estimator is locally unbiased if the following conditions are met for all $i \in [m]$:

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

Lemma 8.5. *Let $\{X_i\}$ be a set of deviation observables, and let $\hat{\rho}$ be the output of an unbiased tomography algorithm, given copies of ρ_θ . Then the corresponding local shadow estimator is a locally unbiased estimator for θ .*

Proof. We verify the two conditions that define a locally unbiased estimator. Firstly, for all $i \in [m]$,

$$\mathbf{E}[\hat{\theta}_i \mid \rho_{\theta^*}] = \theta_i^* + \text{tr}(X_i \cdot \mathbf{E}[\hat{\rho} \mid \rho_{\theta^*}]) = \theta_i^* + \text{tr}(X_i \cdot \rho_{\theta^*}) = \theta_i^*,$$

so that $\mathbf{E}[\hat{\theta} \mid \rho_{\theta^*}] = \theta^*$. Secondly, for all $i, j \in [m]$,

$$\left. \partial_i \mathbf{E}[\hat{\theta}_j \mid \rho_\theta] \right|_{\theta=\theta^*} = \text{tr}(X_j \cdot \partial_i \mathbf{E}[\hat{\rho} \mid \rho_\theta]) \Big|_{\theta=\theta^*} = \text{tr}(X_j \cdot \partial_i \rho_\theta) \Big|_{\theta=\theta^*} = \delta_{ij}. \quad \square$$

Finally, Zhou and Chen's estimator is then the local shadow estimator, constructed from some fixed unbiased tomography algorithm and deviation observables $\{X_i\}_{i \in [m]}$. The $\{X_i\}$ are built from the symmetric logarithm derivative matrices $\{L_i\}_{i \in [m]}$ as in [Theorem 8.3](#), with L_i is given explicitly by [Equation \(48\)](#). By [Theorem 8.5](#), this estimator is locally unbiased. In [\[ZC25\]](#), Zhou and Chen apply this using the bias-corrected single copy estimator of [Section 1.1.1](#).

We are now ready to prove the main result of the section. Our argument is largely based on the proofs of [Theorem 1](#) and [Theorem 4](#) of [\[ZC25\]](#).

Proof of [Theorem 8.1](#). We use Zhou and Chen's construction with the debiased Key's algorithm. Given access to $\rho_\theta^{\otimes n}$, the debiased Key's algorithm produces an unbiased estimator $\hat{\rho}$ with

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

by [Theorems 1.3](#) and [1.4](#). The MSEM is the matrix V_n with entries:

$$(V_n)_{ij} = \mathbf{E}[(\hat{\theta}_i - \theta_i^*) \cdot (\hat{\theta}_j - \theta_j^*) \mid \rho_{\theta^*}] = \mathbf{E}[\text{tr}(X_i \cdot \hat{\rho}) \cdot \text{tr}(X_j \cdot \hat{\rho}) \mid \rho_{\theta^*}] = \text{tr}(X_i \otimes X_j \cdot \mathbf{E}[\hat{\rho} \otimes \hat{\rho} \mid \rho_{\theta^*}]).$$

Plugging in our variance gives:

$$\begin{aligned} (V_n)_{ij} &= \frac{n-1}{n} \text{tr}(X_i \rho_{\theta^*}) \cdot \text{tr}(X_j \rho_{\theta^*}) + \frac{1}{n} (\text{tr}(\rho_{\theta^*} X_i X_j) + \text{tr}(\rho_{\theta^*} X_j X_i)) + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \text{tr}(X_i X_j) \\ &\quad - \text{tr}(X_i \otimes X_j \cdot \text{Lower}_{\rho_{\theta^*}}) \\ &= \frac{1}{n} (\text{tr}(\rho_{\theta^*} X_i X_j) + \text{tr}(\rho_{\theta^*} X_j X_i)) + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \text{tr}(X_i X_j) - \text{tr}(X_i \otimes X_j \cdot \text{Lower}_{\rho_{\theta^*}}). \end{aligned}$$

where we have used $\text{tr}(X_i \rho_{\theta^*}) = 0$, as shown in [Equation \(52\)](#). Since $\text{Lower}_{\rho_{\theta^*}}$ is a positive combination of matrices of the form $P \otimes P \cdot \text{SWAP}$, the last term is a positive combination of matrices of the form $\text{tr}(X_i P X_j P)$, and hence symmetric in i and j . Moreover, recall P is positive semidefinite.

Consider the new matrix V'_n , obtained from V_n by dropping the terms containing $\text{Lower}_{\rho_{\theta^*}}$. That is,

$$(V'_n)_{ij} = \frac{1}{n} (\text{tr}(\rho_{\theta^*} X_i X_j) + \text{tr}(\rho_{\theta^*} X_j X_i)) + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \text{tr}(X_i X_j). \quad (53)$$

Both V_n and V'_n are real and symmetric matrices, and hence Hermitian. We claim $V_n \preceq V'_n$. To show this, let $|\psi\rangle = \sum_{i=1}^m z_i |i\rangle$ be an arbitrary state. We have

$$\langle\psi|(V'_n - V_n)|\psi\rangle = \sum_{i,j} \bar{z}_i z_j (V'_n - V_n)_{ij} = \sum_{i,j} \bar{z}_i z_j \text{tr}(X_i \otimes X_j \cdot \text{Lower}_{\rho_{\theta^*}}) = \text{tr}(Q^\dagger \otimes Q \cdot \text{Lower}_{\rho_{\theta^*}}),$$

where $Q := \sum_{i=1}^m z_i X_i$. However, this trace is a positive combination of terms of the form

$$\text{tr}(Q^\dagger \otimes Q \cdot P \otimes P \cdot \text{SWAP}) = \text{tr}(Q^\dagger P Q \cdot P).$$

Since P is positive semidefinite, so too is $Q^\dagger P Q$. Thus, $\text{tr}(Q^\dagger P Q \cdot P) \geq 0$, and $\langle\psi|(V'_n - V_n)|\psi\rangle \geq 0$. Since $|\psi\rangle$ was arbitrary, $V'_n \succeq V_n$. So, it will suffice to bound V'_n in the PSD order.

Now consider the product $\mathcal{F}V'_n\mathcal{F}^T$. From Equation (53), we have

$$\begin{aligned} (\mathcal{F}V'_n\mathcal{F}^T)_{ij} &= \sum_{k,\ell} \mathcal{F}_{ik}(V'_n)_{k\ell}\mathcal{F}_{j\ell} \\ &= \frac{1}{n} \text{tr}(\rho_{\theta^*} \sum_{k=1}^m \mathcal{F}_{ik} X_k \sum_{\ell=1}^m \mathcal{F}_{j\ell} X_\ell) + \frac{1}{n} \text{tr}(\rho_{\theta^*} \sum_{\ell=1}^m \mathcal{F}_{j\ell} X_\ell \sum_{k=1}^m \mathcal{F}_{ik} X_k) + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \text{tr}(\sum_{k=1}^m \mathcal{F}_{ik} X_k \sum_{\ell=1}^m \mathcal{F}_{j\ell} X_\ell). \end{aligned}$$

However, we have

$$\sum_{k=1}^m \mathcal{F}_{ik} X_k = \sum_{k,\ell} \mathcal{F}_{ik} (\mathcal{F}^{-1})_{k\ell} L_\ell = \sum_{\ell} \delta_{i\ell} L_\ell = L_i,$$

where we have used $X_k := \sum_{\ell} (\mathcal{F}^{-1})_{k\ell} L_\ell$. Thus

$$(\mathcal{F}V'_n\mathcal{F}^T)_{ij} = \frac{1}{n} (\text{tr}(\rho_{\theta^*} L_i L_j) + \text{tr}(\rho_{\theta^*} L_j L_i)) + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \text{tr}(L_i L_j).$$

Now, note that the first term is equal to $2\mathcal{F}_{ij}/n$, by definition of $\mathcal{F}_{ij}$. Thus

$$\mathcal{F}V'_n\mathcal{F}^T = \frac{2}{n} \mathcal{F} + \frac{2\mathbf{E}[\ell(\lambda)]}{n^2} K, \quad (54)$$

where K is the matrix given by

$$K_{ij} := \frac{1}{2} \text{tr}(L_i L_j) = \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{((\partial_i \rho_\theta)_{k\ell} \cdot (\partial_j \rho_\theta)_{\ell k})|_{\theta=\theta^*}}{(\alpha_k + \alpha_\ell)^2},$$

where we have used the explicit expressions for L_i in Equation (48). We now bound K in the PSD order. Let $|\psi\rangle = \sum_{i=1}^m z_i |i\rangle$ again be an arbitrary state. Define $R_{k\ell} := \sum_{i=1}^m z_i \cdot (\partial_i \rho_\theta)_{k\ell}|_{\theta=\theta^*}$, and note that

$$(R^\dagger)_{k\ell} = (R)_{\ell k}^* = \sum_{i=1}^m z_i^* \cdot (\partial_i \rho_\theta)_{\ell k}^*|_{\theta=\theta^*} = \sum_{i=1}^m z_i^* \cdot (\partial_i \rho_\theta^\dagger)_{k\ell}|_{\theta=\theta^*} = \sum_{i=1}^m z_i^* \cdot (\partial_i \rho_\theta)_{k\ell}|_{\theta=\theta^*}.$$

So, we have

$$\langle\psi|K|\psi\rangle = \sum_{\substack{k,\ell \in [d] \\ \alpha_k + \alpha_\ell > 0}} \sum_{i,j \in [m]} z_i^* z_j \cdot \frac{((\partial_i \rho_\theta)_{k\ell} \cdot (\partial_j \rho_\theta)_{\ell k})|_{\theta=\theta^*}}{(\alpha_k + \alpha_\ell)^2} = \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{(R^\dagger)_{k\ell} \cdot R_{\ell k}}{(\alpha_k + \alpha_\ell)^2} = \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{|R_{\ell k}|^2}{(\alpha_k + \alpha_\ell)^2},$$

and similarly,

$$\frac{1}{2} \langle\psi|\mathcal{F}|\psi\rangle = \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \sum_{i,j} z_i^* z_j \cdot \frac{((\partial_i \rho_\theta)_{k\ell} \cdot (\partial_j \rho_\theta)_{\ell k})|_{\theta=\theta^*}}{\alpha_k + \alpha_\ell} = \sum_{\substack{k,\ell \\ \alpha_k + \alpha_\ell > 0}} \frac{|R_{\ell k}|^2}{\alpha_k + \alpha_\ell}.$$

Now since

$$\langle \psi | K | \psi \rangle = \sum_{\substack{k, \ell \\ \alpha_k + \alpha_\ell > 0}} \frac{|R_{\ell k}|^2}{(\alpha_k + \alpha_\ell)^2} \leq \frac{1}{\alpha_r} \sum_{\substack{k, \ell \\ \alpha_k + \alpha_\ell > 0}} \frac{|R_{\ell k}|^2}{\alpha_k + \alpha_\ell} = \frac{1}{2\alpha_r} \langle \psi | \mathcal{F} | \psi \rangle,$$

we have $K \preceq \mathcal{F}/2\alpha_r$. Hence, from Equation (54),

$$\mathcal{F}V'_n\mathcal{F}^T \preceq \left(\frac{2}{n} + \frac{4\mathbf{E}[\ell(\boldsymbol{\lambda})]}{n^2} \cdot \frac{1}{2\alpha_r} \right) \cdot \mathcal{F} \preceq \frac{2}{n} \cdot \left(1 + \frac{r}{n\alpha_r} \right) \cdot \mathcal{F}.$$

Recall that $\mathcal{F}^T = \mathcal{F}$, because $\mathcal{F}$ is symmetric. So, by multiplying on the left and right by $\mathcal{F}^{-1}$, we have

$$V'_n \preceq \frac{2}{n} \cdot \left(1 + \frac{r}{n\alpha_r} \right) \cdot \mathcal{F}^{-1}.$$

Finally, since $V_n \preceq V'_n$, for $n \geq r/(\varepsilon\alpha_r)$ we conclude

$$V_n \preceq (1 + \varepsilon) \cdot \frac{2}{n} \cdot \mathcal{F}^{-1}. \quad \square$$

Part III

The Clebsch-Gordan transform

9 The Clebsch-Gordan transform

The tensor product of two irreps of $U(d)$ is generically reducible. That is, for λ and μ partitions of length at most d , we have the decomposition

$$V_\lambda^d \otimes V_\mu^d \cong^{U(d)} \bigoplus_{\ell(\tau) \leq d} m_\tau \cdot V_\tau^d,$$

for some integers $\{m_\tau\}$. It will suffice for us to consider only the special case where $\mu = (1)$. In this case, the decomposition is multiplicity-free [Har05, Chapter 7.2]:

$$V_\lambda^d \otimes V_{(1)}^d \cong^{U(d)} \bigoplus_{\substack{\lambda \nearrow \mu \\ \ell(\mu) \leq d}} V_\mu^d. \quad (55)$$

Recall that for $\mu = (1)$, our choice of V_μ^d is the defining representation on $\mathbb{C}^d$, and our choice of GT basis is the computational basis $\{|i\rangle\}_{i \in [d]}$, on which $\nu_\mu(U)$ acts as U (see Theorem 2.45). Therefore, by Theorem 2.17, for each λ there exists a unitary $\mathcal{U}_{\text{CG}}^{(\lambda)} : V_\lambda^d \otimes \mathbb{C}^d \rightarrow \bigoplus_\mu V_\mu^d$ that implements the isomorphism in Equation (55). In other words, for all $U \in U(d)$,

$$\mathcal{U}_{\text{CG}}^{(\lambda)} \cdot (\nu_\lambda(U) \otimes U) \cdot \mathcal{U}_{\text{CG}}^{(\lambda)\dagger} = \sum_{\substack{\lambda \nearrow \mu \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \nu_\mu(U). \quad (56)$$

With these unitaries fixed, we can collect them into a single operator.

Definition 9.1 (The Clebsch-Gordan transform). The *Clebsch-Gordan transform* is the unitary operator on $\bigoplus_\lambda \text{Sp}_\lambda \otimes (V_\lambda^d \otimes \mathbb{C}^d) \cong (\mathbb{C}^d)^{\otimes(n+1)}$ given by

$$\mathcal{U}_{\text{CG}}^{(n+1)} := \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \mathcal{U}_{\text{CG}}^{(\lambda)}. \quad (57)$$

The Clebsch-Gordan transform takes the state $|\lambda\rangle \otimes |S\rangle \otimes |T\rangle \otimes |i\rangle$, to a superposition of states of the form $|\lambda\rangle \otimes |S\rangle \otimes |\mu\rangle \otimes |T'\rangle$, where $\lambda \nearrow \mu$, and $T' \in \text{SSYT}(\mu, d)$.

10 Clebsch-Gordan coefficients

Definition 10.1 (Clebsch-Gordan coefficients). Let $\lambda \vdash n$, $T \in \text{SSYT}(\lambda, d)$, and $k \in [d]$. We can write:

$$\mathcal{U}_{\text{CG}}^{(\lambda)} \cdot |T\rangle \otimes |k\rangle = \sum_{\substack{\mu \vdash (n+1, d) \\ T' \in \text{SSYT}(\mu)}} \langle \mu, T' | \mathcal{U}_{\text{CG}}^{(\lambda)} | T, k \rangle \cdot |\mu, T'\rangle.$$

The coefficients $\langle \mu, T' | \mathcal{U}_{\text{CG}}^{(\lambda)} | T, k \rangle$ are known as the *Clebsch-Gordan (CG) coefficients*. We will typically write $\langle T' | T, k \rangle$ as shorthand for $\langle \mu, T' | \mathcal{U}_{\text{CG}}^{(\lambda)} | T, k \rangle$, leaving $\mu = \text{sh}(T')$ and $\mathcal{U}_{\text{CG}}^{(\lambda)}$ implicit.

Our proofs will rely on explicit formulas for the CG coefficients. To state these formulas, we will first need to introduce new notation which will allow us to describe how an SSYT T changes when we restrict it to a specific subset of letters.

Notation 10.2. Let λ be a Young diagram. We define $\lambda_{\leq k} = (\lambda_1, \dots, \lambda_k)$ to be the first k rows of λ .

Notation 10.3 (Restriction of SSYT to alphabet $[p]$). Let T be an SSYT of shape λ and alphabet $[d]$. We use $T^{[p]}$ to denote the SSYT of alphabet $[p]$ that we obtain when we delete all the boxes in T which contain letters strictly larger than p . Since the entries of T are weakly increasing along rows and strictly increasing along columns, $T^{[p]}$ is a valid SSYT of height at most p . We also define $T^{\leq p} := \text{sh}(T^{[p]})$ to be its shape.

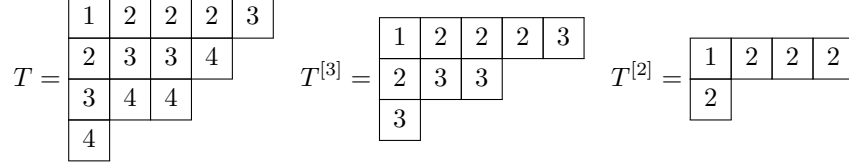


Figure 13: An example of the SSYT notation that we will use in this section. Given the SSYT T over the alphabet $[4]$, $T^{[3]}$ and $T^{[2]}$ are the SSYT we obtain when we delete all the boxes which are bigger than 3 and 2 respectively. Thus $T^{\leq 4} = (5, 4, 3, 1)$, $T^{\leq 3} = (5, 3, 1)$, and $T^{\leq 2} = (4, 1)$.

Notation 10.4 (Horizontal strip). Let T be an SSYT. We define

$$T^{\neq p} := (T^{\leq p}, T^{\leq (p-1)}) = T^{\leq p} \setminus T^{\leq (p-1)}$$

to be the skew shape that contains all boxes with value p . Due to the properties of an SSYT, this skew shape is a *horizontal strip*, that is, a union of continuous horizontal lines of boxes that do not overlap vertically.

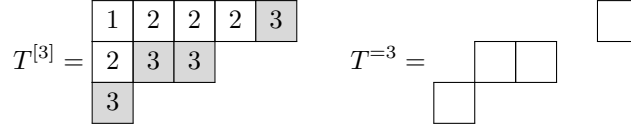


Figure 14: Let T be the SSYT from Figure 13. Above, we have $T^{[3]}$ and $T^{\neq 3}$. The cells of $T^{\neq 3}$ are shaded in $T^{[3]}$ – these are the cells containing the symbol 3.

We introduce two additional pieces of notation that correspond to the shapes that we get when we add boxes to the horizontal strip $T^{\neq p} = (T^{\leq p}, T^{\leq (p-1)})$.

Notation 10.5 (Adding boxes to a horizontal strip $T^{\neq p}$). Let T be an SSYT. We define

$$T_{+i}^{\neq p} := (T^{\leq p} + e_i, T^{\leq (p-1)})$$

to be the shape of the skew diagram that we get when we add a box to the end of the i -th row of $T^{\neq p}$. Similarly,

$$T_{+i, -j}^{\neq p} := (T^{\leq p} + e_i, T^{\leq (p-1)} + e_j)$$

is the shape of the skew diagram we obtain when we add a box to the end of the i -th row and remove a box from the beginning of the j -th row of $T^{\neq p}$.

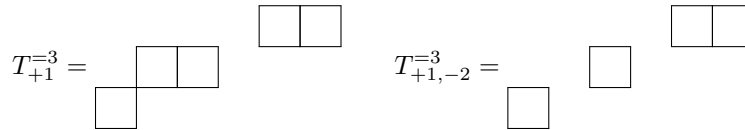


Figure 15: An example of the skew shapes that we get when we add and subtract cells from the skew shape $T^{\neq 3}$ from Figure 14. On the left diagram, we have added a cell to the first row. On the right diagram, we have both added a cell to the first row and removed a cell from the second row.

It is well-known that the CG coefficients can be expressed as a product of numbers that only depend on the shapes of the skew diagrams $\{T^{=p}, \boxed{k}^{=p}, (T')^{=p}\}_{p \in \{1, \dots, d\}}$. In particular, the following was shown by [VK92, Section 18.2.2, Equation 3].

Theorem 10.6 (CG coefficients are products of scalar factors). *When Section 9 is multiplicity-free, the Clebsch-Gordan coefficients can be written as a product:*

$$\langle T' | T, k \rangle = \prod_{p=1}^d (T^{=p}, \boxed{k}^{=p} | (T')^{=p}). \quad (58)$$

The factor $(T^{=p}, \boxed{k}^{=p} | (T')^{=p})$ that appears on the right-hand side is called the $U(d-p)$ -scalar factor of the Clebsch-Gordan coefficient. These coefficients are also known as reduced Wigner operators. Note that each $U(d-p)$ scalar factor only depends on the shapes of the skew diagrams $T^{=p}, \boxed{k}^{=p}, (T')^{=p}$.

Here we use $\boxed{k}$ to denote the SSYT that corresponds to the $|k\rangle$ basis vector of $V_{(1)}^d$, which has a single box filled with the number k . The shape of the restricted SSYT $\boxed{k}^{\leq p}$ has the following simple description:

$$\boxed{k}^{\leq p} = \begin{cases} (1, 0^{p-1}) & p \geq k, \\ (0, 0^{p-1}) & p < k. \end{cases}$$

Thus, the skew diagrams that can appear in the scalar factors of Equation (58) are one of three possible choices:

$$\boxed{k}^{=p} = \begin{cases} e_{11}^{=p} := ((1, 0^{p-1}), (1, 0^{p-2})) & p \geq k+1, \\ e_{10}^{=p} := ((1, 0^{p-1}), (0, 0^{p-2})) & p = k, \\ e_{00}^{=p} := ((0, 0^{p-1}), (0, 0^{p-2})) & p < k. \end{cases}$$

The scalar factors which appear in Equation (58) can be computed using the formulas in the following lemma.

Lemma 10.7 (Formulas for the $U(d)$ -scalar factors). *Let $c_s(T^{\leq k}) = T_s^{\leq k} - s$ be the content of the last box in the s -th row. Define $S(i, j) = 1$ if $i \leq j$, and $S(i, j) = -1$ if $i > j$. Then, for some choice of phases of the GT basis for ν_λ , the following three equations hold.*

(i) For any SSYT T and any p ,

$$(T^{=p}, e_{00}^{=p} | T^{=p}) = 1, \quad (59)$$

and $(T^{=p}, e_{00}^{=p} | (T')^{=p}) = 0$ otherwise.

(ii) Section 18.2.10, equation 3 in [VK92]

$$(T^{=p}, e_{10}^{=p} | T_{+i}^{=p}) = \left| \frac{\prod_{j=1}^{p-1} (c_j(T^{\leq p-1}) - c_i(T^{\leq p}) - 1)}{\prod_{j=1, j \neq i}^p (c_j(T^{\leq p}) - c_i(T^{\leq p}))} \right|^{1/2}, \quad (60)$$

and $(T^{=p}, e_{10}^{=p} | (T')^{=p}) = 0$ otherwise.

(iii) Section 18.2.10, equation 4 in [VK92]

$$\begin{aligned} & (T^{=p}, e_{11}^{=p} | T_{+i, -j}^{=p}) \\ &= S(i, j) \left| \frac{\prod_{k \neq j}^{p-1} (c_k(T^{\leq p-1}) - c_i(T^{\leq p}) - 1) \prod_{k \neq i}^p (c_k(T^{\leq p}) - c_j(T^{\leq p-1}))}{\prod_{k \neq i}^p (c_k(T^{\leq p}) - c_i(T^{\leq p})) \prod_{k \neq j}^{p-1} (c_k(T^{\leq p-1}) - c_j(T^{\leq p-1}) - 1)} \right|^{1/2}, \end{aligned} \quad (61)$$

and $(T^{=p}, e_{11}^{=p} | (T')^{=p}) = 0$ otherwise.

Equation (58) expresses the Clebsch-Gordan coefficients as products of d scalar factors. For the tableaux that we are dealing with in this paper (that are close to highest-weight tableaux), it will be the case that most of these scalar factors are trivial. The following lemma gives a condition for a scalar factor to be trivial that will be sufficient for our purposes.

Lemma 10.8. *Let T be an SSYT, and $i \in [p]$ be a row such that T has no cells with the value p on the i -th row. Then*

$$(T^{=p}, e_{11}^{=p} \mid T_{+i,-i}^{=p}) = 1.$$

Proof. Since T has no p values in the i -th row, it holds that $c_i(T^{\leq p}) = c_i(T^{\leq p-1})$. Then Equation (61) implies that

$$\begin{aligned} (T^{=p}, e_{11}^{=p} \mid T_{+i,-i}^{=p}) &= S(i, i) \left| \frac{\prod_{k \neq i}^{p-1} (c_k(T^{\leq p-1}) - c_i(T^{\leq p}) - 1) \prod_{k \neq i}^p (c_k(T^{\leq p}) - c_i(T^{\leq p-1}))}{\prod_{k \neq i}^p (c_k(T^{\leq p}) - c_i(T^{\leq p})) \prod_{k \neq i}^{p-1} (c_k(T^{\leq p-1}) - c_i(T^{\leq p-1}) - 1)} \right|^{1/2} \\ &= \left| \frac{\prod_{k \neq i}^{p-1} (c_k(T^{\leq p-1}) - c_i(T^{\leq p-1}) - 1) \prod_{k \neq i}^p (c_k(T^{\leq p}) - c_i(T^{\leq p-1}))}{\prod_{k \neq i}^p (c_k(T^{\leq p}) - c_i(T^{\leq p-1})) \prod_{k \neq i}^{p-1} (c_k(T^{\leq p-1}) - c_i(T^{\leq p-1}) - 1)} \right|^{1/2} = 1. \quad \square \end{aligned}$$

10.1 Clebsch-Gordan insertion

In this subsection, we describe a simple algorithm for inserting a new box containing a symbol $k \in [d]$ into an SSYT T . The output is a set of SSYTs, each with a matching coefficient, $\{(T', c_{T'})\}$. The SSYTs in this set turn out to be exactly the SSYTs with nonzero Clebsch-Gordan coefficients, and $c_{T'}$ is the Clebsch-Gordan coefficient $\langle T' | T, k \rangle$. This insertion algorithm, *Clebsch-Gordan insertion*, thus offers a more intuitive perspective on the Clebsch-Gordan transform and coefficients. Clebsch-Gordan insertion is, in fact, a special case of a continuous family of “quantum” insertion rules generalizing RSK insertion (see, for example, [Wri16, Chapter 3]), originally introduced by Sonya Berg in her thesis [Ber12].

Definition 10.9 (Clebsch-Gordan insertion). Given an SSYT T with alphabet $[d]$, we insert a letter $k \in [d]$ into T to obtain a set of output SSYTs as follows:

1. Initialize $T' \leftarrow T$, $c_{T'} \leftarrow 1$, and $j \leftarrow \text{None}$. Here, T' is a tentative output SSYT, $c_{T'}$ its tentative coefficient, and j stores a previous row index.
2. In all possible ways, choose a row $i \leq k$ such that $T^{\leq k} + e_i$ is a valid Young diagram.
3. Find the leftmost cell of T in row i containing a letter ℓ such that $\ell > k$, if such a cell exists. Let T'' be obtained from T' by replacing ℓ with k in this cell. Set

$$T' \leftarrow T'', \quad c_{T'} \leftarrow c_{T'} \cdot \text{SF}(T, k; i, j), \quad k \leftarrow \ell, \quad j \leftarrow i,$$

and return to Step 2. We will say ℓ has been *bumped* by k .

4. If no such cell exists, then append a new cell containing k to the end of row i to obtain T'' . Set

$$T' \leftarrow T'', \quad c_{T'} \leftarrow c_{T'} \cdot \text{SF}(T, k; i, j),$$

append $(T', c_{T'})$ to the output, and terminate.

Here, $\text{SF}(T, k; i, j)$ is a scalar factor, given by

$$\text{SF}(T, k; i, j) = \begin{cases} (T^{=k}, e_{10}^{=k} \mid T_{+i}^{=k}) & j = \text{None}, \\ (T^{=k}, e_{11}^{=k} \mid T_{+i,-j}^{=k}) & \text{otherwise.} \end{cases}$$

See Figure 16 for an example of Clebsch-Gordan insertion.

Lemma 10.10. *Suppose we insert a letter $k \in [d]$ into an SSYT T with alphabet $[d]$ via Clebsch-Gordan insertion, obtaining as output a set $\{(T', c_{T'})\}$. Then the SSYTs $\{T'\}$ are exactly those with nonzero Clebsch-Gordan coefficients, and $c_{T'} = \langle T' | T, k \rangle$.*

Proof. We prove the two directions of the lemma separately.

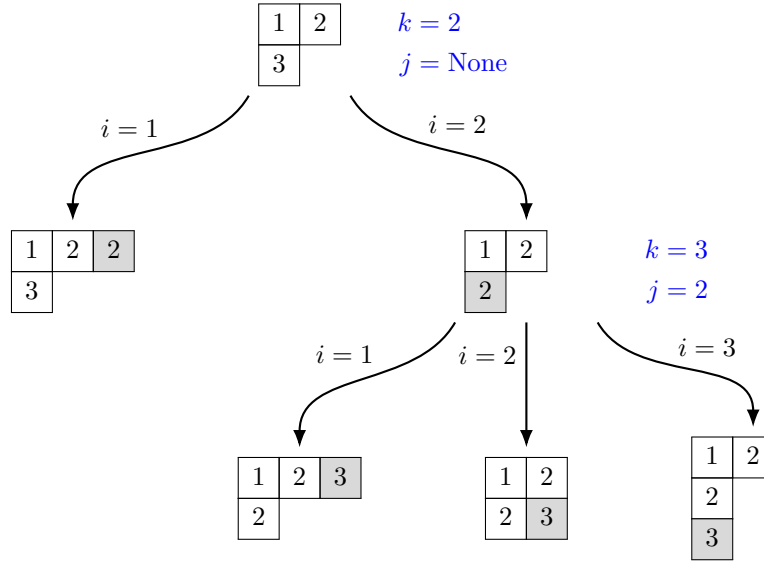


Figure 16: Here we depict Clebsch-Gordan insertion of $k = 2$ into the example SSYT

$$\begin{array}{|c|c|} \hline 1 & 2 \\ \hline 3 & \\ \hline \end{array}.$$

At the start of the insertion process, $k = 2$ is to be inserted into either of the first two rows. On the left branch, 2 is inserted at the end of the first row, and the insertion terminates. On the right branch, 2 bumps 3 from the second row, and we must now re-insert 3 into the tableau. Hence, $k \leftarrow 3$. We also update $j \leftarrow 1$ at this stage, since we had inserted the 2 into the second row. The 3 may now be inserted at the end of any of the first, second, or third rows, and the insertion terminates.

Clebsch-Gordan insertion outputs are valid SSYT with nonzero Clebsch-Gordan coefficients: Say that the Clebsch-Gordan insertion algorithm output includes the pair $(T', c_{T'})$. Moreover, assume that T' was constructed by adding k to row j_k , and then sequentially bumping the letters $\ell_1 < \ell_2 < \dots < \ell_t$, which were inserted into rows $j_{\ell_1}, \dots, j_{\ell_t}$ respectively. The case of $t = 0$ corresponds to k not bumping any letter. For convenience, let us define $\ell_0 = k$ and $\ell_{t+1} = d + 1$.

We show that $c_{T'} = \langle T' | T, k \rangle$. In particular, let us define the sequence $i_k, i_{k+1}, \dots, i_d$, such that

$$j_k = i_k = i_{k+1} = \dots = i_{\ell_1-1}, \quad j_{\ell_m} = i_{\ell_m} = \dots = i_{\ell_{m+1}-1}, \quad (62)$$

for all $m \in [t]$. Then we claim that

$$(T')^{\leq p} = \begin{cases} T^{\leq p} & p < k, \\ T^{\leq p} + e_{i_p} & p \geq k. \end{cases}$$

To see this, note that inserting an x into a tableau can only change the locations of symbols labeled x or greater. Thus, the expression above holds for $p < k$ because k is the smallest symbol inserted into the tableau. For $p \geq k$, suppose that ℓ_m is the final letter that is inserted which satisfies $\ell_m \leq p$. Then inserting the letters $\ell_{m+1}, \dots, \ell_t$ does not change $T^{\leq p}$ since they are all larger than p . Similarly, the letters $\ell_0 < \ell_1 < \dots < \ell_{m-1}$ which are inserted prior to ℓ_m do not change the shape of $T^{\leq p}$ either, because they replace letters which are also $\leq p$. Thus, only ℓ_m will change the shape $T^{\leq p}$, either by adding a new box or by bumping a letter larger than p in row i_{ℓ_m} , and this will update the shape to $T^{\leq p} + e_{i_{\ell_m}} = T^{\leq p} + e_{i_p}$, as claimed. Therefore,

$$(T')^{=p} = \begin{cases} T^{=p} & p < k, \\ T_{+i_p}^{=p} & p = k, \\ T_{+i_p, -i_{p-1}}^{=p} & p > k. \end{cases} \quad (63)$$

Thus, the CG coefficient is

$$\begin{aligned} \langle T' | T, k \rangle &= \prod_{p=1}^d (T^{=p}, \boxed{k}^{=p} | (T')^{=p}) \\ &= \prod_{p \in \{\ell_0, \ell_1, \dots, \ell_t\}} (T^{=p}, \boxed{k}^{=p} | (T')^{=p}) \\ &= (T^{=k}, e_{10}^{=k} | T_{+j_k}^{=k}) \cdot \prod_{m=1}^t (T^{=\ell_m}, e_{11}^{=\ell_m} | T_{j_{\ell_m}, -j_{\ell_{m-1}}}^{=\ell_m}) \\ &= \text{SF}(T, k; j_k, \text{None}) \cdot \prod_{m=1}^t \text{SF}(T, \ell_m; j_{\ell_m}, j_{\ell_{m-1}}) = c_{T'}. \end{aligned}$$

In the first step, we have used [Theorem 10.6](#). In the second step, we have set all scalar products with $p \notin \{\ell_0, \ell_1, \dots, \ell_t\}$ equal to 1. For $p < k = \ell_0$, this holds by item (i) of [Theorem 10.7](#). For $p > k$ and $p \notin \{\ell_1, \dots, \ell_t\}$, this follows by [Theorem 10.8](#), since $i_p = i_{p-1}$ from our definition of the indices i ([Equation \(62\)](#)), and the fact that T contains no letter p in this row. To see this, let ℓ_m be the largest letter that was bumped which is smaller than p . Then ℓ_m was added to row $j_{\ell_m} = i_{p-1} = i_p$ and bumped the letter ℓ_{m+1} , which is larger than p . Thus, there is no letter p in the row i_p , as otherwise that would have been bumped instead.

Valid SSYT with nonzero CG coefficients are output by the Clebsch-Gordan insertion algorithm:

Let T' be a valid SSYT with nonzero CG coefficient $\langle T' | T, k \rangle$. Our first goal is to show that it must satisfy [Equation \(63\)](#), i.e. there exist integers $i_k, i_{k+1}, \dots, i_d$ such that

$$(T')^{=p} = \begin{cases} T^{=p} & p < k, \\ T_{+i_p}^{=p} & p = k, \\ T_{+i_p, -i_{p-1}}^{=p} & p > k. \end{cases} \quad (64)$$

To see this, note that Equation (58) implies that for all $1 \leq p \leq d$,

$$\langle T' | T, k \rangle = \langle T'^{[p]} | T^{[p]}, \boxed{k}^{[p]} \rangle \cdot \prod_{i=p+1}^d (T'^{=i}, \boxed{k}^{=i} | T'^{=i}). \quad (65)$$

Hence, for $\langle T' | T, k \rangle$ to be nonzero, we must have

$$\langle T'^{[p]} | T^{[p]}, \boxed{k}^{[p]} \rangle \neq 0.$$

When $p < k$, $\boxed{k}^{[p]}$ is just an empty tableau. Then we must have $T'^{[p]} = T^{[p]}$. On the other hand, when $p \geq k$, $\boxed{k}^{[p]}$ is a single box with k , and thus

$$\langle T^{[p]}, \boxed{k}^{[p]} | T'^{[p]} \rangle \neq 0 \implies \text{sh}(T'^{[p]}) = \text{sh}(T^{[p]}) + e_{i_p}.$$

for some integer i_p , by Equation (55). Putting these together, we have that for $p < k$,

$$T'^{=p} = (T'^{\leq p}, T'^{\leq p-1}) = (T^{\leq p}, T^{\leq p-1}) = T'^{=p}.$$

Similarly, for $p = k$, we have $T'^{=k} = T'^{=p}_{+i_k}$, and for $p > k$, we have $T'^{=p} = T'^{=p}_{+i_p, -i_{p-1}}$. This proves Equation (64).

Next, we show that Equation (64) implies that T' can be viewed as the result of inserting a letter k into T and performing a sequence of bumps, as in the Clebsch-Gordan insertion algorithm. We do this by explicitly determining the sequence of bumped letters that produces this SSYT T' . The first letter added was k , and it was inserted at row i_k . Thus, we let $j_k = i_k$. The leftmost cell of T in row j_k that contains a letter greater than k determines the first bumped letter ℓ_1 . This is because for any $p \in \{k, \dots, \ell_1 - 1\}$, $T'^{\leq p} = T^{\leq p} + e_{i_p}$, and since there is no letter greater than k and at most p on row i_k , it must hold that $i_p = i_k$. Thus $i_k = i_{k+1} = \dots = i_{\ell_1-1}$, and we let $j_{\ell_1} = i_{\ell_1}$.

Continuing the above process, the leftmost cell of T in row $j_{\ell_{m-1}}$ that contains a letter greater than ℓ_{m-1} determines the bumped letter ℓ_m (and if no such letter exists, the Clebsch-Gordan insertion would terminate). We let $j_{\ell_m} = i_{\ell_m}$ and observe that $i_{\ell_{m-1}} = \dots = i_{\ell_m-1}$.

Having defined the bumped letters $\ell_1 < \dots < \ell_t$ and their respective rows $j_k, j_{\ell_1}, \dots, j_{\ell_t}$, the argument from the first part of this proof shows that the tableau output by the Clebsch-Gordan insertion procedure is the same as T' . This is because by our choice of the j indices, the row indices $i_k, \dots, i_d$ satisfy:

$$j_k = i_k = i_{k+1} = \dots = i_{\ell_1-1}, \quad j_{\ell_m} = i_{\ell_m} = \dots = i_{\ell_{m+1}-1},$$

for all $m \in [t]$ (where we again assume $\ell_{t+1} = d+1$). □

10.2 One-step coefficients

In this section, we obtain a closed-form expression for the Clebsch-Gordan coefficients $\langle T' | T^\lambda, k \rangle$, where we recall that $|T^\lambda\rangle$ is the highest weight vector of V_λ^d . We first show that the SSYTs T' for which this coefficient is nonzero have a very specific form.

Definition 10.11 (One-step semistandard Young tableaux). Let λ be a Young diagram of height at most d , and $i \leq k$ be positive integers at most d such that $\lambda + e_i$ is a valid Young diagram. We use $T_{k \rightarrow i}^\lambda$ to denote the *one-step semistandard Young tableau*, which is the SSYT we obtain by starting with the highest weight SSYT T^λ , and adding a new box containing the value k to the end of the i -th row.

Definition 10.12 (One-step Clebsch-Gordan coefficient). Let $T_{k \rightarrow i}^\lambda$ be a one-step SSYT. We define the *one-step Clebsch-Gordan coefficient* to be

$$c_{k \rightarrow i}^\lambda := \langle T_{k \rightarrow i}^\lambda | T^\lambda, k \rangle,$$

when $T_{k \rightarrow i}^\lambda$ is well-defined (i.e. a valid SSYT), and 0 otherwise.

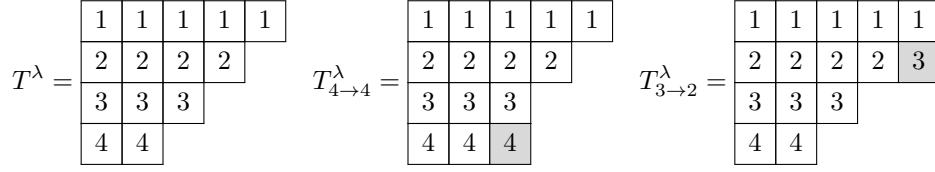


Figure 17: Some examples of one-step SSYT when $\lambda = (5, 4, 3, 2)$.

Lemma 10.13. *The only basis vectors T' for which the coefficient $\langle T' | T^\lambda, k \rangle$ is nonzero correspond to the one-step SSYT $T_{k \rightarrow i}^\lambda$, for any $i \leq k$ such that $\lambda + e_i$ is a valid Young diagram. Moreover, the one-step Clebsch-Gordan coefficient satisfies the following expression:*

$$c_{k \rightarrow i}^\lambda = \left| \prod_{j=1, j \neq i}^{k-1} \left(\frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \right) \cdot \frac{1}{(\lambda_i - \lambda_k) + (k - i) + \delta_{ik}} \right|^{1/2}. \quad (66)$$

Proof. Consider the Clebsch-Gordan insertion procedure from [Theorem 10.9](#). When inserting k into row $i \leq k$ during Step 2, the highest weight SSYT T^λ only contains the letter i , which is not greater than k . Thus, the insertion procedure does not reach Step 3, but always jumps to Step 4. See [Figure 18](#). Step 4 will now append a new cell containing k to the end of row i to obtain $T_{k \rightarrow i}^\lambda$, and will set $c_{k \rightarrow i}^\lambda$ to be

$$c_{k \rightarrow i}^\lambda \leftarrow \text{SF}(T^\lambda, k; i, \text{None}) = ((T^\lambda)^{=k}, e_{10}^{=k} \mid (T^\lambda)_{+i}^{=k}).$$

Finally, the procedure appends $(T_{k \rightarrow i}^\lambda, c_{k \rightarrow i}^\lambda)$ to the output and terminates. We will now use [Equation \(60\)](#) to compute an expression for the one-step Clebsch-Gordan coefficient.

$$\begin{aligned} c_{k \rightarrow i}^\lambda &= ((T^\lambda)^{=k}, e_{10}^{=k} \mid (T^\lambda)_{+i}^{=k}) \\ &= \left| \frac{\prod_{j=1}^{k-1} ((\lambda_j - \lambda_i) + (i - j) - 1)}{\prod_{j \neq i}^k ((\lambda_j - \lambda_i) + (i - j))} \right|^{1/2} && (T^\lambda \text{ is highest weight} \implies T_i^\lambda = \lambda_i) \\ &= \left| \prod_{j=1, j \neq i}^{k-1} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \cdot \frac{-1}{(\lambda_k - \lambda_i) + (i - k) + \delta_{ik}} \right|^{1/2}. \end{aligned}$$

The last equality holds because if $i = k$, then the denominator of the additional factor is equal to 1, and otherwise equal to $(\lambda_k - \lambda_i) + (i - k)$. $\square$

Before we proceed, let us build some intuition about these coefficients. We first observe that all possible one-step SSYT have a nonzero coefficient.

Observation 10.14. *Every valid one-step SSYT $T_{k \rightarrow i}^\lambda$ has a nonzero one-step Clebsch-Gordan coefficient $c_{k \rightarrow i}^\lambda$.*

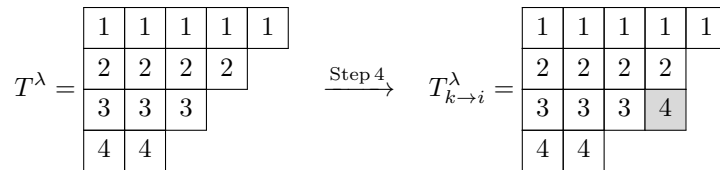


Figure 18: An example of how T' evolves during the Clebsch-Gordan insertion when $k = 4$, and $i = 3$, for $\lambda = (5, 4, 3, 2)$.

Looking at Equation (66), we can see that the only way for $c_{k \rightarrow i}^\lambda$ to be zero is when the expression in the numerator $(\lambda_j - \lambda_i) + (i - j) - 1$ is equal to zero for some j in $\{1, \dots, k-1\} \setminus \{i\}$. Rewriting this, it means that $\lambda_i = \lambda_j + (i - j - 1)$. We consider the following two cases:

1. If $j < i$, this means that $i - j - 1 \geq 0$ and thus $\lambda_i \geq \lambda_j$. Since the j -th row is above the i -th row, we conclude that the two row lengths must be equal. Then we cannot add another box to row i , since that would make this row longer than row j , which is above it.
2. If $j > i$, this means that $i - j - 1 \leq -2$ and thus $\lambda_i \leq \lambda_j - 2$. This is impossible, since row i is above row j , and for λ to be a valid Young diagram, it must hold that $\lambda_i \geq \lambda_j$.

We conclude that the only way for the one-step Clebsch-Gordan coefficients to vanish is by adding a box to a row i that satisfies $\lambda_{i-1} = \lambda_i$, which means that $\lambda + e_i$ is not a valid Young diagram.

10.3 Two-step coefficients

In the previous section, we introduced the one-step Clebsch-Gordan coefficients to understand the state we obtain when we tensor a qudit $|k\rangle$ to a highest weight vector. This suffices to study the first moment of our estimator $\hat{\rho}$. For the second moment of the estimator, we will need to understand the state we obtain when we tensor two qudits, i.e. $|k\rangle \otimes |\ell\rangle$, to a highest weight vector. We have

$$\begin{aligned} |T^\lambda\rangle \otimes |k\rangle \otimes |\ell\rangle &= \left(\sum_{i \in [d]} c_{k \rightarrow i}^\lambda \cdot |\lambda + e_i\rangle \otimes |T_{k \rightarrow i}^\lambda\rangle \right) \otimes |\ell\rangle \\ &= \sum_{i, j \in [d]} \sum_{T'} c_{k \rightarrow i}^\lambda \cdot \langle T' | T_{k \rightarrow i}^\lambda, \ell \rangle \cdot |\lambda + e_i\rangle \otimes |\lambda + e_i + e_j\rangle \otimes |T'\rangle, \end{aligned} \quad (67)$$

where, in the inner sum, $T' \in \text{SSYT}(\lambda + e_i + e_j, d)$. Let us first try to understand which SSYTs T' appear in this superposition with nonnegative amplitude. From Theorem 10.14, the coefficient $c_{k \rightarrow i}^\lambda$ in Equation (67) is always nonzero if $T_{k \rightarrow i}^\lambda$ is a valid SSYT. Thus, it suffices to determine for which T' the coefficient $\langle T' | T_{k \rightarrow i}^\lambda, \ell \rangle$ is nonzero. We will need the following definition.

Definition 10.15 (Two-step semistandard Young tableaux). Let λ be a Young diagram of length at most d , and $i \leq k$ and $j \leq \ell$ be positive integers at most d such that $\lambda + e_i + e_j$ is a valid Young diagram.

There is at most one valid SSYT that can be obtained from T^λ by adding a k to row i and an ℓ to row j . If such an SSYT exists, we denote it $T_{k\ell \rightarrow ij}^\lambda$ and refer to it as a *two-step semistandard Young tableau*.

Note that the ordering of the value-row pairs (k, i) and (ℓ, j) does not matter. In particular, $T_{k\ell \rightarrow ij}^\lambda = T_{\ell k \rightarrow ji}^\lambda$. Moreover, we can obtain $T_{k\ell \rightarrow ij}^\lambda$ from T^λ by first appending a k into row i , then appending an ℓ into row j . However, in the case where $k > \ell$, and $i = j$, in which case we must also swap the boxes containing k and ℓ , so that the SSYT is valid. See Figure 19 for some examples of two-step SSYTs.

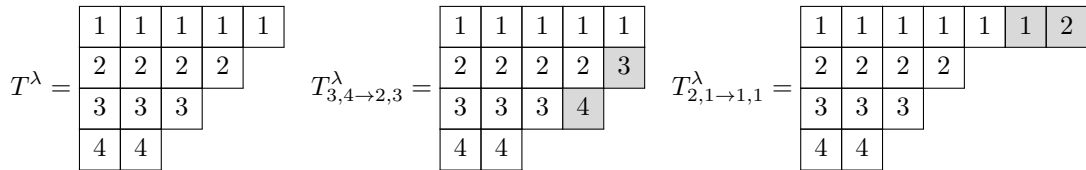


Figure 19: A partition $\lambda = (5, 4, 3, 2)$ and two examples of two-step SSYTs $T_{k\ell \rightarrow ij}^\lambda$. The added boxes are shaded.

Lemma 10.16. *The coefficient $\langle T' | T_{k\ell \rightarrow ij}^\lambda, \ell \rangle$ is nonzero only when T' satisfies one of the following two cases:*

- (i) $T' = T_{k\ell \rightarrow ij}^\lambda$, or

(ii) $T' = T_{\ell k \rightarrow ij}^\lambda$ and $k > \ell$,

where j is any valid row.

Proof. Recall the Clebsch-Gordan insertion procedure from [Theorem 10.9](#), during which we attempt to insert ℓ into a row of $T_{k \rightarrow i}^\lambda$ during Step 2. We consider the following two cases:

1. If we insert ℓ into row $j \leq \ell$ that is different than i , the one-step SSYT $T_{k \rightarrow i}^\lambda$ only contains the letter j in that row. Since j is at most ℓ , the insertion procedure does not reach Step 3, but always jumps to Step 4. Step 4 will now append a new cell containing ℓ to the end of row j to obtain $T_{k\ell \rightarrow ij}^\lambda$, append it (along with the respective coefficient) to the output, and terminate.

$$T_{k \rightarrow i}^\lambda = \begin{array}{|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 \\ \hline 2 & 2 & 2 & 2 & 3 \\ \hline 3 & 3 & 3 & & \\ \hline 4 & 4 & & & \\ \hline \end{array} \xrightarrow{\text{Step 4}} T_{k\ell \rightarrow ij}^\lambda = \begin{array}{|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 \\ \hline 2 & 2 & 2 & 2 & 3 \\ \hline 3 & 3 & 3 & 4 & \\ \hline 4 & 4 & & & \\ \hline \end{array}$$

Figure 20: An example of how T' evolves during the Clebsch-Gordan insertion in the first case from the proof of [Theorem 10.16](#) when $(k, \ell) = (3, 4)$, and $(i, j) = (2, 3)$, for $\lambda = (5, 4, 3, 2)$.

2. If we insert ℓ into row i (assuming $i \leq \ell$), the one-step SSYT $T_{k \rightarrow i}^\lambda$ contains the letter i in the first λ_i cells, and the letter k in the final cell of the row. If $k \leq \ell$, then the insertion procedure does not reach Step 3, but always jumps to Step 4 and will append a new cell containing ℓ to the end of row i to obtain $T_{k\ell \rightarrow ii}^\lambda$.

If, on the other hand, $k > \ell$, then ℓ will replace k at the end of row i . The current tableau becomes $T_{\ell \rightarrow i}^\lambda$, and the procedure will repeat Step 2 with the bumped letter k , which will be added in a valid row $j \leq k$. The current SSYT will contain in row j the letters j , and possibly ℓ , which are not greater than k . Thus, we proceed to Step 4, during which a new cell with the letter k is appended at the end of row j . The final SSYT is thus $T_{\ell k \rightarrow ij}^\lambda$.

$$T_{k \rightarrow i}^\lambda = \begin{array}{|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 \\ \hline 2 & 2 & 2 & 2 & 4 \\ \hline 3 & 3 & 3 & & \\ \hline 4 & 4 & & & \\ \hline \end{array} \xrightarrow{\text{Step 3}} T_{\ell \rightarrow i}^\lambda = \begin{array}{|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 \\ \hline 2 & 2 & 2 & 2 & 3 \\ \hline 3 & 3 & 3 & & \\ \hline 4 & 4 & & & \\ \hline \end{array} \xrightarrow{\text{Step 4}} T_{\ell k \rightarrow ij}^\lambda = \begin{array}{|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 \\ \hline 2 & 2 & 2 & 2 & 3 \\ \hline 3 & 3 & 3 & & \\ \hline 4 & 4 & 4 & & \\ \hline \end{array}$$

Figure 21: An example of how T' evolves during the Clebsch-Gordan insertion in the second case from the proof of [Theorem 10.16](#) when $(k, \ell) = (4, 3)$, and $(i, j) = (2, 4)$, for $\lambda = (5, 4, 3, 2)$.

We conclude that the only SSYTs T' with a nonzero coefficient are either $T_{k\ell \rightarrow ij}^\lambda$, or $T_{\ell k \rightarrow ij}^\lambda$ if $k > \ell$. $\square$

In light of [Theorem 10.16](#), we define the two-step Clebsch-Gordan coefficients as follows.

Definition 10.17 (Two-step Clebsch-Gordan coefficients). Let λ be a Young diagram. The *two-step Clebsch-Gordan coefficients* are:

$$a_{k\ell \rightarrow ij}^\lambda := \langle T_{k \rightarrow i}^\lambda | T^\lambda, k \rangle \cdot \langle T_{k\ell \rightarrow ij}^\lambda | T_{k \rightarrow i}^\lambda, \ell \rangle, \quad b_{k\ell \rightarrow ij}^\lambda := \langle T_{k \rightarrow i}^\lambda | T^\lambda, k \rangle \cdot \langle T_{\ell k \rightarrow ij}^\lambda | T_{k \rightarrow i}^\lambda, \ell \rangle.$$

Whenever a one-step or two-step SSYT is not valid, we define the corresponding coefficient to be 0.

Remark 10.18. The scalar factors that do not involve a bumped letter are always nonnegative, since they come from items (i) and (ii) of [Theorem 10.7](#). Therefore, the two-step Clebsch-Gordan coefficients $a_{k\ell \rightarrow ij}^\lambda$ are always nonnegative.

11 Constructing the Schur transform

The Clebsch-Gordan transform can be used to give a recursive construction of the Schur transform. That is, we can construct a unitary $\mathcal{U}_{\text{SW}}^{(n)}$ that acts as in [Equation \(17\)](#): for all $\pi \in S_n$ and $U \in U(d)$,

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \left(\mathcal{P}^{(n)}(\pi) \mathcal{Q}^{(n)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\pi) \otimes \nu_\lambda(U).$$

This construction is originally due to [\[Har05, BCH05\]](#). Note that in these works, the Schur transform was only required to compute a Young-Yamanouchi basis in the permutation register. We will prove that the construction can be used to prepare Young's orthogonal form as well.

For intuition, we now sketch how we will do a recursive step using the Clebsch-Gordan transform. Assume we have constructed $\mathcal{U}_{\text{SW}}^{(n)}$ already, so that we can rewrite vectors in $(\mathbb{C}^d)^{\otimes n}$ in the Schur basis. We want to give a construction of $\mathcal{U}_{\text{SW}}^{(n+1)}$.

Start with a tensor product $|\lambda\rangle \otimes |S\rangle \otimes |T\rangle \otimes |i\rangle$, obtained by using $\mathcal{U}_{\text{SW}}^{(n)}$ on the first n qudits. By applying the Clebsch-Gordan transform, $\mathcal{U}_{\text{CG}}^{(n+1)}$, we take this state to a superposition of vectors of the form $|\lambda\rangle \otimes |S\rangle \otimes |\mu\rangle \otimes |T'\rangle$, where T' is of shape μ . If we now rearrange factors, we can obtain $|\mu\rangle \otimes (|\lambda\rangle \otimes |S\rangle) \otimes |T'\rangle$. In this product, $|\lambda\rangle \otimes |S\rangle$ is exactly the labeling of the Young-Yamanouchi basis vector corresponding to $|S'\rangle$, where $S' \in \text{SYT}(\mu)$ is obtained from S by appending a box labeled $(n+1)$ at the location of the single box of $\mu \setminus \lambda$. Formally, the rearrangement is performed by the following unitary.

Definition 11.1 (Rearrangement unitary). Let $\mathcal{U}_{\text{R}}^{(n+1)} : \bigoplus_\lambda \text{Sp}_\lambda \otimes \left(\bigoplus_\mu V_\mu^d \right) \rightarrow \bigoplus_\mu \left(\bigoplus_\lambda \text{Sp}_\lambda \otimes V_\mu^d \right)$ mapping

$$|\lambda\rangle \otimes |S\rangle \otimes |\mu\rangle \otimes |T'\rangle \rightarrow |\mu\rangle \otimes |\lambda\rangle \otimes |S\rangle \otimes |T'\rangle,$$

for all $\lambda \vdash n$, $S \in \text{SYT}(\lambda)$, μ with $\lambda \nearrow \mu$ and $T' \in \text{SSYT}(\mu, d)$.

The structure of the new basis vectors suggests that the product $\hat{\mathcal{U}}_{\text{SW}}^{(n+1)} := \mathcal{U}_{\text{R}}^{(n+1)} \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \cdot (\mathcal{U}_{\text{SW}}^{(n)} \otimes I)$ may be performing the Schur transform on $(\mathbb{C}^d)^{\otimes(n+1)}$, i.e.

$$\hat{\mathcal{U}}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(\pi) \mathcal{Q}^{(n+1)}(U) \right) \cdot \hat{\mathcal{U}}_{\text{SW}}^{(n+1)\dagger} \stackrel{?}{=} \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \kappa_\mu(\pi) \otimes \nu_\mu(U). \quad (68)$$

Indeed, loosely speaking, the “unitary register” (the register containing $|T'\rangle$) transforms according to the ν_μ representation, by the definition of the Clebsch-Gordan transform. Moreover, by the uniqueness of the decomposition into irreps by Schur-Weyl duality, the “permutation register” (the register containing $|S'\rangle = |\lambda\rangle \otimes |S\rangle$) transforms according to some representation of S_{n+1} isomorphic to κ_μ , which we recall is Young's orthogonal representation. Moreover, this representation's restriction to S_n is κ_λ (as we will show in [Theorem A.4](#)), so that the permutation register's basis vectors constitute a Young-Yamanouchi basis.

It turns out that [Equation \(68\)](#) *does* hold. Based on this, we make the following definition.

Definition 11.2 (Schur transform). We define the *Schur transform*, $\mathcal{U}_{\text{SW}}^{(n)} : (\mathbb{C}^d)^{\otimes n} \rightarrow \bigoplus_{\lambda \vdash (n,d)} \text{Sp}_\lambda \otimes V_\lambda^d$, by the following recursive construction. We explicitly define $\mathcal{U}_{\text{SW}}^{(1)}$ as the unitary such that

$$\mathcal{U}_{\text{SW}}^{(1)} |i\rangle := |\square\rangle \otimes |\square\rangle \otimes |i\rangle,$$

for all $i \in [d]$. Then, for $n > 1$,

$$\mathcal{U}_{\text{SW}}^{(n)} := \mathcal{U}_{\text{R}}^{(n)} \cdot \mathcal{U}_{\text{CG}}^{(n)} \cdot (\mathcal{U}_{\text{SW}}^{(n-1)} \otimes I).$$

Theorem 11.3. *The unitary in [Theorem 11.2](#) is a Schur transform. That is, it satisfies*

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \left(\mathcal{P}^{(n)}(\pi) \mathcal{Q}^{(n)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\pi) \otimes \nu_\lambda(U), \quad (69)$$

for all $\pi \in S_n$ and $U \in U(d)$.

This theorem is useful for us, because we will eventually want to make use of explicit formulas for matrix elements of Young's orthogonal form. If the permutation registers computed a generic Young-Yamanouchi basis only, then the corresponding matrix elements may contain unknown phases, making them harder to use.

Proving [Theorem 11.3](#) is a little involved, and we defer its proof to [Part V](#).

12 Dual Clebsch-Gordan transform

It will be convenient for us to use the *dual Clebsch-Gordan transform* in our proofs as well. We now give a minimal introduction to what we will need.

Definition 12.1 (Rational representations). A finite-dimensional representation (μ, V) of a matrix group G is *rational*, if, expressed in some basis of V , every matrix element $\mu(g)_{ij}$ is a rational function in the entries of $g \in G$.

A representation being rational is a basis-independent notion, since if the entries of $\mu(g)$ are rational in one basis, then they are also rational in any other basis. Rational representations generalize polynomial representations. The rational irreps can be classified as follows.

Theorem 12.2. *The rational irreps of $U(d)$ are indexed by tuples $\lambda \in \mathbb{Z}^d$, such that $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d$. The irrep corresponding to λ is denoted $(\nu_\lambda, V_\lambda^d)$.*

Unlike Young diagrams, the tuples indexing rational irreps do not necessarily have nonnegative row lengths. These are sometimes called *staircases* or *rational tableaux* in the literature [\[Ngu24\]](#). One new rational representation is particularly important.

Definition 12.3 (Antifundamental representation). The *antifundamental representation* of $U(d)$ is the representation $(\nu_{-\square}, V_{-\square}^d)$ defined by $V_{-\square}^d := \mathbb{C}^d$ and $\nu_{-\square}(U) = U^*$. This is the irrep corresponding to $\lambda = (0^{d-1}, -1)$. The computational basis elements are denoted $\{|\bar{k}\rangle\}_{k \in [d]}$ in this context.

It turns out (see for example [\[Ngu24, Equation \(49\)\]](#)) that the tensor product of a polynomial irrep V_λ^d of $U(d)$ with $V_{-\square}^d$ decomposes into rational irreps in a way similar to [Equation \(55\)](#).

$$V_\lambda^d \otimes V_{-\square}^d \cong \bigoplus_{\substack{\mu = \lambda - \square \\ \ell(\mu) \leq d}}^{U(d)} V_\mu^d. \quad (70)$$

Here by $\lambda - \square$ we mean a staircase tableau of the form $\lambda - e_i$, for any $i \in [d]$. So the right-hand side of the above expression iterates over all valid tuples that can be obtained from λ by removing a box. Note that with tuples, we may also remove boxes from rows with zero or fewer boxes.

Therefore, there exists a unitary $\mathcal{U}_{\text{dCG}}^{(\lambda)} : V_\lambda^d \otimes \mathbb{C}^d \rightarrow \bigoplus_\mu V_\mu^d$ that implements the isomorphism in [Equation \(70\)](#). In other words, for all $U \in U(d)$,

$$\mathcal{U}_{\text{dCG}}^{(\lambda)} \cdot (\nu_\lambda(U) \otimes U^*) \cdot \mathcal{U}_{\text{dCG}}^{(\lambda)\dagger} = \sum_{\substack{\mu = \lambda - \square \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \nu_\mu(U). \quad (71)$$

We observe that the choice of $\mathcal{U}_{\text{dCG}}^{(\lambda)}$ is not unique. Indeed, one can see from Schur's Lemma that any unitary of the form

$$\left(\sum_{\mu = \lambda - \square} \alpha^\mu \cdot |\mu\rangle\langle\mu| \otimes I_{V_\mu^d} \right) \cdot \mathcal{U}_{\text{dCG}}^{(\lambda)}$$

where α^μ is a complex scalar of unit norm will also satisfy [Equation \(71\)](#). The above transformation defines the dual Clebsch-Gordan coefficients.

Definition 12.4 (Dual Clebsch-Gordan coefficients). Let $\lambda \in \mathbb{Z}^d$, $T \in \text{SSYT}(\lambda, d)$, and $k \in [d]$. We can write

$$\mathcal{U}_{\text{dCG}}^{(\lambda)} \cdot |T\rangle \otimes |\bar{k}\rangle = \sum_{\mu = \lambda - \square} \sum_{T'} \langle \mu, T' | \mathcal{U}_{\text{dCG}}^{(\lambda)} | T, \bar{k} \rangle \cdot |\mu, T'\rangle.$$

Here, T' indexes a basis of the irrep corresponding to μ .² Then coefficients $\langle \mu, T' | \mathcal{U}_{\text{dCG}}^{(\lambda)} | T, \bar{k} \rangle$ are known as the *dual Clebsch-Gordan coefficients*. We will typically write $\langle T' | T, \bar{k} \rangle$ as shorthand for $\langle \mu, T' | \mathcal{U}_{\text{dCG}}^{(\lambda)} | T, \bar{k} \rangle$, leaving $\mu = \text{sh}(T')$ and $\mathcal{U}_{\text{dCG}}^{(\lambda)}$ implicit.

We are interested in the relationship between the Clebsch-Gordan and the dual Clebsch-Gordan coefficients when both λ and μ correspond to polynomial irreps. In particular, we use the following result from Equation (13) of Section 18.2.1 of [VK92].

Theorem 12.5 ([VK92]). *Let $\lambda, \mu \in \mathbb{Z}^d$ be Young diagrams with nonnegative row lengths such that $\mu = \lambda - \square$. Then there exists a unitary $\mathcal{U}_{\text{dCG}}^{(\lambda)}$ that performs the dual Clebsch-Gordan transform such that:*

$$\langle \mu, T' | \mathcal{U}_{\text{dCG}}^{(\lambda)} | T, \bar{k} \rangle = \sqrt{\frac{\dim(V_\mu^d)}{\dim(V_\lambda^d)}} \cdot \langle \lambda, T | \mathcal{U}_{\text{CG}}^{(\mu)} | T', k \rangle.$$

²We have already seen that for polynomial representations, T' can be thought of as an SSYT. For rational representations, there is also a connection to tableaux. However, we will only make use of this in the case when λ and μ are polynomial representations, and so we will not need the more general case. See, e.g. [Ngu24], for more details.

Part IV

The moments of the debiased Keyl's estimator

13 The first moment

In this section, we formally prove [Theorem 1.3](#) from the introduction. That is, we show that the debiased Keyl's algorithm, given in [Theorem 1.2](#), is an unbiased estimator. We start by developing some tools related to the one-step Clebsch-Gordan coefficients.

13.1 Technical lemmas concerning one-step Clebsch-Gordan coefficients

In this subsection, we collect some identities about the one-step Clebsch-Gordan coefficients that we will use in our proofs of the first and second moments of our estimator. Recall that for a Young diagram $\lambda = (\lambda_1, \dots, \lambda_d)$, we use $\lambda_{\leq j} = (\lambda_1, \dots, \lambda_j)$ to denote the truncation of this tableau to its first j rows.

Notation 13.1. Let λ be a Young diagram, and let $i \in [d]$. We write $D_{k \rightarrow i}^\lambda$ as shorthand for the ratio of dimensions

$$D_{k \rightarrow i}^\lambda := \frac{\dim(V_{\lambda_{\leq k} + e_i}^k)}{\dim(V_{\lambda_{\leq k}}^k)},$$

when $i \leq k \leq d$, and define $D_{k \rightarrow i}^\lambda = 0$ when $0 \leq k < i$.

Notation 13.2. Let λ be a Young diagram, and let $i \in [d]$. We write C_i^λ as shorthand for the quantity $C_i^\lambda := \lambda_i - i$, which is also the content of the rightmost box in row i in λ . We will omit λ when clear from context.

The first lemma relates the one-step Clebsch-Gordan coefficients to the dimensions of truncated tableaux.

Lemma 13.3. *Let λ be a Young diagram, and $k, i \in [d]$. Then*

$$|c_{k \rightarrow i}^\lambda|^2 = D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda, \quad (72)$$

and

$$(C_i - C_k) \cdot |c_{k \rightarrow i}^\lambda|^2 = D_{k-1 \rightarrow i}^\lambda. \quad (73)$$

These equations together immediately imply:

$$(C_i - C_k) \cdot D_{k \rightarrow i}^\lambda = (C_i - C_k + 1) \cdot D_{k-1 \rightarrow i}^\lambda. \quad (74)$$

Proof. The proof follows by direct calculation. If $k < i$, each side of both [Equations \(72\)](#) and [\(73\)](#) is equal to zero, so the equations hold. When $k \geq i$, recall the Clebsch-Gordan coefficient expression from [Theorem 10.13](#):

$$|c_{k \rightarrow i}^\lambda|^2 = \left(\prod_{j < i} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \right) \cdot \left(\prod_{j=i+1}^{k-1} \frac{(\lambda_i - \lambda_j) + (j - i) + 1}{(\lambda_i - \lambda_j) + (j - i)} \right) \cdot \frac{1}{(\lambda_i - \lambda_k) + (k - i) + \delta_{ik}}.$$

Now suppose $k > i$. We use the Weyl dimension formula from [Equation \(15\)](#) to write

$$\begin{aligned} \frac{\dim(V_{\lambda_{\leq k} + e_i}^k)}{\dim(V_{\lambda_{\leq k}}^k)} &= \left(\prod_{j < i} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \right) \cdot \left(\prod_{j=i+1}^k \frac{(\lambda_i - \lambda_j) + (j - i) + 1}{(\lambda_i - \lambda_j) + (j - i)} \right) \\ &= \left(\prod_{j < i} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \right) \cdot \left(\prod_{j=i+1}^{k-1} \frac{(\lambda_i - \lambda_j) + (j - i) + 1}{(\lambda_i - \lambda_j) + (j - i)} \right) \cdot \left(\frac{1}{(\lambda_i - \lambda_k) + (k - i)} + 1 \right) \\ &= |c_{k \rightarrow i}^\lambda|^2 + \frac{\dim(V_{\lambda_{\leq k-1} + e_i}^{k-1})}{\dim(V_{\lambda_{\leq k-1}}^{k-1})}. \end{aligned}$$

Thus, $D_{k \rightarrow i}^\lambda = |c_{k \rightarrow i}^\lambda|^2 + D_{k-1 \rightarrow i}^\lambda$. Moreover, the second equality above implies that for $k > i$,

$$D_{k \rightarrow i}^\lambda = D_{k-1 \rightarrow i}^\lambda \cdot \left(\frac{1}{C_i - C_k} + 1 \right) \implies (C_i - C_k) \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) = D_{k-1 \rightarrow i}^\lambda.$$

Lastly, suppose $k = i$. Then

$$\begin{aligned} \frac{\dim(V_{\lambda_{\leq k} + e_i}^k)}{\dim(V_{\lambda_{\leq k}}^k)} &= \prod_{j < i} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \\ &= \left(\prod_{j < i} \frac{(\lambda_j - \lambda_i) + (i - j) - 1}{(\lambda_j - \lambda_i) + (i - j)} \right) \cdot \frac{1}{(\lambda_i - \lambda_k) + (k - i) + \delta_{ik}} = |c_{k \rightarrow i}^\lambda|^2. \end{aligned}$$

Since $D_{i-1 \rightarrow i}^\lambda$ is defined to be 0, we conclude that $D_{i \rightarrow i}^\lambda = |c_{i \rightarrow i}^\lambda|^2 + D_{i-1 \rightarrow i}^\lambda$. Finally, both sides of Equation (73) are equal to zero when $i = k$. $\square$

Telescoping the above sum in Equation (72) gives the following corollary.

Corollary 13.4. *The one-step Clebsch-Gordan coefficients satisfy the following relation:*

$$\sum_{k=1}^s |c_{k \rightarrow i}^\lambda|^2 = D_{s \rightarrow i}^\lambda.$$

The next lemma shows that the one-step CG coefficient $c_{k \rightarrow i}^\lambda$, regarded as a function of k alone (i.e. for fixed λ , d , and i), is constant on the blocks of λ .

Lemma 13.5. *Let λ be a Young diagram, and let $\{k, k+1, \dots, \ell\}$ be a block of λ , i.e. $\lambda_k = \lambda_{k+1} = \dots = \lambda_\ell$. Then, for each $i \leq d$, the one-step Clebsch-Gordan coefficients that correspond to adding any value in $\{k, \dots, \ell\}$ satisfy:*

$$c_{k \rightarrow i}^\lambda = c_{k+1 \rightarrow i}^\lambda = \dots = c_{\ell \rightarrow i}^\lambda. \quad (75)$$

Proof. We first observe that if $i \in \{k, \dots, \ell\}$, then we must have $i = k$ if $\lambda + e_i$ is a valid Young diagram. If $k+1 \leq i \leq \ell$, then Equation (75) holds trivially as all coefficients are zero, since $\lambda + e_i$ is invalid. Furthermore, if $i > \ell$, even in the case where $\lambda + e_i$ is a valid Young diagram, $T_{m \rightarrow i}^\lambda$ is an invalid SSYT for all $m \in \{k, \dots, \ell\}$, since $m < i$. In this case too, Equation (75) holds trivially.

It remains to show that $c_{m+1 \rightarrow i}^\lambda = c_{m \rightarrow i}^\lambda$ for all $m \in \{k, \dots, \ell-1\}$ whenever $i \leq k$. We consider the case when $i = m$ (which happens when they are both equal to k) separately.

1. If $i = m = k$, then, from Theorem 10.13, we have

$$\begin{aligned} c_{k+1 \rightarrow i}^\lambda &= \left| \frac{1}{(\lambda_i - \lambda_{k+1}) + (k+1-i) + \delta_{i,k+1}} \cdot \prod_{j=1, j \neq i}^k \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} \\ &= \left| \frac{1}{(\lambda_i - \lambda_k) + (k+1-i)} \cdot \prod_{j=1, j \neq k}^k \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} \quad (\lambda_k = \lambda_{k+1}) \\ &= \left| \frac{1}{(\lambda_i - \lambda_k) + (k-i) + \delta_{ik}} \cdot \prod_{j=1, j \neq i}^{k-1} \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} = c_{k \rightarrow i}^\lambda. \end{aligned}$$

2. If $i < m$, then

$$\begin{aligned}
c_{m+1 \rightarrow i}^\lambda &= \left| \frac{1}{(\lambda_i - \lambda_{m+1}) + (m+1-i) + \delta_{i,m+1}} \cdot \prod_{j=1, j \neq i}^m \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} \\
&= \left| \frac{1}{(\lambda_i - \lambda_{m+1}) + (m+1-i)} \cdot \frac{(\lambda_m - \lambda_i) + (i-m) - 1}{(\lambda_m - \lambda_i) + (i-m)} \cdot \prod_{j=1, j \neq i}^{m-1} \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} \\
&= \left| \frac{1}{(\lambda_i - \lambda_m) + (m-i) + 1} \cdot \frac{(\lambda_i - \lambda_m) + (m-i) + 1}{(\lambda_i - \lambda_m) + (m-i)} \cdot \prod_{j=1, j \neq i}^{m-1} \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} \\
&= \left| \frac{1}{(\lambda_i - \lambda_m) + (m-i) + \delta_{im}} \cdot \prod_{j=1, j \neq i}^{m-1} \frac{(\lambda_j - \lambda_i) + (i-j) - 1}{(\lambda_j - \lambda_i) + (i-j)} \right|^{1/2} = c_{m \rightarrow i}^\lambda,
\end{aligned}$$

where we used the fact that $\lambda_m = \lambda_{m+1}$ in the third step. $\square$

13.2 Proof of the first moment

Theorem 13.6 (Theorem 1.3, restated). *Let $\hat{\rho}$ be the output of the debiased Keyl's algorithm when run on $\rho^{\otimes n}$. Then $\mathbf{E}[\hat{\rho}] = \rho$.*

Let us interpret what it means for a tomography algorithm to give an unbiased estimator of a density matrix ρ . A tomography algorithm is determined by a POVM $M = \{M_i\}$ on $(\mathbb{C}^d)^{\otimes n}$ and a corresponding set of possible outputs $\{\hat{\rho}_i\}$: the algorithm measures $\rho^{\otimes n}$ using M , and outputs $\hat{\rho}_i$ upon obtaining outcome M_i . If the tomography algorithm is unbiased, then

$$\sum_i \hat{\rho}_i \cdot \text{tr}(M_i \cdot \rho^{\otimes n}) = \rho.$$

Some elementary manipulations of the left-hand side give

$$\sum_i \hat{\rho}_i \cdot \text{tr}(M_i \cdot \rho^{\otimes n}) = \sum_i \text{tr}_{[n]}(M_i \otimes \hat{\rho}_i \cdot \rho^{\otimes n} \otimes I) = \text{tr}_{[n]} \left(\left(\sum_i M_i \otimes \hat{\rho}_i \right) \cdot \rho^{\otimes n} \otimes I \right) = \text{tr}_{[n]}(M_{\text{avg}}^{(1)} \cdot \rho^{\otimes n} \otimes I),$$

where we define

$$M_{\text{avg}}^{(1)} := \sum_i M_i \otimes \hat{\rho}_i.$$

So our goal is to satisfy $\text{tr}_{[n]}(M_{\text{avg}}^{(1)} \cdot \rho^{\otimes n} \otimes I) = \rho$. This would happen, for example, if $M_{\text{avg}}^{(1)} = (1, n+1)$, where $(1, n+1) := \mathcal{P}(1, n+1)$ and $\mathcal{P}(\cdot)$ is the representation of the symmetric group from Theorem 2.46, since

$$\text{tr}_{[n]}((1, n+1) \cdot \rho^{\otimes n} \otimes I) = \text{tr}_{[2]}((1, 2) \cdot \rho \otimes I) \cdot \text{tr}(\rho)^{n-1} = \rho \cdot \text{tr}(\rho)^{n-1} = \rho.$$

Similarly, it would also be true if $M_{\text{avg}}^{(1)} = (k, n+1)$ for any $k \in [n]$, and also, more generally, if

$$M_{\text{avg}}^{(1)} = \sum_{k=1}^n \alpha_k \cdot (k, n+1),$$

for any constants with $\alpha_1 + \dots + \alpha_n = 1$. One might then expect that any unbiased algorithm that treats its n registers symmetrically should satisfy

$$M_{\text{avg}}^{(1)} = \frac{1}{n} \cdot \left((1, n+1) + \dots + (n, n+1) \right) = \frac{1}{n} \cdot X_{n+1}, \quad (76)$$

where X_{n+1} is the Jucys-Murphy element as in Theorem 2.35. We show the debiased Keyl's algorithm satisfies Equation (76) below in Theorem 13.7, from which Theorem 1.3 follows as an immediate corollary.

Lemma 13.7. Let $M = \{M_{\lambda,U}\}$ be the POVM that corresponds to Keyl's algorithm. Consider a legal estimator, which outputs $\hat{\rho}_{\lambda,U} := U \cdot f(\lambda) \cdot U^\dagger / n$ when the POVM element indexed by λ and U is obtained. Let $M_{\text{avg}}^{(1)} = \sum_{\lambda,U} M_{\lambda,U} \otimes \hat{\rho}_{\lambda,U}$, where $\hat{\rho}$. Then, for any legal estimator,

$$M_{\text{avg}}^{(1)} = \frac{1}{n} \cdot \left((1, n+1) + \cdots + (n, n+1) \right) = \frac{1}{n} \cdot X_{n+1}.$$

Proof. The POVM used by any legal estimator is the same as in the original Keyl's algorithm. The POVM elements are indexed by $\lambda \vdash n$ and $U \in U(d)$, and, from Equation (22), is given by:

$$M_{\lambda,U} = \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(\dim(V_\lambda^d) \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(U) \cdot |T^\lambda\rangle\langle T^\lambda| \cdot \nu_\lambda(U^\dagger) \cdot dU \right) \cdot \mathcal{U}_{\text{SW}}^{(n)}.$$

Upon receiving λ and U , the state we output is $\hat{\rho}_{\lambda,U} = U \cdot f(\lambda) \cdot U^\dagger / n$. We can therefore rewrite $M_{\text{avg}}^{(1)}$ as:

$$\begin{aligned} M_{\text{avg}}^{(1)} &= \frac{1}{n} \sum_{\lambda} \int_U M_{\lambda,U} \otimes \hat{\rho}_{\lambda,U} \\ &= \frac{1}{n} \sum_{\lambda} \dim(V_\lambda^d) \int_U \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(U) \cdot |T^\lambda\rangle\langle T^\lambda| \cdot \nu_\lambda(U^\dagger) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \otimes U f(\lambda) U^\dagger \cdot dU \\ &= \frac{1}{n} \sum_{\lambda} \dim(V_\lambda^d) \sum_S \sum_{k=1}^d f(\lambda)_k \int_U \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes \nu_\lambda(U) \cdot |T^\lambda\rangle\langle T^\lambda| \cdot \nu_\lambda(U^\dagger) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \otimes U |k\rangle\langle k| U^\dagger \cdot dU. \end{aligned} \tag{77}$$

Here, $S \in \text{SYT}(\lambda)$, and $k \in [d]$. We can manipulate the integrand as follows:

$$\begin{aligned} &\mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes \nu_\lambda(U) \cdot |T^\lambda\rangle\langle T^\lambda| \cdot \nu_\lambda(U^\dagger) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \otimes U |k\rangle\langle k| U^\dagger \\ &= U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes |T^\lambda\rangle\langle T^\lambda| \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot U^{\dagger, \otimes n} \otimes U |k\rangle\langle k| U^\dagger \\ &= U^{\otimes(n+1)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes |T^\lambda\rangle\langle T^\lambda| \otimes |k\rangle\langle k| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+1)}. \end{aligned}$$

Here, we are using the notation

$$\mathcal{U} := \mathcal{U}_{\text{SW}}^{(n)} \otimes I,$$

as shorthand. We have also used

$$U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(U).$$

We proceed by using the Clebsch-Gordan transform to rewrite the state inside parentheses using:

$$\begin{aligned} \mathcal{U}_{\text{R}}^{(n+1)} \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \cdot |\lambda\rangle \otimes |S\rangle \otimes |T^\lambda\rangle \otimes |k\rangle &= \mathcal{U}_{\text{R}}^{(n+1)} \cdot \sum_{i=1}^d c_{k \rightarrow i}^\lambda |\lambda\rangle \otimes |S\rangle \otimes |\lambda + e_i\rangle \otimes |T_{k \rightarrow i}^\lambda\rangle \\ &= \sum_{i=1}^d c_{k \rightarrow i}^\lambda |\lambda + e_i\rangle \otimes |S_i\rangle \otimes |T_{k \rightarrow i}^\lambda\rangle. \end{aligned}$$

Here, S_i is the SYT obtained from S , by adding a box labeled $n+1$ to the end of the i -th row. This is well defined for all λ such that $\lambda + e_i$ is a valid Young diagram. Recalling that the Clebsch-Gordan coefficients are real, and our definition

$$\mathcal{U}_{\text{SW}}^{(n+1)} := \mathcal{U}_{\text{R}}^{(n+1)} \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \cdot (\mathcal{U}_{\text{SW}}^{(n)} \otimes I) = \mathcal{U}_{\text{R}}^{(n+1)} \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \cdot \mathcal{U},$$

we have

$$\begin{aligned}
& U^{\otimes(n+1)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes |T_\lambda\rangle\langle T_\lambda| \otimes |k\rangle\langle k| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+1)} \\
&= U^{\otimes(n+1)} \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\sum_{i,j=1}^d c_{k \rightarrow i}^\lambda c_{k \rightarrow j}^\lambda |\lambda + e_i\rangle\langle\lambda + e_j| \otimes |S_i\rangle\langle S_j| \otimes |T_{k \rightarrow i}^\lambda\rangle\langle T_{k \rightarrow j}^\lambda| \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)} \cdot U^{\dagger, \otimes(n+1)} \\
&= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\sum_{i,j=1}^d c_{k \rightarrow i}^\lambda c_{k \rightarrow j}^\lambda |\lambda + e_i\rangle\langle\lambda + e_j| \otimes |S_i\rangle\langle S_j| \otimes \nu_{\lambda+e_i}(U) |T_{k \rightarrow i}^\lambda\rangle\langle T_{k \rightarrow j}^\lambda| \nu_{\lambda+e_j}(U)^\dagger \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)}.
\end{aligned}$$

After substituting this back into Equation (77), the integral over unitaries in $M_{\text{avg}}^{(1)}$ can then be pushed under the sum and onto the unitary register. The resulting integral can be evaluated using Schur's lemma:

$$\int_U \nu_{\lambda+e_i}(U) |T_{k \rightarrow i}^\lambda\rangle\langle T_{k \rightarrow j}^\lambda| \nu_{\lambda+e_j}(U)^\dagger \cdot dU = \frac{1}{\dim(V_{\lambda+e_i}^d)} \cdot I_{\dim(V_{\lambda+e_i}^d)} \cdot \delta_{ij},$$

since the integral is an intertwining operator between the irreducible representations $\nu_{\lambda+e_i}$ and $\nu_{\lambda+e_j}$, which are non-isomorphic unless $i = j$. Therefore,

$$\begin{aligned}
& \int_U U^{\otimes(n+1)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes |T_\lambda\rangle\langle T_\lambda| \otimes |k\rangle\langle k| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+1)} dU \\
&= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\sum_{i=1}^d \frac{|c_{k \rightarrow i}^\lambda|^2}{\dim(V_{\lambda+e_i}^d)} \cdot |\lambda + e_i\rangle\langle\lambda + e_i| \otimes |S_i\rangle\langle S_i| \otimes I_{\dim(V_{\lambda+e_i}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)}.
\end{aligned}$$

Our expression for $M_{\text{avg}}^{(1)}$ thus becomes:

$$\begin{aligned}
M_{\text{avg}}^{(1)} &= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\frac{1}{n} \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k=1}^d f(\lambda)_k \sum_S \sum_{i=1}^d \frac{|c_{k \rightarrow i}^\lambda|^2}{\dim(V_{\lambda+e_i}^d)} \cdot |\lambda + e_i\rangle\langle\lambda + e_i| \otimes |S_i\rangle\langle S_i| \otimes I_{\dim(V_{\lambda+e_i}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)} \\
&= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\frac{1}{n} \sum_{\lambda} \sum_S \sum_{i=1}^d \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda+e_i}^d)} \cdot \left(\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 \right) \cdot |\lambda + e_i\rangle\langle\lambda + e_i| \otimes |S_i\rangle\langle S_i| \otimes I_{\dim(V_{\lambda+e_i}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)}.
\end{aligned}$$

We have shown so far that $M_{\text{avg}}^{(1)}$ is diagonal in the $(n+1)$ -copy Schur basis, and the expression below gives each diagonal entry:

$$\frac{1}{n} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda+e_i}^d)} \cdot \left(\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 \right).$$

We prove in Theorem 13.9 below, using our lemmas for the one-step Clebsch-Gordan coefficients, that

$$\frac{1}{n} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda+e_i}^d)} \cdot \left(\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 \right) = \frac{1}{n} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda+e_i}^d)} \cdot (\lambda_i + 1 - i) \cdot \frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_{\lambda}^d)} = \frac{1}{n} \cdot (\lambda_i + 1 - i). \quad (78)$$

Finally,

$$\begin{aligned}
M_{\text{avg}}^{(1)} &= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\frac{1}{n} \sum_{\lambda} \sum_S \sum_{i=1}^d (\lambda_i + 1 - i) \cdot |\lambda + e_i\rangle\langle\lambda + e_i| \otimes |S_i\rangle\langle S_i| \otimes I_{\dim(V_{\lambda+e_i}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)} \\
&= \mathcal{U}_{\text{SW}}^{(n+1)\dagger} \cdot \left(\frac{1}{n} \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} \sum_{S' \in \text{SYT}(\mu)} \text{cont}_{S'}(n+1) \cdot |\mu\rangle\langle\mu| \otimes |S'\rangle\langle S'| \otimes I_{\dim(V_{\mu}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)}.
\end{aligned}$$

Comparing with Equation (19), we conclude:

$$M_{\text{avg}}^{(1)} = \frac{1}{n} \cdot X_{n+1}.$$

□

To complete the proof of our first moment, it suffices to prove that the following identity from [Equation \(78\)](#) holds, for all legal Young diagram transformations f :

$$\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 = (\lambda_i + 1 - i) \cdot D_{d \rightarrow i}^\lambda.$$

We first show in [Theorem 13.8](#) below that this identity holds in the special case when f corresponds to the staircase transformation. Then we use the property that the one-step Clebsch-Gordan coefficients are constant on the blocks of a diagram to extend this result to all f 's in [Theorem 13.9](#).

Lemma 13.8. *Let λ be a Young diagram, and $\lambda^{\uparrow\uparrow}$ be the outcome of the staircase transformation on λ . Then*

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k \rightarrow i}^\lambda|^2 = (\lambda_i + 1 - i) \cdot D_{d \rightarrow i}^\lambda. \quad (79)$$

Proof. Fix λ and i . We prove the lemma by induction on d . Note first though that $\lambda_k^{\uparrow\uparrow} = \lambda_k + d + 1 - 2k$ implicitly depends on d . When we would like to emphasize d , we will write $\lambda_k^{\uparrow\uparrow}(d)$. We are also taking the convention that $\lambda_k = 0$ for all $k > \ell(\lambda)$.

The base case is $d = 1$. If $i = 1$, then

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow}(d=1) \cdot |c_{k \rightarrow i}^\lambda|^2 = \lambda_1^{\uparrow\uparrow}(d=1) \cdot |c_{1 \rightarrow 1}^\lambda|^2 = \lambda_1 \cdot |c_{1 \rightarrow 1}^\lambda|^2 = \lambda_1 \cdot D_{1 \rightarrow i}^\lambda = (\lambda_i + 1 - i) \cdot D_{d \rightarrow i}^\lambda,$$

where we have used [Theorem 13.4](#) in the second-last step and the fact that $d = i = 1$ in the final step. If instead $i > 1$, then both sides of [Equation \(79\)](#) vanish.

Now assume the lemma holds for $d - 1$. We use the fact that $\lambda_k^{\uparrow\uparrow}(d) = \lambda_k^{\uparrow\uparrow}(d-1) + 1$ and the inductive hypothesis, to get:

$$\begin{aligned} \sum_{k=1}^d \lambda_k^{\uparrow\uparrow}(d) \cdot |c_{k \rightarrow i}^\lambda|^2 &= \left(\sum_{k=1}^{d-1} (\lambda_k^{\uparrow\uparrow}(d-1) + 1) \cdot |c_{k \rightarrow i}^\lambda|^2 \right) + \lambda_d^{\uparrow\uparrow}(d) \cdot |c_{d \rightarrow i}^\lambda|^2 \\ &= (\lambda_i + 1 - i) \cdot D_{d-1 \rightarrow i}^\lambda + \left(\sum_{k=1}^{d-1} |c_{k \rightarrow i}^\lambda|^2 \right) + (\lambda_d - d + 1) \cdot |c_{d \rightarrow i}^\lambda|^2. \end{aligned} \quad (80)$$

We now rewrite everything in terms of C 's and D 's using [Theorem 13.4](#):

$$\begin{aligned} (80) &= (C_i + 1) \cdot D_{d-1 \rightarrow i}^\lambda + D_{d-1 \rightarrow i}^\lambda + (C_d + 1) \cdot (D_{d \rightarrow i}^\lambda - D_{d-1 \rightarrow i}^\lambda) \\ &= (C_i - C_d + 1) \cdot D_{d-1 \rightarrow i}^\lambda + (C_d + 1) \cdot D_{d \rightarrow i}^\lambda. \end{aligned} \quad (81)$$

Finally, from [Theorem 13.3](#), we have

$$D_{d-1 \rightarrow i}^\lambda = (C_i - C_d) \cdot |c_{d \rightarrow i}^\lambda|^2 = (C_i - C_d) \cdot (D_{d \rightarrow i}^\lambda - D_{d-1 \rightarrow i}^\lambda) \implies (C_i - C_d + 1) \cdot D_{d-1 \rightarrow i}^\lambda = (C_i - C_d) \cdot D_{d \rightarrow i}^\lambda,$$

and substituting this gives us

$$(81) = (C_i - C_d) \cdot D_{d \rightarrow i}^\lambda + (C_d + 1) \cdot D_{d \rightarrow i}^\lambda = (C_i + 1) \cdot D_{d \rightarrow i}^\lambda. \quad \square$$

We are now ready to extend the above identity to all legal transformations f .

Lemma 13.9. *Let λ be a Young diagram, and f be a legal transformation of Young diagrams. Then*

$$\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 = (\lambda_i + 1 - i) \cdot D_{d \rightarrow i}^\lambda. \quad (82)$$

Proof. For a generic legal transformation f , we use [Theorem 13.5](#), which states that the one-step Clebsch-Gordan coefficients $c_{k \rightarrow i}^\lambda$ are constant (as functions of k) on indices in the same blocks of λ . Write $c_{k \rightarrow i}^\lambda = c_{B \rightarrow i}^\lambda$, for all $k \in B$, for B a block of λ . Then, for any estimator, we have

$$\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 = \sum_{\text{blocks } B} |c_{B \rightarrow i}^\lambda|^2 \sum_{k \in B} f(\lambda)_k = \sum_{\text{blocks } B} |c_{B \rightarrow i}^\lambda|^2 \sum_{k \in B} \lambda_k^\uparrow,$$

where the second step is by the definition of a legal estimator, and recall that $\lambda^\uparrow$ refers to the box donation transformation from [Theorem 1.1](#). Thus, the left-hand side of [Equation \(82\)](#) is the same for *any* legal transformation, and so

$$\sum_{k=1}^d f(\lambda)_k \cdot |c_{k \rightarrow i}^\lambda|^2 = \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k \rightarrow i}^\lambda|^2 = (\lambda_i + 1 - i) \cdot D_{d \rightarrow i}^\lambda. \quad \square$$

14 The second moment

In this section, we formally prove [Theorem 1.4](#) from the introduction, which obtains an expression for the second moment of the output $\hat{\rho}$ of the debiased Keyl's algorithm.

14.1 Strategy overview

Since our proof for the second moment is more involved than the first moment, we start with a high-level overview of the main ingredients.

A central part of the first moment proof was the matrix $M_{\text{avg}}^{(1)}$ and its relation to the expected output of the debiased Keyl's algorithm on n copies of ρ :

$$\mathbf{E}[\hat{\rho}] = \text{tr}_{[n]}(M_{\text{avg}}^{(1)} \cdot \rho^{\otimes n} \otimes I).$$

Using Schur-Weyl duality and the Clebsch-Gordan transform, we showed that $M_{\text{avg}}^{(1)}$ is diagonal in the Schur basis, and that it is equal to the Jucys-Murphy element $\frac{1}{n} \cdot X_{n+1}$, which implies that our estimator is unbiased. That is, $\mathbf{E}[\hat{\rho}] = \rho$.

For the second moment, we define the matrix $M_{\text{avg}}^{(2)}$, which captures the second moment of the debiased Keyl's estimator on n copies of the state ρ . We define, for an estimator that measures with POVM $\{M_i\}$ and outputs $\hat{\rho}_i$ upon obtaining M_i :

$$M_{\text{avg}}^{(2)} := \sum_i M_i \otimes \hat{\rho}_i \otimes \hat{\rho}_i.$$

Then

$$\mathbf{E}[\hat{\rho} \otimes \hat{\rho}] = \text{tr}_{[n]}(M_{\text{avg}}^{(2)} \cdot \rho^{\otimes n} \otimes I \otimes I).$$

We show in [Section 14.4](#) the Main Lemma ([Theorem 14.15](#)) for the second moment, which says that the matrix $M_{\text{avg}}^{(2)}$ satisfies the equation

$$M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP} = \frac{1}{n^2} \cdot X_{n+1} X_{n+2}. \quad (83)$$

Throughout this section we use SWAP for $\text{SWAP}_{n+1, n+2}$, and $M_{\text{corr}}^{(2)}$ is a “correction” term for which we have an exact formula. Later in this subsection, we use [Equation \(83\)](#) and the exact form of $M_{\text{corr}}^{(2)}$ to prove [Theorem 1.4](#), which is restated below.

The remainder of this section is dedicated to proving the Main Lemma. In [Section 14.2](#), we use the Clebsch-Gordan transform to deduce that $M_{\text{avg}}^{(2)}$ and $M_{\text{corr}}^{(2)}$ have the same block-diagonal structure in the Schur basis as SWAP, with submatrices of size 2×2 and 1×1 on their main diagonal. The entries of these submatrices are given by expressions involving sums of the two-step Clebsch-Gordan coefficients. The Jucys-Murphy element $X_{n+1} X_{n+2}$ is a diagonal matrix in the Schur basis whose entries are related to the

contents of boxes of standard Young tableaux. Therefore, to prove Equation (83), it suffices to compute the exact value of $M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP}$ for each submatrix in the Schur basis, and show that it matches the values of $\frac{1}{n^2} \cdot X_{n+1} X_{n+2}$. This reduces to proving an expression for the diagonal terms in each submatrix, the Diagonal Expression Lemma (Theorem 14.17), and this we do in Section 14.5.

For this purpose, we develop several tools for the two-step Clebsch-Gordan coefficients in Section 14.3. The two most important ones are Theorem 14.10, which gives a closed-form formula for the partial sums of the two-step Clebsch-Gordan coefficients, and Theorem 14.14, which relates the two-step Clebsch-Gordan coefficients and the blocks of the underlying Young diagram. Both of these technical lemmas are extensions of Theorems 13.4 and 13.5 from the one-step to the two-step coefficients, respectively.

Theorem 14.1 (Theorem 1.4, restated). *Let $\hat{\rho}$ be the output of the debiased Keyl's algorithm when run on $\rho^{\otimes n}$. Then*

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

Here, Lower_ρ refers to the “lower order terms” and has the following characterization: it can be written as a positive linear combination of terms of the form $(P \otimes P) \cdot \text{SWAP}$, where each $P \in \mathbb{C}^{d \times d}$ is a Hermitian, positive semidefinite matrix.

Proof. Recall from the proof of unbiasedness that we interpreted our tomography algorithm as measuring $\rho^{\otimes n}$ using the POVM $M = \{M_i\}$. Upon obtaining outcome M_i , the algorithm outputs $\hat{\rho}_i$. It holds that

$$\begin{aligned} \mathbf{E}[\hat{\rho} \otimes \hat{\rho}] &= \sum_i \hat{\rho}_i \otimes \hat{\rho}_i \cdot \text{tr}(M_i \cdot \rho^{\otimes n}) \\ &= \sum_i \text{tr}_{[n]} \left(M_i \otimes \hat{\rho}_i \otimes \hat{\rho}_i \cdot \rho^{\otimes n} \otimes I \otimes I \right) \\ &= \text{tr}_{[n]} \left(\sum_i M_i \otimes \hat{\rho}_i \otimes \hat{\rho}_i \cdot \rho^{\otimes n} \otimes I \otimes I \right) \\ &= \text{tr}_{[n]} \left(M_{\text{avg}}^{(2)} \cdot \rho^{\otimes n} \otimes I \otimes I \right), \end{aligned}$$

where we define

$$M_{\text{avg}}^{(2)} := \sum_i M_i \otimes \hat{\rho}_i \otimes \hat{\rho}_i.$$

For the debiased Keyl's algorithm, the POVM elements are labeled by $\lambda \vdash n$ and $U \in U(d)$. See Equation (22). Similar calculations as the first moment section provide the following expression for $M_{\text{avg}}^{(2)}$:

$$\begin{aligned} M_{\text{avg}}^{(2)} &= \sum_\lambda \int_U M_{\lambda,U} \otimes U \hat{\alpha} U^\dagger \otimes U \hat{\alpha} U^\dagger \\ &= \sum_\lambda \dim(V_\lambda^d) \int_U \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \nu_\lambda(U) |T^\lambda\rangle\langle T^\lambda| \nu_\lambda(U)^\dagger \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \otimes U \hat{\alpha} U^\dagger \otimes U \hat{\alpha} U^\dagger \cdot dU \\ &= \sum_\lambda \dim(V_\lambda^d) \int_U U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^\lambda\rangle\langle T^\lambda| \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot U^{\dagger, \otimes n} \otimes U \hat{\alpha} U^\dagger \otimes U \hat{\alpha} U^\dagger \cdot dU \\ &= \sum_\lambda \dim(V_\lambda^d) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^\lambda\rangle\langle T^\lambda| \otimes \hat{\alpha} \otimes \hat{\alpha} \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \sum_\lambda \dim(V_\lambda^d) \sum_{k, \ell=1}^d \hat{\alpha}_k \hat{\alpha}_\ell \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^\lambda\rangle\langle T^\lambda| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \frac{1}{n^2} \sum_\lambda \dim(V_\lambda^d) \sum_{k, \ell} \lambda_k^\dagger \lambda_\ell^\dagger \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^\lambda\rangle\langle T^\lambda| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU, \end{aligned}$$

where here $\mathcal{U} := \mathcal{U}_{\text{SW}}^{(n)} \otimes I \otimes I$. As mentioned above (Equation (83)), we show in Theorem 14.15 that

$$M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP} = \frac{1}{n^2} \cdot X_{n+1} X_{n+2},$$

where $M_{\text{corr}}^{(2)}$ is a “correction” term that is given by

$$\begin{aligned} M_{\text{corr}}^{(2)} &:= \frac{1}{n} \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k,\ell} \hat{\alpha}_{\max(k,\ell)} \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k,\ell} \lambda_{\max(k,\ell)}^{\dagger} \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU. \end{aligned}$$

Note that this is identical to our expression for $M_{\text{avg}}^{(2)}$ above except with the $\lambda_k^{\dagger} \lambda_{\ell}^{\dagger}$ part replaced by a $\lambda_{\max(k,\ell)}^{\dagger}$.

From [Theorem 14.2](#) below, we can write $M_{\text{corr}}^{(2)}$ in the following way, with $\{m_j^{\lambda}\}_{j \in [d]}$ a set of nonnegative integers, and $\Pi_{\leq j} := \sum_{i \leq j} |i\rangle\langle i|$ the projector onto the first j computational basis vectors:

$$\begin{aligned} M_{\text{corr}}^{(2)} &= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{j=1}^d m_j^{\lambda} \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes \Pi_{\leq j} \otimes \Pi_{\leq j} \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &\quad - \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \cdot \ell(\lambda) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes I \otimes I \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU. \end{aligned}$$

We will call the first term above the negative correction $M_{\text{corr},-}^{(2)}$ and the second term as the positive correction $M_{\text{corr},+}^{(2)}$, from the sign of the term they contribute to $\mathbf{E}[\hat{\rho} \otimes \hat{\rho}]$. Therefore

$$\begin{aligned} \mathbf{E}[\hat{\rho} \otimes \hat{\rho}] &= \text{tr}_{[n]}(M_{\text{avg}}^{(2)} \cdot \rho^{\otimes n} \otimes I \otimes I) \\ &= \frac{1}{n^2} \text{tr}_{[n]}(X_{n+1} X_{n+2} \cdot \rho^{\otimes n} \otimes I \otimes I) + \text{tr}_{[n]}(M_{\text{corr},+}^{(2)} \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I) \\ &\quad - \text{tr}_{[n]}(M_{\text{corr},-}^{(2)} \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I). \end{aligned}$$

Let us compute each term separately. For the first term, we expand the product of Jucys-Murphy elements as

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

Thus

$$\begin{aligned} &\text{tr}_{[n]}(X_{n+1} X_{n+2} \cdot \rho^{\otimes n} \otimes I \otimes I) \\ &= \sum_{i \neq j} \text{tr}_{[n]}((i, n+1)(j, n+2) \cdot \rho^{\otimes n} \otimes I \otimes I) + \sum_{i=1}^n \text{tr}_{[n]}((i, n+1, n+2) \cdot \rho^{\otimes n} \otimes I \otimes I) \\ &\quad + \sum_{i=1}^n \text{tr}_{[n]}((i, n+2, n+1) \cdot \rho^{\otimes n} \otimes I \otimes I) \\ &= n(n-1) \cdot (\rho \otimes \rho) + n \cdot (\rho \otimes I + I \otimes \rho) \cdot \text{SWAP}. \end{aligned}$$

Now we compute the second term. First, we simplify $M_{\text{corr},+}^{(2)}$ as follows

$$\begin{aligned}
M_{\text{corr},+}^{(2)} &= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \cdot \ell(\lambda) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes I \otimes I \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\
&= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \cdot \ell(\lambda) \cdot \left(\int_U U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \right) U^{\dagger, \otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot dU \right) \otimes I \otimes I \\
&= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \cdot \ell(\lambda) \cdot \left(\mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \int_U \nu_{\lambda}(U) |T^{\lambda}\rangle\langle T^{\lambda}| \nu_{\lambda}(U)^{\dagger} \cdot dU \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \right) \otimes I \otimes I \\
&= \frac{1}{n^2} \sum_{\lambda} \dim(V_{\lambda}^d) \cdot \ell(\lambda) \cdot \left(\mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes \frac{I_{\dim(V_{\lambda}^d)}}{\dim(V_{\lambda}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \right) \otimes I \otimes I \quad (\text{Schur's Lemma}) \\
&= \frac{1}{n^2} \sum_{\lambda} \ell(\lambda) \cdot \Pi_{\lambda} \otimes I \otimes I. \tag{Equation (18)}
\end{aligned}$$

Hence

$$\begin{aligned}
\text{tr}_{[n]} (M_{\text{corr},+}^{(2)} \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I) &= \frac{1}{n^2} \sum_{\lambda} \ell(\lambda) \cdot \text{tr}_{[n]} \left(\Pi_{\lambda} \otimes I \otimes I \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I \right) \\
&= \frac{1}{n^2} \left(\sum_{\lambda} \ell(\lambda) \cdot \text{tr}(\Pi_{\lambda} \cdot \rho^{\otimes n}) \right) \cdot \text{SWAP} \\
&= \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \cdot \text{SWAP}.
\end{aligned}$$

Finally, we prove that the last term corresponds to Lower_{ρ} . In particular, $M_{\text{corr},-}^{(2)}$ can be written as a positive linear combination of the following terms:

$$\int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes \Pi_{\leq j} \otimes \Pi_{\leq j} \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU.$$

Hence $\text{tr}_{[n]} (M_{\text{corr},-}^{(2)} \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I)$ can also be written as a positive linear combination of the terms:

$$\begin{aligned}
&\text{tr}_{[n]} \left(\int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes \Pi_{\leq j} \otimes \Pi_{\leq j} \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot \text{SWAP} \cdot \rho^{\otimes n} \otimes I \otimes I \cdot dU \right) \\
&= \int_U \text{tr}_{[n]} \left(U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes \Pi_{\leq j} \otimes \Pi_{\leq j} \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot \rho^{\otimes n} \otimes I \otimes I \right) \cdot \text{SWAP} \cdot dU \\
&= \int_U \text{tr} \left(U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot U^{\dagger, \otimes n} \cdot \rho^{\otimes n} \right) \cdot (U \Pi_{\leq j} U^{\dagger} \otimes U \Pi_{\leq j} U^{\dagger}) \cdot \text{SWAP} \cdot dU.
\end{aligned}$$

The trace inside the integrand is nonnegative since $\text{tr}(AB) \geq 0$ if A and B are positive semidefinite. We conclude that each term itself can be further written as a positive linear combination of terms $(P \otimes P) \cdot \text{SWAP}$, where each P is of the form $U \Pi_{\leq j} U^{\dagger}$, and is thus a Hermitian, positive semidefinite matrix. $\square$

Lemma 14.2. *Let $n \geq 1$ and $\lambda \vdash n$. It holds that*

$$\sum_{k, \ell} \lambda_{\max(k, \ell)}^{\dagger} \cdot |k\rangle\langle k| \otimes |\ell\rangle\langle \ell| = \left(\sum_{j=1}^d m_j^{\lambda} \cdot \Pi_{\leq j} \otimes \Pi_{\leq j} \right) - \ell(\lambda) \cdot I \otimes I,$$

where $\Pi_{\leq j} := \sum_{i \leq j} |i\rangle\langle i|$ is the projector onto the first j computational basis vectors, and $\{m_j^{\lambda}\}_{j \in [d]}$ is a set of nonnegative integers defined as follows. Regard λ as a tuple of length d , and let $B_a = \{i_a, \dots, j_a\}$ be the a -th block of λ , for $a \in [A]$, where A is the total number of blocks. The coefficients $\{m_j^{\lambda}\}_{j \in [d]}$ are given by:

$$m_j^{\lambda} = \begin{cases} \lambda_{B_a} - \lambda_{B_{a+1}} + |B_a| + |B_{a+1}|, & j = j_a \text{ for some } a \in [A-1], \\ \lambda_{B_A} + |B_A|, & j = j_A \text{ and } \lambda_A > 0 \\ 0, & \text{otherwise.} \end{cases}$$

Proof. We have

$$\sum_{j=1}^d m_j^\lambda \cdot \Pi_{\leq j} \otimes \Pi_{\leq j} = \sum_{j=1}^d m_j^\lambda \cdot \left(\sum_{k=1}^j |k\rangle\langle k| \right) \otimes \left(\sum_{\ell=1}^j |\ell\rangle\langle \ell| \right) = \sum_{k,\ell=1}^d \left(\sum_{j=\max(k,\ell)}^d m_j^\lambda \right) \cdot |k\rangle\langle k| \otimes |\ell\rangle\langle \ell|.$$

First, we consider the case where $\lambda_A = 0$. In this case, $\ell(\lambda) = \sum_{a'=1}^{A-1} |B_{a'}|$. As a subcase, suppose $\max(k, \ell) \in B_a$, with $a < A$, then

$$\begin{aligned} \sum_{j=\max(k,\ell)}^d m_j^\lambda &= \lambda_{B_a} + |B_a| + 2 \left(\sum_{a'=a+1}^{A-1} |B_{a'}| \right) + |B_A| \\ &= \lambda_{B_a} + |B_a| + 2 \left(\sum_{a'=a+1}^{A-1} |B_{a'}| \right) + |B_A| + \left(\ell(\lambda) - \sum_{a'=1}^{A-1} |B_{a'}| \right) \\ &= \left(\lambda_{B_a} - \sum_{a' < a} |B_{a'}| + \sum_{a' > a} |B_{a'}| \right) + \ell(\lambda) \\ &= \lambda_{B_a}^\uparrow + \ell(\lambda). \end{aligned}$$

In the other subcase, with $\max(k, \ell) \in B_A$, we have $\lambda_{B_A} = 0$ and

$$\sum_{j=\max(k,\ell)}^d m_j^\lambda = 0 = \lambda_{B_A} + \left(\ell(\lambda) - \sum_{a'=1}^{A-1} |B_{a'}| \right) = \left(\lambda_{B_A} - \sum_{a'=1}^{A-1} |B_{a'}| \right) + \ell(\lambda) = \lambda_{B_A}^\uparrow + \ell(\lambda).$$

Thus, when $\lambda_A = 0$, we have

$$\begin{aligned} \left(\sum_{j=1}^d m_j^\lambda \cdot \Pi_{\leq j} \otimes \Pi_{\leq j} \right) - \ell(\lambda) \cdot I \otimes I &= \sum_{k,\ell} \left(\sum_{j=\max(k,\ell)}^d m_j^\lambda - \ell(\lambda) \right) \cdot |k\rangle\langle k| \otimes |\ell\rangle\langle \ell| \\ &= \sum_{k,\ell} \lambda_{\max(k,\ell)}^\uparrow \cdot |k\rangle\langle k| \otimes |\ell\rangle\langle \ell|. \end{aligned}$$

The second case, when $\lambda_A > 0$, is similar, but now $\ell(\lambda) = \sum_{a'=1}^A |B_{a'}|$. If we add a formal $(A+1)$ -th block with $|B_{A+1}| = 0$ and $\lambda_{A+1} = 0$, then this case reduces to the first subcase of the $\lambda_A = 0$ case above. $\square$

14.2 Block-diagonal expressions in the Schur basis

Having defined the expressions $M_{\text{avg}}^{(2)}$ and $M_{\text{corr}}^{(2)}$, in this section, we show that they can be written in the Schur basis as block-diagonal matrices, using the Clebsch-Gordan transform.

Notation 14.3. For convenience, we will adopt the following notation and terminology.

- We will abbreviate $\lambda + e_i + e_j$ as λ_{ij} , the Young diagram obtained when we add boxes to λ in rows i and j . Note that $\lambda_{ij} = \lambda_{ji}$.
- If $\lambda + e_i$ and $\lambda + e_i + e_j$ are valid Young diagrams, and S is an SYT of shape λ , we let S_{ij} denote the SYT obtained by starting with S , and adding $n+1$ to the end of the i -th row, and $n+2$ to the end of the j -th row. Note that $S_{ij} \neq S_{ji}$, unless $i = j$. It can also be the case that S_{ij} is valid, whereas S_{ji} is not, or vice versa. For example, if $\lambda = (1, 1)$, then we may insert into the first row, and then the second, but not in the reverse order.
- For S an SYT of shape λ , we can use the pair $(S, \{i, j\})$ to index the subspace of $\text{Sp}_{\lambda_{ij}}$ spanned by the valid SYTs among $|S_{ij}\rangle$ and $|S_{ji}\rangle$.
- We refer to a block with both S_{ij} and S_{ji} valid as *swappable*, since the new boxes in rows i and j may be inserted in either order. Otherwise, we call a block *non-swappable*. There are two non-swappable subcases: if $j = i$, then we refer to this as the *horizontal* case; and if $j = i + 1$, we refer to this as the *vertical* case. See [Figure 22](#) for examples.

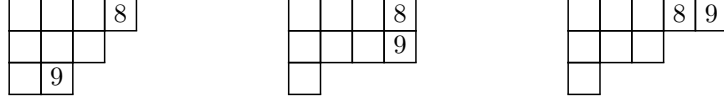


Figure 22: In the above examples, we have an SYT of shape $\lambda = (3, 3, 1)$ (not all of its values are visible because they are irrelevant), and we add two new boxes in rows i and j with values 8 and 9 respectively. The three resulting SYTs correspond to the cases when i, j are swappable, vertical, and horizontal, respectively.

Lemma 14.4. *Any expression M of the form*

$$M = \sum_{\lambda} \dim(V_{\lambda}^d) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU,$$

can be written in the Schur basis as a block-diagonal matrix:

$$\mathcal{U}_{\text{SW}}^{(n+2)} \cdot M \cdot \mathcal{U}_{\text{SW}}^{(n+2)\dagger} = \sum_{\lambda \vdash n} \sum_{i \leq j} |\lambda_{ij}\rangle\langle\lambda_{ij}| \otimes \sum_{S \in \lambda} |S, \{i, j\}\rangle\langle S, \{i, j\}| \otimes M_{\{i, j\}}^S \otimes I_{\dim(V_{\lambda_{ij}}^d)}.$$

Here, the vectors $|S, \{i, j\}\rangle$ index into the blocks (swappable, horizontal and vertical), and each $M_{\{i, j\}}^S$ is a 2×2 matrix if S_{ij} and S_{ji} are valid and distinct SYTs supported on the block $\text{span}(|S_{ij}\rangle, |S_{ji}\rangle)$ (i.e. the swappable case), and is a 1×1 matrix otherwise supported on the block $\text{span}(|S_{ij}\rangle)$ (i.e. the non-swappable case). The diagonal entries of $M_{\{i, j\}}^S$ are equal to

$$M_{ij, ij}^S := \langle S_{ij} | M_{\{i, j\}}^S | S_{ij} \rangle = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot (|a_{k\ell \rightarrow ij}^{\lambda}|^2 + |b_{k\ell \rightarrow ij}^{\lambda}|^2).$$

Whenever $M_{\{i, j\}}^S$ is a 2×2 matrix, its off-diagonal entries are equal to

$$M_{ij, ji}^S := \langle S_{ij} | M_{\{i, j\}}^S | S_{ji} \rangle = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot (a_{k\ell \rightarrow ij}^{\lambda} b_{k\ell \rightarrow ji}^{\lambda} + a_{k\ell \rightarrow ji}^{\lambda} b_{k\ell \rightarrow ij}^{\lambda} + |a_{kk \rightarrow ij}^{\lambda}|^2 \cdot \delta_{k\ell}).$$

Proof. We will prove the first part of the lemma by applying the Clebsch-Gordan transform to M . For brevity, we will use $\lambda_{ij} = \lambda + e_i + e_j$ to denote the Young diagram when we add boxes in rows i and j to λ , and analogously, S_{ij} (though note $S_{ij} \neq S_{ji}$). The Clebsch-Gordan transform gives

$$\begin{aligned} & |\lambda\rangle \otimes |S\rangle \otimes |T^{\lambda}\rangle \otimes |k\rangle \otimes |\ell\rangle \\ &= \left(\mathcal{U}_{\text{R}}^{(n+2)} \cdot \mathcal{U}_{\text{CG}}^{(n+2)} \cdot \mathcal{U}_{\text{R}}^{(n+1)} \otimes I \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \otimes I \right)^{\dagger} \cdot \sum_{i, j} |\lambda_{ij}\rangle \otimes |S_{ij}\rangle \otimes \left(a_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ij}^{\lambda}\rangle + b_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ji}^{\lambda}\rangle \right) \\ &= \left(\mathcal{U}_{\text{SW}}^{(n+2)} \cdot \mathcal{U}^{\dagger} \right)^{\dagger} \cdot \sum_{i, j} |\lambda_{ij}\rangle \otimes |S_{ij}\rangle \otimes \left(a_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ij}^{\lambda}\rangle + b_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ji}^{\lambda}\rangle \right). \end{aligned}$$

Thus

$$\begin{aligned} & |\lambda\rangle\langle\lambda| \otimes |S\rangle\langle S| \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \\ &\cong \sum_{i, j} \sum_{i', j'} |\lambda_{ij}\rangle\langle\lambda_{i'j'}| \otimes |S_{ij}\rangle\langle S_{i'j'}| \otimes \left(a_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ij}^{\lambda}\rangle + b_{k\ell \rightarrow ij}^{\lambda} |T_{k\ell \rightarrow ji}^{\lambda}\rangle \right) \left(a_{k\ell \rightarrow i'j'}^{\lambda} \langle T_{k\ell \rightarrow i'j'}^{\lambda}| + b_{k\ell \rightarrow i'j'}^{\lambda} \langle T_{k\ell \rightarrow j'i'}^{\lambda}| \right), \end{aligned} \tag{84}$$

where $\cong$ indicates equality up to conjugation by $\mathcal{U}_{\text{SW}}^{(n+2)} \cdot \mathcal{U}^{\dagger}$. When we compute the integral of Equation (84), due to Schur-Weyl duality (Theorem 2.47), we will get a linear combination of expressions of the form:

$$\begin{aligned} & \int_U U^{\otimes(n+2)} \cdot \mathcal{U}_{\text{SW}}^{(n+2)\dagger} \cdot \left(|\lambda_{ij}\rangle\langle\lambda_{i'j'}| \otimes |S_{ij}\rangle\langle S_{i'j'}| \otimes |T_{k\ell \rightarrow pq}^{\lambda}\rangle\langle T_{k\ell \rightarrow p'q'}^{\lambda}| \right) \cdot \mathcal{U}_{\text{SW}}^{(n+2)} \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \mathcal{U}_{\text{SW}}^{(n+2)\dagger} \cdot \left(|\lambda_{ij}\rangle\langle\lambda_{i'j'}| \otimes |S_{ij}\rangle\langle S_{i'j'}| \otimes \int_U \nu_{\lambda_{pq}}(U) |T_{k\ell \rightarrow pq}^{\lambda}\rangle\langle T_{k\ell \rightarrow p'q'}^{\lambda}| \nu_{\lambda_{p'q'}}(U)^{\dagger} \cdot dU \right) \cdot \mathcal{U}_{\text{SW}}^{(n+2)}, \end{aligned}$$

where p, q and p', q' are such that $\{p, q\} = \{i, j\}$ and $\{p', q'\} = \{i', j'\}$. The above integral defines an intertwining operator between the irreps $\nu_{\lambda_{pq}}$ and $\nu_{\lambda_{p'q'}}$, and so can be computed using Schur's lemma. In particular, if the Young diagrams λ_{pq} and $\lambda_{p'q'}$ are different, then the integral is zero. This happens if and only if $\{p, q\} \neq \{p', q'\}$. Otherwise, when $\{p, q\} = \{p', q'\}$ (so that $\{i, j\} = \{p, q\} = \{p', q'\} = \{i', j'\}$) the integral evaluates to a multiple of the identity $c \cdot I_{\dim(V_{\lambda_{ij}}^d)}$, where the constant c is given by [Theorem 2.16](#):

$$c = \frac{1}{\dim(V_{\lambda_{ij}}^d)} \cdot \text{tr} \left(\int_U \nu_{\lambda_{ij}}(U) |T_{k\ell \rightarrow pq}^\lambda \rangle \langle T_{k\ell \rightarrow p'q'}^\lambda| \nu_{\lambda_{ij}}(U)^\dagger \cdot dU \right) = \frac{\langle T_{k\ell \rightarrow p'q'}^\lambda | T_{k\ell \rightarrow pq}^\lambda \rangle}{\dim(V_{\lambda_{ij}}^d)}.$$

Thus, in the Schur basis, M is a linear combination of terms of the form

$$|\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes |S_{ij}\rangle \langle S_{i'j'}| \otimes I_{\dim(V_{\lambda_{ij}}^d)}.$$

Since $\{i, j\} = \{i', j'\}$, we see that each standard Young tableau $|S_{ij}\rangle$ can only possibly be paired with itself and the tableau $|S_{ji}\rangle$ that we obtain if we swap i and j , provided this is a valid tableau. Thus, we obtain the block structure claimed in the first part of the lemma.

Let us now compute the diagonal entries of $M_{\{i,j\}}^S$. A term with $|S_{ij}\rangle \langle S_{ij}|$ in the second register arises from terms as in [Equation \(84\)](#) with $(i, j) = (i', j')$:

$$|\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes |S_{ij}\rangle \langle S_{ij}| \otimes \left(a_{k\ell \rightarrow ij}^\lambda |T_{k\ell \rightarrow ij}^\lambda\rangle + b_{k\ell \rightarrow ij}^\lambda |T_{k\ell \rightarrow ji}^\lambda\rangle \right) \left(a_{k\ell \rightarrow ij}^\lambda \langle T_{k\ell \rightarrow ij}^\lambda| + b_{k\ell \rightarrow ij}^\lambda \langle T_{k\ell \rightarrow ji}^\lambda| \right).$$

With such a term, after integrating over the unitary group and using Schur's lemma, we are left with terms of the form

$$\frac{1}{\dim(V_{\lambda_{ij}}^d)} \cdot |\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes |S_{ij}\rangle \langle S_{ij}| \otimes \left(|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 \right) \cdot I_{\dim(V_{\lambda_{ij}}^d)}.$$

This is because $T_{k\ell \rightarrow ij}^\lambda$ and $T_{k\ell \rightarrow ji}^\lambda$ are distinct SSYTs when $k \neq \ell$, and when $k = \ell$, we have that $b_{k\ell \rightarrow ij}^\lambda = 0$.

We proceed to the off-diagonal entries of $M_{\{i,j\}}^S$, which occur only in the swappable case, when necessarily $i \neq j$. A term with $|S_{ji}\rangle \langle S_{ji}|$ in the second register arises from terms as in [Equation \(84\)](#) with $(i, j) = (j', i')$:

$$|\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes |S_{ji}\rangle \langle S_{ji}| \otimes \left(a_{k\ell \rightarrow ij}^\lambda |T_{k\ell \rightarrow ij}^\lambda\rangle + b_{k\ell \rightarrow ij}^\lambda |T_{k\ell \rightarrow ji}^\lambda\rangle \right) \left(a_{k\ell \rightarrow ji}^\lambda \langle T_{k\ell \rightarrow ji}^\lambda| + b_{k\ell \rightarrow ji}^\lambda \langle T_{k\ell \rightarrow ij}^\lambda| \right).$$

With such a term, after integrating over the unitary group and using Schur's lemma, we are left with

$$\frac{1}{\dim(V_{\lambda_{ij}}^d)} \cdot |\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes |S_{ji}\rangle \langle S_{ji}| \otimes \left(a_{k\ell \rightarrow ij}^\lambda b_{k\ell \rightarrow ji}^\lambda + a_{k\ell \rightarrow ji}^\lambda b_{k\ell \rightarrow ij}^\lambda + |a_{k\ell \rightarrow ij}^\lambda|^2 \cdot \delta_{k\ell} \right) \cdot I_{\dim(V_{\lambda_{ij}}^d)}.$$

This is because $T_{k\ell \rightarrow ij}^\lambda$ and $T_{k\ell \rightarrow ji}^\lambda$ are distinct SSYTs when $k \neq \ell$, but not when $k = \ell$. $\square$

In addition to the block-diagonal structure of the expressions from [Theorem 14.4](#), we show below that each 2×2 block matrix is a linear combination of the identity and SWAP when the expression is real and symmetric in k and ℓ .

Notation 14.5. Following [Theorem 2.34](#), SWAP is block-diagonal in the Schur basis, and preserves horizontal, vertical, and swappable blocks. In particular, we can write:

$$\begin{aligned} \text{SWAP} &= \mathcal{P}((n+1, n+2)) \\ &= \mathcal{U}_{\text{SW}}^{(n+2)\dagger} \cdot \left(\sum_{\lambda \vdash n} \sum_{i \leq j} |\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes \kappa_{\lambda_{ij}}(n+1, n+2) \otimes I_{\dim(V_{\lambda_{ij}}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+2)} \\ &= \mathcal{U}_{\text{SW}}^{(n+2)\dagger} \cdot \left(\sum_{\lambda \vdash n} \sum_{i \leq j} |\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes \sum_{S \in \text{SYT}(\lambda)} |S, \{i, j\}\rangle \langle S, \{i, j\}| \otimes \text{SWAP}_{\{i,j\}}^S \otimes I_{\dim(V_{\lambda_{ij}}^d)} \right) \cdot \mathcal{U}_{\text{SW}}^{(n+2)}. \end{aligned}$$

Moreover, $\text{SWAP}_{\{i,j\}}^S$ is 1 in the horizontal case, -1 in the vertical case, and in the swappable case,

$$\text{SWAP}_{\{i,j\}}^S = \begin{pmatrix} \frac{1}{\Delta_{ji}} & \sqrt{1 - \frac{1}{\Delta_{ji}^2}} \\ \sqrt{1 - \frac{1}{\Delta_{ji}^2}} & -\frac{1}{\Delta_{ji}} \end{pmatrix}.$$

Here, $\Delta_{ji} := \Delta_S(n+1) = \text{cont}_S(n+2) - \text{cont}_S(n+1)$, is the difference of contents between the two new cells in rows j and i . We will also use Δ_{ji} in the horizontal and vertical cases, where $\Delta_{ji} = +1$ and $\Delta_{ji} = -1$ respectively. Note that in the swappable case, $|\Delta_{ji}| \geq 2$.

Lemma 14.6. *Any expression M of the form*

$$M = \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell=1}^d f(k, \ell) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU, \quad (85)$$

with f is a real symmetric function in k, ℓ , can be written in the Schur basis as a block-diagonal matrix

$$\mathcal{U}_{\text{SW}}^{(n+2)} \cdot M \cdot \mathcal{U}_{\text{SW}}^{(n+2)\dagger} = \sum_{\lambda \vdash n} |\lambda_{ij}\rangle\langle\lambda_{ij}| \otimes \sum_{S \in \lambda} \sum_{i \leq j} |S, \{i, j\}\rangle\langle S, \{i, j\}| \otimes M_{\{i, j\}}^S \otimes I_{\dim(V_{\lambda+e_i+e_j}^d)},$$

where each $M_{\{i, j\}}^S$ of dimension 2 must have the form $x_{\{i, j\}}^S \cdot I + y_{\{i, j\}}^S \cdot \text{SWAP}_{\{i, j\}}^S$, where $x_{\{i, j\}}^S$ and $y_{\{i, j\}}^S$ are scalars.

Proof. Observe that if $f(k, \ell) = f(\ell, k)$, then M commutes with SWAP, since M is equal to

$$\begin{aligned} & \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell} f(k, \ell) \int_U U^{\otimes(n+2)} \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell} f(k, \ell) \int_U \left(U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot U^{\dagger, \otimes n} \right) \otimes \left(U |k\rangle\langle k| U^{\dagger} \otimes U |\ell\rangle\langle\ell| U^{\dagger} \right) \cdot dU. \end{aligned}$$

Then $\text{SWAP} \cdot M \cdot \text{SWAP}$ becomes

$$\begin{aligned} & \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell} f(k, \ell) \int_U \left(U^{\otimes n} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \right) \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot U^{\dagger, \otimes n} \right) \otimes \left(U |\ell\rangle\langle\ell| U^{\dagger} \otimes U |k\rangle\langle k| U^{\dagger} \right) \cdot dU \\ &= \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell} f(k, \ell) \int_U U^{\otimes(n+2)} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |\ell\rangle\langle\ell| \otimes |k\rangle\langle k| \right) \cdot U^{\dagger, \otimes(n+2)} \cdot dU \\ &= \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell} f(\ell, k) \int_U U^{\otimes(n+2)} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot U^{\dagger, \otimes(n+2)} \cdot dU, \end{aligned}$$

which is equal to M if $f(k, \ell) = f(\ell, k)$. So, if f is a real symmetric function, then M and SWAP commute. In this case, M and SWAP are simultaneously diagonalizable as they are both Hermitian. By writing both of them in the Schur basis, where they have the same block diagonal structure, we conclude that their principal submatrices that correspond to each $(S, \{i, j\})$ block must also commute.

For a swappable block $(S, \{i, j\})$, the projectors on the two eigenspaces of $\text{SWAP}_{\{i, j\}}^S$ are $\frac{1}{2}(I \pm \text{SWAP}_{\{i, j\}}^S)$. Thus, we conclude that any $M_{\{i, j\}}^S$ can be written as a linear combination of $\frac{1}{2}(I \pm \text{SWAP}_{\{i, j\}}^S)$, or equivalently, as a linear combination of I and $\text{SWAP}_{\{i, j\}}^S$. $\square$

14.3 Technical lemmas concerning two-step Clebsch-Gordan coefficients

In this subsection, we collect some identities about the two-step Clebsch-Gordan coefficients that we will use in our proof for the second moment of our estimator. Recall from [Theorem 13.1](#) that we use $D_{k \rightarrow i}^{\lambda}$ to refer to the ratio of dimensions $\dim(V_{\lambda_{\leq k}+e_i}^k) / \dim(V_{\lambda_{\leq k}}^k)$ when $i \leq k \leq d$, and $D_{k \rightarrow i}^{\lambda} = 0$ otherwise.

Notation 14.7. We will often consider expressions of the form $|a_{k \rightarrow i, j}^{\lambda}|^2 + |b_{k \rightarrow i, j}^{\lambda}|^2$. To simplify notation, we introduce the following shorthand:

$$|c_{k \rightarrow i, j}^{\lambda}|^2 := |a_{k \rightarrow i, j}^{\lambda}|^2 + |b_{k \rightarrow i, j}^{\lambda}|^2.$$

Notation 14.8. We will use the following shorthand for the partial sum of two-step CG coefficients:

$$F(s, t) := \sum_{k=1}^s \sum_{\ell=1}^t |c_{k \rightarrow i, j}^{\lambda}|^2.$$

Notation 14.9. We will also use the shorthand

$$D_{s \rightarrow ij}^\lambda := D_{s \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} = \frac{\dim(V_{\lambda \leq s+e_i}^s)}{\dim(V_{\lambda \leq s}^s)} \cdot \frac{\dim(V_{\lambda \leq s+e_i+e_j}^s)}{\dim(V_{\lambda \leq s+e_i}^s)} = \frac{\dim(V_{\lambda \leq s+e_i+e_j}^s)}{\dim(V_{\lambda \leq s}^s)}.$$

Note that $D_{s \rightarrow ij}^\lambda = D_{s \rightarrow ji}^\lambda$.

Lemma 14.10 (Partial sums of two-step Clebsch-Gordan coefficients). *Let λ be a Young diagram, and $i, j, k, \ell \in [d]$. Then*

$$F(s, t) = \sum_{k=1}^s \sum_{\ell=1}^t |c_{k\ell \rightarrow ij}^\lambda|^2 = \begin{cases} D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i} & s \leq t, \\ D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} (D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} - D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j}) & s \geq t. \end{cases}$$

Proof. We prove the lemma in three steps.

Step 1: First, we prove the result in the case when $s \leq t$, using the dual Clebsch-Gordan transform. We observe that $|c_{k\ell \rightarrow ij}^\lambda|^2$ can be written as:

$$\begin{aligned} |c_{k\ell \rightarrow ij}^\lambda|^2 &= |a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 \\ &= |\langle T^\lambda, k | T_{k \rightarrow i}^\lambda \rangle \cdot \langle T_{k \rightarrow i}^\lambda, \ell | T_{k\ell \rightarrow ij}^\lambda \rangle|^2 + |\langle T^\lambda, k | T_{k \rightarrow i}^\lambda \rangle \cdot \langle T_{k \rightarrow i}^\lambda, \ell | T_{k\ell \rightarrow ji}^\lambda \rangle|^2 \\ &= |\langle T^\lambda, k | T_{k \rightarrow i}^\lambda \rangle|^2 \cdot (|\langle T_{k \rightarrow i}^\lambda, \ell | T_{k\ell \rightarrow ij}^\lambda \rangle|^2 + |\langle T_{k \rightarrow i}^\lambda, \ell | T_{k\ell \rightarrow ji}^\lambda \rangle|^2) \\ &= |c_{k \rightarrow i}^\lambda|^2 \sum_{T'} |\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle|^2, \end{aligned}$$

where $T' \in \text{SSYT}(\lambda_{ij}, d)$. The last equality follows because the only SSYTs T' for which $\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle$ is nonzero are $T_{k\ell \rightarrow ij}^\lambda$ and $T_{k\ell \rightarrow ji}^\lambda$. Now $F(s, t)$ can be written as

$$\sum_{k=1}^s \sum_{\ell=1}^t |c_{k\ell \rightarrow ij}^\lambda|^2 = \sum_{k=1}^s |c_{k \rightarrow i}^\lambda|^2 \sum_{\ell=1}^t \sum_{T'} |\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle|^2.$$

We proceed by showing that for $k \leq t$, the value of the inner summation does not depend on k , and has a quite simple expression:

$$\sum_{\ell=1}^t \sum_{T'} |\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle|^2 = D_{t \rightarrow j}^{\lambda+e_i}. \quad (86)$$

From the description of the Clebsch-Gordan insertion algorithm in [Theorem 10.9](#), since $k, \ell \leq t$, no letter greater than t is bumped during the insertion of the letter ℓ in $T_{k \rightarrow i}^\lambda$ to obtain T' . Thus, the Clebsch-Gordan coefficient $\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle$ is a product of scalar factors that only depend on the restrictions of $T_{k \rightarrow i}^\lambda$ and T' to the alphabet $[t]$. We conclude that in the case of $k, \ell \leq t$, it holds that:

$$\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle = \langle (T_{k \rightarrow i}^\lambda)^{[t]}, \ell | (T')^{[t]} \rangle.$$

This gives an alternative expression for the left-hand side of [Equation \(86\)](#):

$$\sum_{\ell=1}^t \sum_{T'} |\langle T_{k \rightarrow i}^\lambda, \ell | T' \rangle|^2 = \sum_{\ell=1}^t \sum_{T''} |\langle (T_{k \rightarrow i}^\lambda)^{[t]}, \ell | T'' \rangle|^2.$$

The variable T'' on the right-hand side is iterating over the SSYTs of shape $(\lambda_{ij})_{\leq t}$, and thus has height at most t . From [Theorem 12.5](#), we can rewrite the above coefficient using the dual Clebsch-Gordan coefficients:

$$\sum_{\ell=1}^t \sum_{T''} |\langle (T_{k \rightarrow i}^\lambda)^{[t]}, \ell | T'' \rangle|^2 = \frac{\dim(V_{\lambda_{ij}}^t)}{\dim(V_{\lambda+e_i}^t)} \sum_{\ell=1}^t \sum_{T''} |\langle (T_{k \rightarrow i}^\lambda)^{[t]} | T'', \bar{\ell} \rangle|^2.$$

We observe that since ℓ iterates over $[t]$, then $\bar{\ell}$ iterates over all basis elements of the space $V_{-\square}^t$. Hence

$$\begin{aligned}
\frac{\dim(V_{\lambda_{ij}}^t)}{\dim(V_{\lambda+e_i}^t)} \sum_{\ell=1}^t \sum_{T''} |\langle (T_{k \rightarrow i}^\lambda)^{[t]} | T'', \bar{\ell} \rangle|^2 &= D_{t \rightarrow j}^{\lambda+e_i} \sum_{\ell=1}^t \sum_{T''} \langle (T_{k \rightarrow i}^\lambda)^{[t]} | T'', \bar{\ell} \rangle \cdot \langle T'', \bar{\ell} | (T_{k \rightarrow i}^\lambda)^{[t]} \rangle \quad (\text{Theorem 13.1}) \\
&= D_{t \rightarrow j}^{\lambda+e_i} \langle (T_{k \rightarrow i}^\lambda)^{[t]} | \left(\sum_{\ell=1}^t \sum_{T''} |T'', \bar{\ell}\rangle \langle T'', \bar{\ell}| \right) | (T_{k \rightarrow i}^\lambda)^{[t]} \rangle \\
&= D_{t \rightarrow j}^{\lambda+e_i} \cdot \langle (T_{k \rightarrow i}^\lambda)^{[t]} | (I_{V_{\lambda_{ij}}^t} \otimes I_{V_{-\square}^t}) | (T_{k \rightarrow i}^\lambda)^{[t]} \rangle \\
&= D_{t \rightarrow j}^{\lambda+e_i} \cdot \langle (T_{k \rightarrow i}^\lambda)^{[t]} | \left(\bigoplus_{\mu=\lambda_{ij}-\square} I_{V_\mu^t} \right) | (T_{k \rightarrow i}^\lambda)^{[t]} \rangle \quad (\text{Equation (70)}) \\
&= D_{t \rightarrow j}^{\lambda+e_i}.
\end{aligned}$$

The last equation follows because the space $V_{\lambda+e_i}^t$ appears with unit multiplicity in the direct sum, and $|(T_{k \rightarrow i}^\lambda)^{[t]} \rangle$ is a vector of unit norm in this space. We conclude that when $s \leq t$:

$$\begin{aligned}
\sum_{k=1}^s \sum_{\ell=1}^t |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{k=1}^s |c_{k \rightarrow i}^\lambda|^2 \sum_{\ell=1}^t \sum_{T''} |\langle T_{k \rightarrow i}^\lambda, \ell | T'' \rangle|^2 \\
&= \sum_{k=1}^s |c_{k \rightarrow i}^\lambda|^2 \cdot D_{t \rightarrow j}^{\lambda+e_i} \quad (\text{Equation (86)}) \\
&= D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i}. \quad (\text{Theorem 13.4})
\end{aligned}$$

Step 2: Next, we note that $F(s, t)$ is related to certain entries of a particular operator in the Schur basis. Specifically, consider the operator M given by

$$M_{st} := \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k, \ell=1}^d f_{st}(k, \ell) \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^\dagger \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^\lambda\rangle\langle T^\lambda| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU,$$

for f_{st} given by

$$f_{st}(k, \ell) := \begin{cases} 1 & k \leq s, \text{ and } \ell \leq t, \\ 0 & \text{otherwise.} \end{cases}$$

Previous results, [Theorems 14.4](#) and [14.6](#), have characterized such operators in the Schur basis. Specifically, M_{st} is block-diagonal in the blocks of SWAP. Fix S and a pair of indices $\{i, j\}$ with $i \leq j$. By [Theorem 14.4](#) and linearity, the $i \leq j$ entry is given by:

$$(M_{st})_{ij, ij}^S = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k=1}^s \sum_{\ell=1}^t (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot F(s, t). \quad (87)$$

The other diagonal entry, if it exists, is given by swapping i and j in the above expression (noting that F itself generally also has implicit dependence on i and j). The off-diagonal entries, if they exist, are given by

$$(M_{st})_{ij, ji}^S = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k=1}^s \sum_{\ell=1}^t (a_{k\ell \rightarrow ij}^\lambda b_{k\ell \rightarrow ji}^\lambda + a_{k\ell \rightarrow ji}^\lambda b_{k\ell \rightarrow ij}^\lambda + |a_{kk \rightarrow ij}^\lambda|^2 \cdot \delta_{k\ell}). \quad (88)$$

We will now restrict to the case we already understand from the previous step: $s \leq t$. We have two cases:

- (i) In the non-swappable case, [Equation \(87\)](#) is the sole entry of the 1×1 matrix, and, by the previous step, is equal to

$$(M_{st})_{ij, ij}^S = \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i}.$$

(ii) In the swappable case, we have a 2×2 matrix. The diagonal entries are likewise given by

$$(M_{st})_{ij,ij}^S = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i}, \quad (M_{st})_{ji,ji}^S = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot D_{s \rightarrow j}^\lambda \cdot D_{t \rightarrow i}^{\lambda+e_j}. \quad (89)$$

The off-diagonal entries, we claim, are zero.

We show this first for the case when $s = t$. In this case, f_{st} is symmetric, and hence M_{st} commutes with SWAP. Thus $(M_{st})_{\{i,j\}}^S = x \cdot I + y \cdot \text{SWAP}_{\{i,j\}}^S$, for some coefficients $x, y \in \mathbb{C}$. However, when $s = t$, we also have that the two diagonal entries given in Equation (89) are equal, and hence $y = 0$. Thus, the off-diagonal entries are zero in this case.

We now extend this to the case $s < t$. Abbreviate

$$g(k, \ell) := a_{k\ell \rightarrow ij}^\lambda b_{k\ell \rightarrow ji}^\lambda + a_{k\ell \rightarrow ji}^\lambda b_{k\ell \rightarrow ij}^\lambda + |a_{kk \rightarrow ij}^\lambda|^2 \cdot \delta_{k\ell}.$$

Note that $g(k, \ell) = 0$ for $k < \ell$, because $\delta_{k\ell} = 0$, and because no bumping can occur when CG inserting ℓ after k , so that $b_{k\ell \rightarrow ij}^\lambda = b_{k\ell \rightarrow ji}^\lambda = 0$. From the $s = t$ case, we also know that, for all s ,

$$\sum_{k=1}^s \sum_{\ell=1}^s g(k, \ell) = 0,$$

since this sum is proportional to the off-diagonal entries, by Equation (88). Thus, the off-diagonal entries are

$$(M_{st})_{ij,ji}^S = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k=1}^s \sum_{\ell=1}^t g(k, \ell) = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \left(\sum_{k=1}^s \sum_{\ell=1}^s g(k, \ell) + \sum_{k=1}^s \sum_{\ell=s+1}^t g(k, \ell) \right) = 0,$$

since the first sum vanishes, and each entry of the second sum vanishes. Thus, the off-diagonal entries are zero for $s < t$ too.

Thus, in the case where $s \leq t$ and i and j are swappable, we have

$$(M_{st})_{\{i,j\}}^S = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \begin{pmatrix} D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i} & 0 \\ 0 & D_{s \rightarrow j}^\lambda \cdot D_{t \rightarrow i}^{\lambda+e_j} \end{pmatrix}. \quad (90)$$

Step 3: In the last step, we show the result for $s \geq t$. We begin by noting that since

$$\text{SWAP} \cdot \left(\sum_{k=1}^t \sum_{\ell=1}^s |k\rangle\langle k| \otimes |\ell\rangle\langle \ell| \right) \cdot \text{SWAP} = \sum_{k=1}^t \sum_{\ell=1}^s |\ell\rangle\langle \ell| \otimes |k\rangle\langle k| = \sum_{k=1}^s \sum_{\ell=1}^t |k\rangle\langle k| \otimes |\ell\rangle\langle \ell|,$$

we have $M_{st} = \text{SWAP} \cdot M_{ts} \cdot \text{SWAP}$. Thus, in the block determined by S and $\{i, j\}$, we have the identity:

$$(M_{st})_{\{i,j\}}^S = \text{SWAP}_{\{i,j\}}^S \cdot (M_{ts})_{\{i,j\}}^S \cdot \text{SWAP}_{\{i,j\}}^S. \quad (91)$$

Now, take $t \leq s$. We consider two cases:

(i) In the non-swappable case, the matrices in Equation (91) are 1×1 , and hence all commute. So,

$$(M_{st})_{\{i,j\}}^S = (M_{ts})_{\{i,j\}}^S \implies \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot F(s, t) = \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot F(t, s). \quad (\text{Equation (87)})$$

Therefore,

$$F(s, t) = F(t, s) = D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} = D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} (D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} - D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j}),$$

where the second step holds since $t \leq s$, and the third step holds since $\Delta_{ji}^2 = 1$.

(ii) In the swappable case, we can explicitly compute the right-hand side of Equation (91) in the $s \geq t$ case using Equation (90). We get:

$$\begin{aligned}
& (M_{st})_{\{i,j\}}^S \\
&= \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \begin{pmatrix} \frac{1}{\Delta_{ji}} & \sqrt{1 - \frac{1}{\Delta_{ji}^2}} \\ \sqrt{1 - \frac{1}{\Delta_{ji}^2}} & -\frac{1}{\Delta_{ji}} \end{pmatrix} \cdot \begin{pmatrix} D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} & 0 \\ 0 & D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} \end{pmatrix} \cdot \begin{pmatrix} \frac{1}{\Delta_{ji}} & \sqrt{1 - \frac{1}{\Delta_{ji}^2}} \\ \sqrt{1 - \frac{1}{\Delta_{ji}^2}} & -\frac{1}{\Delta_{ji}} \end{pmatrix} \\
&= \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \begin{pmatrix} D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} (D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} - D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j}) & * \\ * & * \end{pmatrix},
\end{aligned}$$

where the asterisks represent entries we will not need. Finally, comparing this computation with Equation (87) gives:

$$F(s, t) = D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} (D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i} - D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j}),$$

as claimed. $\square$

There are many useful corollaries of Theorem 14.10. We use the following corollaries in our proof of the Diagonal Expression Lemma, Theorem 14.17 below, which itself is important for proving the Main Lemma, Theorem 14.15.

Corollary 14.11. $M_{\text{corr}}^{(2)}$ is diagonal in the Schur basis.

Proof. Fix a 2×2 block indexed by S and $\{i, j\}$, with $i < j$. From Theorem 14.6, it suffices to show the two diagonal entries of $M_{\text{corr}}^{(2)}$ in this block are equal. We start by considering the following sum:

$$\begin{aligned}
\sum_{k=1}^d \sum_{\ell=1}^d \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{k=1}^d \lambda_k^\uparrow \cdot \left(\sum_{k', \ell=1}^k |c_{k'\ell \rightarrow ij}^\lambda|^2 - \sum_{k', \ell=1}^{k-1} |c_{k'\ell \rightarrow ij}^\lambda|^2 \right) \\
&= \sum_{k=1}^d \lambda_k^\uparrow \cdot \left(D_{k \rightarrow i}^\lambda \cdot D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow i}^\lambda \cdot D_{k-1 \rightarrow j}^{\lambda+e_i} \right) \quad (\text{Theorem 14.10}) \\
&= \sum_{k=1}^d \lambda_k^\uparrow \cdot \left(D_{k \rightarrow ij}^\lambda - D_{k-1 \rightarrow ij}^\lambda \right). \quad (\text{Theorem 14.9})
\end{aligned}$$

Note that the remaining sum is symmetric in exchanging $i \leftrightarrow j$, by Theorem 14.9. That is,

$$\begin{aligned}
\sum_{k=1}^d \sum_{\ell=1}^d \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{k=1}^d \lambda_k^\uparrow \cdot \left(D_{k \rightarrow ij}^\lambda - D_{k-1 \rightarrow ij}^\lambda \right) \\
&= \sum_{k=1}^d \lambda_k^\uparrow \cdot \left(D_{k \rightarrow ji}^\lambda - D_{k-1 \rightarrow ji}^\lambda \right) = \sum_{k=1}^d \sum_{\ell=1}^d \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{k\ell \rightarrow ji}^\lambda|^2.
\end{aligned}$$

Finally, from the expression for the diagonal entries given in Theorem 14.4,

$$M_{ij, ij}^S = \frac{1}{n} \cdot \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k, \ell=1}^d \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = \frac{1}{n} \cdot \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ji}}^d)} \sum_{k, \ell=1}^d \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{k\ell \rightarrow ji}^\lambda|^2 = M_{ji, ji}^S. \quad \square$$

Corollary 14.12. We have the following equations:

- (i) $\sum_{k, \ell=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1) \cdot D_{d \rightarrow ij}^\lambda.$
- (ii) $\sum_{k, \ell=1}^d \left(\lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda.$

Proof. Our general strategy for showing both equations will be to rewrite it in terms of partial sums of two-step CG coefficients, and then to use [Theorem 14.10](#) and [Theorem 13.8](#). We start with the first equation. We have:

$$\begin{aligned}
\sum_{k=1}^d \sum_{\ell=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot \left(\sum_{k'=1}^k \sum_{\ell=1}^d |c_{k'\ell \rightarrow ij}^\lambda|^2 - \sum_{k'=1}^{k-1} \sum_{\ell=1}^d |c_{k'\ell \rightarrow ij}^\lambda|^2 \right) \\
&= \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot (D_{k \rightarrow i}^\lambda \cdot D_{d \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow i}^\lambda \cdot D_{d \rightarrow j}^{\lambda+e_i}) \quad (\text{Theorem 14.10}) \\
&= \left(\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \right) \cdot D_{d \rightarrow j}^{\lambda+e_i} \\
&= \left(\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k \rightarrow i}^\lambda|^2 \right) \cdot D_{d \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.3}) \\
&= (C_i^\lambda + 1) \cdot D_{d \rightarrow i}^\lambda \cdot D_{d \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.8}) \\
&= (C_i^\lambda + 1) \cdot D_{d \rightarrow ij}^\lambda. \quad (\text{Theorem 14.9})
\end{aligned}$$

This proves the first equation.

The second is similar, but slightly more complicated, since when we appeal to [Theorem 14.10](#), we will be using the more complicated case for when $s \geq t$. We have:

$$\begin{aligned}
\sum_{k=1}^d \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot \left(\sum_{k=1}^d \sum_{\ell'=1}^\ell |c_{k\ell' \rightarrow ij}^\lambda|^2 - \sum_{k=1}^d \sum_{\ell'=1}^{\ell-1} |c_{k\ell' \rightarrow ij}^\lambda|^2 \right) \\
&= \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot \left(\left(1 - \frac{1}{\Delta_{ji}^2}\right) (D_{\ell \rightarrow j}^\lambda - D_{\ell-1 \rightarrow j}^\lambda) \cdot D_{d \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} \cdot (D_{\ell \rightarrow i}^\lambda - D_{\ell-1 \rightarrow i}^\lambda) \cdot D_{d \rightarrow j}^{\lambda+e_i} \right) \quad (\text{Theorem 14.10}) \\
&= \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot \left(\left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot |c_{\ell \rightarrow j}^\lambda|^2 \cdot D_{d \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} \cdot |c_{\ell \rightarrow i}^\lambda|^2 \cdot D_{d \rightarrow j}^{\lambda+e_i} \right) \quad (\text{Theorem 13.3}) \\
&= (C_j^\lambda + 1) \cdot \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{d \rightarrow j}^\lambda \cdot D_{d \rightarrow i}^{\lambda+e_j} + (C_i^\lambda + 1) \cdot \frac{1}{\Delta_{ji}^2} \cdot D_{d \rightarrow i}^\lambda \cdot D_{d \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.8}) \\
&= (C_j^\lambda + 1) \cdot \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{d \rightarrow ij}^\lambda + (C_i^\lambda + 1) \cdot \frac{1}{\Delta_{ji}^2} \cdot D_{d \rightarrow ij}^\lambda. \quad (\text{Theorem 14.9})
\end{aligned}$$

We also have:

$$\sum_{k=1}^d \sum_{\ell=1}^d \frac{1}{\Delta_{ji}^2} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = \frac{1}{\Delta_{ji}^2} \cdot D_{d \rightarrow ij}^\lambda = \frac{(C_j^{\lambda+e_i} + 1) - (C_i^\lambda + 1)}{\Delta_{ji}^2} \cdot D_{d \rightarrow ij}^\lambda.$$

Combining these gives:

$$\sum_{k=1}^d \sum_{\ell=1}^d \left(\lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}^2} \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = \left(\left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (C_j^\lambda + 1) + \frac{1}{\Delta_{ji}^2} \cdot (C_j^{\lambda+e_i} + 1) \right) \cdot D_{d \rightarrow ij}^\lambda.$$

In the horizontal subcase, we have $\Delta_{ji}^2 = 1$. In the non-horizontal case, $C_j^\lambda = C_j^{\lambda+e_i}$. In both cases, the sum resolves to the second equation:

$$\sum_{k=1}^d \sum_{\ell=1}^d \left(\lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}^2} \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda. \quad \square$$

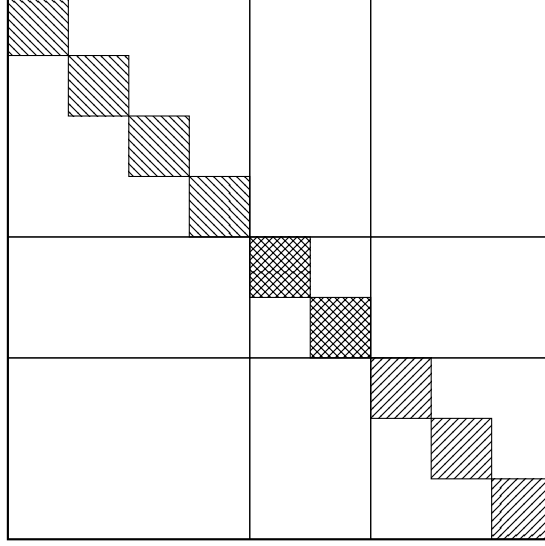


Figure 23: An example of the relation of the two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$ with respect to the blocks of $\lambda = (3, 3, 3, 2, 2, 1, 1, 1)$ for some fixed $i, j \in [d]$. The horizontal axis corresponds to the value of $k \in [9]$, and the vertical axis to the value of $\ell \in [9]$. The Young diagram has 3 blocks, consisting of the indices $\{1, 2, 3, 4\}$, $\{5, 6\}$, and $\{7, 8, 9\}$. The blocks define 9 regions in which the two-step coefficients are equal. The only exceptions are the diagonal entries $|c_{kk \rightarrow ij}^\lambda|^2$ in a block, which are equal to a distinct constant.

Lemma 14.13. *We have the following equation:*

$$\sum_{\substack{k, \ell \\ \max(k, \ell) = d}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \lambda_d^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) \cdot D_{d-1 \rightarrow ij}^\lambda.$$

The proof of [Theorem 14.13](#) is more involved, but requires no new ideas, and we defer it to [Section 14.6.1](#). We also prove a two-step analog to [Theorem 13.5](#), which shows how the two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$, regarded as a function of k and ℓ , are related with respect to the blocks of λ .

Lemma 14.14. *Let λ be a Young diagram, and $i, j \in [d]$. The two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$ satisfy:*

(i) *Let B be a block of λ . Then for each pair of $k, \ell \in B$, with $k \neq \ell$,*

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{BB \rightarrow ij}^\lambda|^2.$$

In addition, when $k = \ell$,

$$|c_{kk \rightarrow ij}^\lambda|^2 = \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot |c_{BB \rightarrow ij}^\lambda|^2.$$

(ii) *Let B_1, B_2 be two distinct blocks of λ . Then the two-step Clebsch-Gordan coefficient $|c_{k\ell \rightarrow ij}^\lambda|^2$ is equal for all $k \in B_1$ and $\ell \in B_2$, that is:*

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{B_1 B_2 \rightarrow ij}^\lambda|^2.$$

The proof of this lemma is also deferred; see [Section 14.6.2](#).

14.4 Main Lemma

Lemma 14.15 (Main Lemma). *It holds that*

$$M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP} = \frac{1}{n^2} \cdot X_{n+1} X_{n+2}, \quad (92)$$

where

$$M_{\text{corr}}^{(2)} = \frac{1}{n} \sum_{\lambda} \dim(V_{\lambda}^d) \sum_{k,\ell} \lambda_{\max(k,\ell)}^{\uparrow} \int_U U^{\otimes(n+2)} \cdot \mathcal{U}^{\dagger} \cdot \left(|\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)} \otimes |T^{\lambda}\rangle\langle T^{\lambda}| \otimes |k\rangle\langle k| \otimes |\ell\rangle\langle\ell| \right) \cdot \mathcal{U} \cdot U^{\dagger, \otimes(n+2)} \cdot dU.$$

Proof. [Theorem 14.4](#) above shows that $M_{\text{avg}}^{(2)}$ and $M_{\text{corr}}^{(2)}$ have the same block structure in the Schur basis as SWAP. In particular, their blocks are indexed by an SYT S with n boxes, and a set of indices $\{i, j\} \in [d]^2$. Let S_{ij} be the SYT obtained from S by inserting a box containing $(n+1)$ into the i -th row, and then a box containing $(n+2)$ into the j -th row. However, we note that S_{ij} may not be a valid SYT. For example, if $\lambda = (1, 1)$, and S is the SYT of shape λ , then S_{12} is valid, but S_{21} is not. Necessarily, however, at least one of S_{ij} and S_{ji} must be valid. Whenever both S_{ij} and S_{ji} are valid SYTs, the block corresponding to $(S, \{i, j\})$ is 2×2 . Otherwise, if one of S_{ij} or S_{ji} is not a valid SYT, the corresponding block is 1×1 , a scalar.

We know from [Theorem 2.36](#) that $X_{n+1}X_{n+2}$ is a diagonal matrix in the Schur basis with the entry $(C_i^{\lambda} + 1)(C_j^{\lambda+e_i} + 1)$ at $|S_{ij}\rangle\langle S_{ij}|$. Thus it suffices to prove [Equation \(92\)](#) for each block $(S, \{i, j\})$:

$$[M_{\text{avg}}^{(2)}]_{\{i,j\}}^S + [M_{\text{corr}}^{(2)} \cdot \text{SWAP}]_{\{i,j\}}^S = \frac{1}{n^2} \cdot [X_{n+1}X_{n+2}]_{\{i,j\}}^S. \quad (93)$$

To show the above equality, we will use the following identity:

$$\sum_{k,\ell} (\lambda_k^{\uparrow} \lambda_{\ell}^{\uparrow} + 1/\Delta_{ji} \cdot \lambda_{\max(k,\ell)}^{\uparrow}) \cdot (|a_{k\ell \rightarrow ij}^{\lambda}|^2 + |b_{k\ell \rightarrow ij}^{\lambda}|^2) = (C_i^{\lambda} + 1)(C_j^{\lambda+e_i} + 1) \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_{\lambda}^d)}, \quad (94)$$

which we prove in [Theorem 14.17](#) below. For the rest of this proof, we fix any such block $(S, \{i, j\})$ and consider the case it corresponds to a 2×2 block, or a 1×1 block separately. Moreover, for notational simplicity, we will omit the superscript S and subscript $\{i, j\}$ for these cases.

1. If both S_{ij} and S_{ji} are distinct, valid SYTs, then $i \neq j$, and we are dealing with a 2×2 block. [Theorem 14.6](#) below implies that on each block $(S, \{i, j\})$, our matrices have the following form:

$$M_{\text{avg}}^{(2)} = x_{\text{avg}} \cdot I + y_{\text{avg}} \cdot \text{SWAP},$$

$$M_{\text{corr}}^{(2)} = x_{\text{corr}} \cdot I + y_{\text{corr}} \cdot \text{SWAP}.$$

Hence

$$M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP} = (x_{\text{avg}} + y_{\text{corr}}) \cdot I + (y_{\text{avg}} + x_{\text{corr}}) \cdot \text{SWAP}.$$

So we would like to show that

$$x_{\text{avg}} + y_{\text{corr}} = \frac{1}{n^2} \cdot (C_i^{\lambda} + 1)(C_j^{\lambda+e_i} + 1), \quad (95)$$

$$y_{\text{avg}} + x_{\text{corr}} = 0. \quad (96)$$

From the structure of SWAP, we can compute x_{avg} and y_{avg} as follows:

$$x_{\text{avg}} = \frac{1}{2} ((M_{\text{avg}}^{(2)})_{ij,ij} + (M_{\text{avg}}^{(2)})_{ji,ji}), \quad y_{\text{avg}} = \frac{\Delta_{ji}}{2} ((M_{\text{avg}}^{(2)})_{ij,ij} - (M_{\text{avg}}^{(2)})_{ji,ji}),$$

and similarly

$$x_{\text{corr}} = \frac{1}{2} ((M_{\text{corr}}^{(2)})_{ij,ij} + (M_{\text{corr}}^{(2)})_{ji,ji}), \quad y_{\text{corr}} = \frac{\Delta_{ji}}{2} ((M_{\text{corr}}^{(2)})_{ij,ij} - (M_{\text{corr}}^{(2)})_{ji,ji}).$$

The left-hand side of Equation (95) becomes

$$\begin{aligned}
& x_{\text{avg}} + y_{\text{corr}} \\
&= \frac{1}{2} \cdot (M_{\text{avg}}^{(2)})_{ij,ij} + \frac{1}{2} \cdot (M_{\text{avg}}^{(2)})_{ji,ji} + \frac{\Delta_{ji}}{2} ((M_{\text{corr}}^{(2)})_{ij,ij} - (M_{\text{corr}}^{(2)})_{ji,ji}) \\
&= \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \lambda_k^\uparrow \lambda_\ell^\uparrow (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) \\
&\quad + \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \lambda_k^\uparrow \lambda_\ell^\uparrow (|a_{k\ell \rightarrow ji}^\lambda|^2 + |b_{k\ell \rightarrow ji}^\lambda|^2) \\
&\quad + \frac{\Delta_{ji}}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \lambda_{\max(k,\ell)}^\uparrow (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 - |a_{k\ell \rightarrow ji}^\lambda|^2 - |b_{k\ell \rightarrow ji}^\lambda|^2) \quad (\text{Theorem 14.4}) \\
&= \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \left(\lambda_k^\uparrow \lambda_\ell^\uparrow + \frac{\lambda_{\max(k,\ell)}^\uparrow}{\Delta_{ji}} \right) (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) \\
&\quad + \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \left(\lambda_k^\uparrow \lambda_\ell^\uparrow + \frac{\lambda_{\max(k,\ell)}^\uparrow}{\Delta_{ji}} \right) (|a_{k\ell \rightarrow ji}^\lambda|^2 + |b_{k\ell \rightarrow ji}^\lambda|^2) \\
&\quad + \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \left(\Delta_{ji} - \frac{1}{\Delta_{ji}} \right) \sum_{k,\ell} \lambda_{\max(k,\ell)}^\uparrow (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 - |a_{k\ell \rightarrow ji}^\lambda|^2 - |b_{k\ell \rightarrow ji}^\lambda|^2).
\end{aligned}$$

From Equation (94), the first two summations are equal to

$$\frac{1}{2n^2} (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) + \frac{1}{2n^2} (C_j^\lambda + 1)(C_i^{\lambda+e_j} + 1) = \frac{1}{n^2} (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1),$$

where we used that the fact that $C_j^{\lambda+e_i} = C_j^\lambda$ and $C_i^{\lambda+e_j} = C_i^\lambda$ when $i \neq j$. It remains to show that the last summation is zero. We will prove that for any $s \in [d]$,

$$\sum_{\substack{k,\ell \\ \max(k,\ell)=s}} (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 - |a_{k\ell \rightarrow ji}^\lambda|^2 - |b_{k\ell \rightarrow ji}^\lambda|^2) = 0.$$

This follows from the partial sums of the two-step Clebsch-Gordan coefficients. In particular, Theorem 14.10 implies that for any $s \in [d]$:

$$\sum_{k=1}^s \sum_{\ell=1}^s (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) = D_{s \rightarrow i}^\lambda D_{s \rightarrow j}^{\lambda+e_i} = \frac{\dim(V_{\lambda_{\leq s} + e_i + e_j}^s)}{\dim(V_{\lambda_{\leq s}}^s)}.$$

Exchanging i and j in the above equation does not change the right-hand side, and thus

$$\begin{aligned}
& \sum_{k=1}^s \sum_{\ell=1}^s (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) = \sum_{k=1}^s \sum_{\ell=1}^s (|a_{k\ell \rightarrow ji}^\lambda|^2 + |b_{k\ell \rightarrow ji}^\lambda|^2) \\
& \implies \sum_{k=1}^s \sum_{\ell=1}^s (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 - |a_{k\ell \rightarrow ji}^\lambda|^2 - |b_{k\ell \rightarrow ji}^\lambda|^2) = 0. \quad (97)
\end{aligned}$$

Subtracting Equation (97) for $s-1$ from the expression above indeed gives:

$$\sum_{\substack{k,\ell \\ \max(k,\ell)=s}} (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2 - |a_{k\ell \rightarrow ji}^\lambda|^2 - |b_{k\ell \rightarrow ji}^\lambda|^2) = 0.$$

The left-hand side of Equation (96) becomes

$$\begin{aligned}
& y_{\text{avg}} + x_{\text{corr}} \\
&= \frac{1}{2} \cdot ((M_{\text{corr}}^{(2)})_{ij,ij} + \Delta_{ji} \cdot (M_{\text{avg}}^{(2)})_{ij,ij}) + \frac{1}{2} \cdot ((M_{\text{corr}}^{(2)})_{ji,ji} - \Delta_{ji} \cdot (M_{\text{avg}}^{(2)})_{ji,ji}) \\
&= \frac{1}{2} \cdot ((M_{\text{corr}}^{(2)})_{ij,ij} + \Delta_{ji} \cdot (M_{\text{avg}}^{(2)})_{ij,ij}) + \frac{1}{2} \cdot ((M_{\text{corr}}^{(2)})_{ji,ji} + \Delta_{ij} \cdot (M_{\text{avg}}^{(2)})_{ji,ji}) \\
&= \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell=1}^d (\lambda_{\max(k,\ell)}^\uparrow + \Delta_{ji} \cdot \lambda_k^\uparrow \lambda_\ell^\uparrow) (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) \\
&\quad + \frac{1}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell=1}^d (\lambda_{\max(k,\ell)}^\uparrow + \Delta_{ij} \cdot \lambda_k^\uparrow \lambda_\ell^\uparrow) (|a_{k\ell \rightarrow ji}^\lambda|^2 + |b_{k\ell \rightarrow ji}^\lambda|^2) \\
&= \frac{\Delta_{ji}}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell=1}^d \left(\frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^\uparrow + \lambda_k^\uparrow \lambda_\ell^\uparrow \right) (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) \\
&\quad + \frac{\Delta_{ij}}{2n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell=1}^d \left(\frac{1}{\Delta_{ij}} \cdot \lambda_{\max(k,\ell)}^\uparrow + \lambda_k^\uparrow \lambda_\ell^\uparrow \right) (|a_{k\ell \rightarrow ji}^\lambda|^2 + |b_{k\ell \rightarrow ji}^\lambda|^2) \\
&= \frac{\Delta_{ji}}{2n^2} \cdot (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) + \frac{\Delta_{ij}}{2n^2} \cdot (C_j^\lambda + 1)(C_i^{\lambda+e_j} + 1) = 0,
\end{aligned}$$

where we used the fact that $C_j^{\lambda+e_i} = C_j^\lambda$ and $C_i^{\lambda+e_j} = C_i^\lambda$ when $i \neq j$.

2. If S_{ij} is a valid SYT, but S_{ji} is either not valid, or not a distinct SYT, then we are dealing with a 1×1 block. First, we obtain the following expressions for the diagonal entries of $M_{\text{avg}}^{(2)}$ and $M_{\text{corr}}^{(2)}$ from Theorem 14.4:

$$\begin{aligned}
M_{\text{avg}}^{(2)} &= \frac{1}{n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2), \\
M_{\text{corr}}^{(2)} &= \frac{1}{n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell} \lambda_{\max(k,\ell)}^\uparrow \cdot (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2).
\end{aligned}$$

The SWAP matrix has the scalar $\frac{1}{\Delta_{ji}}$ in the $(S, \{i, j\})$ block. Hence

$$\begin{aligned}
M_{\text{avg}}^{(2)} + M_{\text{corr}}^{(2)} \cdot \text{SWAP} &= \frac{1}{n^2} \cdot \frac{\dim(V_{\lambda}^d)}{\dim(V_{\lambda_{ij}}^d)} \cdot \sum_{k,\ell} (\lambda_k^\uparrow \lambda_\ell^\uparrow + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^\uparrow) \cdot (|a_{k\ell \rightarrow ij}^\lambda|^2 + |b_{k\ell \rightarrow ij}^\lambda|^2) \\
&= \frac{1}{n^2} \cdot (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \\
&= \frac{1}{n^2} \cdot X_{n+1} X_{n+2}.
\end{aligned}$$

□

14.5 Proof of the Diagonal Expression Lemma

In this section, we prove the Diagonal Expression Lemma, which states that:

$$\sum_{k,\ell} (\lambda_k^\uparrow \lambda_\ell^\uparrow + 1/\Delta_{ji} \cdot \lambda_{\max(k,\ell)}^\uparrow) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_{\lambda}^d)}.$$

This identity can be thought of as the second moment analog of the identity

$$\sum_{k,\ell} \lambda_k^\uparrow \cdot |c_{k \rightarrow i}^\lambda|^2 = (C_i^\lambda + 1) \cdot \frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_{\lambda}^d)}$$

from the first moment. Similar to the first moment, we will first show this identity for the staircase transformation $\lambda^{\uparrow\uparrow}$ in Equation (98), and then “average” the $\lambda^{\uparrow\uparrow}$ ’s over the blocks of λ to obtain the desired statement in Theorem 14.17. In contrast to the first moment, this identity does not hold for all legal estimators.

Lemma 14.16 (Diagonal Expression Lemma for the staircase estimator). *Let λ be a Young diagram. Then the Diagonal Expression Lemma holds for the staircase estimator:*

$$\sum_{k=1}^d \sum_{\ell=1}^d (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_\lambda^d)}. \quad (98)$$

Proof. Fix λ , i and j . We prove the lemma by induction on d . Our base case is $d = 1$. If either i or j are greater than 1, both sides of the desired equation vanish, since $|c_{k\ell \rightarrow ij}^\lambda|^2 = 0$ and $\dim(V_{\lambda_{ij}}^d)/\dim(V_\lambda^d) = D_{d \rightarrow ij}^\lambda = 0$ in cases where $d < i, j$. So, assume $i = j = 1$. Note that in this case, $\Delta_{ji} = 1$, $C_i^\lambda = \lambda_1 - 1$, and $C_j^{\lambda+e_i} = \lambda_1$. The only nonzero CG coefficient is $c_{11 \rightarrow 11}^\lambda$, which must therefore have unit norm. Since $d = 1$, we have $\dim(V_{\lambda_{ij}}^d) = \dim(V_\lambda^d) = 1$, as the only valid SSYT is filled with 1’s. Lastly, we have $\lambda_1^{\uparrow\uparrow} = \lambda_1$. Thus,

$$((\lambda_1^{\uparrow\uparrow})^2 + \lambda_1^{\uparrow\uparrow}) \cdot |c_{11 \rightarrow 11}^\lambda|^2 = \lambda_1^{\uparrow\uparrow}(\lambda_1^{\uparrow\uparrow} + 1) = \lambda_1(\lambda_1 + 1) = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_\lambda^d)},$$

which proves the base case.

Now assume the result holds for $d - 1$. We split the left-hand side of Equation (98) up into the following terms:

$$\begin{aligned} & \sum_{k=1}^d \sum_{\ell=1}^d (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 \\ &= \sum_{k=1}^{d-1} \sum_{\ell=1}^{d-1} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{\substack{k,\ell \\ \max(k,\ell)=d}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2. \end{aligned} \quad (99)$$

We now consider the first term of Equation (99). Recall that $\lambda_k^{\uparrow\uparrow}(d) = \lambda_k^{\uparrow\uparrow}(d-1) + 1$. We have

$$\begin{aligned} & \sum_{k,\ell=1}^{d-1} \left(\lambda_k^{\uparrow\uparrow}(d) \lambda_\ell^{\uparrow\uparrow}(d) + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}(d) \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 \\ &= \sum_{k,\ell=1}^{d-1} \left((\lambda_k^{\uparrow\uparrow}(d-1) + 1)(\lambda_\ell^{\uparrow\uparrow}(d-1) + 1) + \frac{1}{\Delta_{ji}} \cdot (\lambda_{\max(k,\ell)}^{\uparrow\uparrow}(d-1) + 1) \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 \\ &= \sum_{k,\ell=1}^{d-1} \left(\lambda_k^{\uparrow\uparrow}(d-1) \lambda_\ell^{\uparrow\uparrow}(d-1) + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}(d-1) \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{k,\ell=1}^{d-1} \left(\lambda_k^{\uparrow\uparrow}(d-1) + \lambda_\ell^{\uparrow\uparrow}(d-1) + 1 + \frac{1}{\Delta_{ji}} \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2. \end{aligned}$$

We can use the inductive hypothesis on the first sum here, and Theorem 14.12 and Theorem 14.10 for the second. We get:

$$\begin{aligned} & \sum_{k,\ell=1}^{d-1} \left(\lambda_k^{\uparrow\uparrow}(d) \lambda_\ell^{\uparrow\uparrow}(d) + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}(d) \right) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d-1 \rightarrow ij}^\lambda + (C_i^\lambda + C_j^{\lambda+e_i} + 3) \cdot D_{d-1 \rightarrow ij}^\lambda \\ &= (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) \cdot D_{d-1 \rightarrow ij}^\lambda. \end{aligned} \quad (100)$$

We now turn to the second term in Equation (99). This term was calculated in Theorem 14.13, and we have:

$$\sum_{\substack{k,\ell \\ \max(k,\ell)=d}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) \cdot D_{d-1 \rightarrow ij}^\lambda. \quad (101)$$

Finally, substituting [Equations \(100\)](#) and [\(101\)](#) back into [Equation \(99\)](#) we obtain the claimed expression:

$$\sum_{k=1}^d \sum_{\ell=1}^d (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda. \quad \square$$

We next show that the Diagonal Expression Lemma holds for the debiased Keyl's estimator as well. We remind the reader that our general approach is to “average” the analogous result for the staircase estimator.

Lemma 14.17 (Diagonal Expression Lemma for the debiased Keyl's estimator). *Let λ be a Young diagram. Then the Diagonal Expression Lemma holds for the debiased Keyl's estimator:*

$$\sum_{k=1}^d \sum_{\ell=1}^d (\lambda_k^{\uparrow} \lambda_\ell^{\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_\lambda^d)}.$$

Proof. To prove the lemma, we will make use of the results in [Theorem 14.14](#). Our starting point is the Diagonal Expression Lemma for the staircase estimator ([Theorem 14.16](#)), which we rewrite in terms of the blocks $\{B\}$ of λ :

$$\begin{aligned} & \sum_{k=1}^d \sum_{\ell=1}^d (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 \\ &= \sum_{B \neq B'} \sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 + \sum_B \sum_{\substack{k, \ell \in B \\ k \neq \ell}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ & \quad + \sum_B \sum_{k \in B} ((\lambda_k^{\uparrow\uparrow})^2 + \frac{1}{\Delta_{ji}} \cdot \lambda_k^{\uparrow\uparrow}) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \cdot (1 + \frac{1}{\Delta_{ji}}). \end{aligned} \quad (102)$$

We aim to show that we can replace $\lambda_k^{\uparrow\uparrow}$ to $\lambda_k^{\uparrow}$ without changing the value of the combined expression. We informally say that an expression where we can make such a replacement is *averageable*. We begin by showing the first sum by itself is averageable. Assume $B < B'$, meaning that all rows of B occur above all rows of B' . Then

$$\begin{aligned} \sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 &= \left(\sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \right) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 \\ &= \left(\sum_{k \in B} (\lambda_k^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}}) \right) \cdot \left(\sum_{\ell \in B'} \lambda_\ell^{\uparrow\uparrow} \right) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 \\ &= \left(\sum_{k \in B} (\lambda_k^{\uparrow} + \frac{1}{\Delta_{ji}}) \right) \cdot \left(\sum_{\ell \in B'} \lambda_\ell^{\uparrow} \right) \cdot |c_{BB' \rightarrow ij}^\lambda|^2. \end{aligned}$$

The last step holds since the average value on any block of a legal estimator is the same. Undoing the factorizations thus gives:

$$\sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 = \sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow} \lambda_\ell^{\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2.$$

The result also holds if $B > B'$, by swapping $k \leftrightarrow \ell$ in the above reasoning. So,

$$\sum_{B \neq B'} \sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2 = \sum_{B \neq B'} \sum_{k \in B} \sum_{\ell \in B'} (\lambda_k^{\uparrow} \lambda_\ell^{\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow}) \cdot |c_{BB' \rightarrow ij}^\lambda|^2. \quad (103)$$

We now rewrite the remaining two sums of [Equation \(102\)](#) in the following way:

$$\sum_B \sum_{k, \ell \in B} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \cdot \lambda_{\max(k,\ell)}^{\uparrow\uparrow}) \cdot |c_{BB \rightarrow ij}^\lambda|^2 + \sum_B \sum_{k \in B} \frac{1}{\Delta_{ji}} \cdot ((\lambda_k^{\uparrow\uparrow})^2 + \lambda_k^{\uparrow\uparrow}) \cdot |c_{BB \rightarrow ij}^\lambda|^2. \quad (104)$$

The first part of the first sum in Equation (104) is individually averageable:

$$\begin{aligned} \sum_B \sum_{k, \ell \in B} \lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} \cdot |c_{BB \rightarrow ij}^\lambda|^2 &= \sum_B \left(\sum_{k \in B} \lambda_k^{\uparrow\uparrow} \right) \left(\sum_{\ell \in B} \lambda_\ell^{\uparrow\uparrow} \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \sum_B \left(\sum_{k \in B} \lambda_k^\uparrow \right) \left(\sum_{\ell \in B} \lambda_\ell^\uparrow \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 = \sum_B \sum_{k, \ell \in B} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{BB \rightarrow ij}^\lambda|^2. \end{aligned}$$

Similarly the second part of the second sum:

$$\frac{1}{\Delta_{ji}^2} \sum_{k \in B} \lambda_k^{\uparrow\uparrow} \cdot |c_{BB \rightarrow ij}^\lambda|^2 = \frac{1}{\Delta_{ji}^2} \sum_{k \in B} \lambda_k^\uparrow \cdot |c_{BB \rightarrow ij}^\lambda|^2.$$

The remaining two pieces of Equation (104) combine to give an averageable expression. We have

$$\begin{aligned} &\frac{1}{\Delta_{ji}} \sum_B \sum_{k, \ell \in B} \lambda_{\max(k, \ell)}^{\uparrow\uparrow} \cdot |c_{BB \rightarrow ij}^\lambda|^2 + \frac{1}{\Delta_{ji}} \sum_B \sum_{k \in B} (\lambda_k^{\uparrow\uparrow})^2 \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \left(\sum_{k \in B} (N_B(k) \cdot \lambda_k^{\uparrow\uparrow} + (\lambda_k^{\uparrow\uparrow})^2) \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \left(\sum_{k \in B} \lambda_k^{\uparrow\uparrow} \cdot (N_B(k) + \lambda_k^{\uparrow\uparrow}) \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2, \end{aligned} \tag{105}$$

where $N_B(k) = |\{(k_1, k_2) \in B \times B \mid \max(k_1, k_2) = k\}|$. That is, $N_B(k)$ is the number of times k appears in the maximum of a pair of indices in B . We have $N_B(k) = 2k - 1 - 2 \sum_{B' < B} |B'|$, for all $k \in B$. Thus,

$$N_B(k) + \lambda_k^{\uparrow\uparrow} = \left(2k - 1 - 2 \sum_{B' < B} |B'| \right) + (\lambda_B + d + 1 - 2k) = \lambda_B + d - 2 \sum_{B' < B} |B'| =: N'_B,$$

where N'_B is independent of $k \in B$. Therefore, using the fact that $\lambda_k^\uparrow = \lambda_B^\uparrow$ for all $k \in B$, we have:

$$\begin{aligned} (105) &= \frac{1}{\Delta_{ji}} \sum_B N'_B \cdot \left(\sum_{k \in B} \lambda_k^{\uparrow\uparrow} \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B N'_B \cdot \left(\sum_{k \in B} \lambda_k^\uparrow \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \lambda_B^\uparrow \cdot \left(\sum_{k \in B} N_B(k) + \lambda_k^{\uparrow\uparrow} \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \lambda_B^\uparrow \cdot \left(\sum_{k \in B} N_B(k) + \lambda_k^\uparrow \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \left(\sum_{k \in B} (N_B(k) \cdot \lambda_k^\uparrow + (\lambda_k^\uparrow)^2) \right) \cdot |c_{BB \rightarrow ij}^\lambda|^2 \end{aligned}$$

That is,

$$\begin{aligned} &\frac{1}{\Delta_{ji}} \sum_B \sum_{k, \ell \in B} \lambda_{\max(k, \ell)}^{\uparrow\uparrow} \cdot |c_{BB \rightarrow ij}^\lambda|^2 + \frac{1}{\Delta_{ji}} \sum_B \sum_{k \in B} (\lambda_k^{\uparrow\uparrow})^2 \cdot |c_{BB \rightarrow ij}^\lambda|^2 \\ &= \frac{1}{\Delta_{ji}} \sum_B \sum_{k, \ell \in B} \lambda_{\max(k, \ell)}^\uparrow \cdot |c_{BB \rightarrow ij}^\lambda|^2 + \frac{1}{\Delta_{ji}} \sum_B \sum_{k \in B} (\lambda_k^\uparrow)^2 \cdot |c_{BB \rightarrow ij}^\lambda|^2. \end{aligned}$$

Thus, Equation (104) is averageable, and hence Equation (102) is averageable too. $\square$

14.6 Deferred proofs

14.6.1 Proof of Theorem 14.13

The proofs in this section involve elementary, but lengthy calculations. We begin with a few helper lemmas.

Lemma 14.18. *Let λ be a Young diagram of length at most d , and let $i \in [d]$. We have*

$$\sum_{k=1}^d (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot D_{k \rightarrow i}^\lambda = (C_i^\lambda + 1) \cdot D_{d \rightarrow i}^\lambda,$$

where we take $\lambda_{d+1}^{\uparrow\uparrow} = 0$.

Proof. We have

$$\sum_{k=1}^d (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot D_{k \rightarrow i}^\lambda = \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) = \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{k \rightarrow i}^\lambda|^2 = (C_i^\lambda + 1) \cdot D_{d \rightarrow i}^\lambda.$$

In the second sum, $D_{0 \rightarrow i}^\lambda = 0$ by convention. The second equality is by Theorem 13.3; the third is by Theorem 13.8. $\square$

Notation 14.19. In the following lemmas and their proofs, it will be useful to abbreviate the following quantity:

$$B_{ij}^\lambda := (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda - C_d^\lambda + 1) \cdot D_{d-1 \rightarrow ij}^\lambda.$$

We will drop the superscript when λ is clear from context.

Lemma 14.20. *Let λ be a Young diagram of length at most d , and let $i, j \in [d]$. We have*

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{kd \rightarrow ij}^\lambda|^2 = B_{ij}^\lambda + (C_d^\lambda + 1) \cdot (F(d, d) - F(d, d-1)).$$

Proof. Write $G_1(k) := F(k, d) - F(k, d-1)$, and note that $|c_{kd \rightarrow ij}^\lambda|^2 = G_1(k) - G_1(k-1)$, where we take $G_1(0) = 0$. We have

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{kd \rightarrow ij}^\lambda|^2 = \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot (G_1(k) - G_1(k-1)) = \left(\sum_{k=1}^{d-1} (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot G_1(k) \right) + \lambda_d^{\uparrow\uparrow} \cdot G_1(d). \quad (106)$$

The last term can be rewritten as

$$\lambda_d^{\uparrow\uparrow} \cdot G_1(d) = (C_d^\lambda + 1) \cdot (F(d, d) - F(d, d-1)). \quad (107)$$

In the remaining sum, we have $k \leq d-1$, so that from Theorem 14.10, we have

$$G_1(k) = D_{k \rightarrow i}^\lambda \cdot (D_{d \rightarrow j}^{\lambda+e_i} - D_{d-1 \rightarrow j}^{\lambda+e_i}).$$

Then, by Theorem 14.18, we have

$$\begin{aligned} \sum_{k=1}^{d-1} (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot D_{k \rightarrow i}^\lambda &= \left(\sum_{k=1}^d (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot D_{k \rightarrow i}^\lambda \right) - \lambda_d^{\uparrow\uparrow} \cdot D_{d \rightarrow i}^\lambda \\ &= (C_i^\lambda + 1) \cdot D_{d \rightarrow i}^\lambda - (C_d^\lambda + 1) \cdot D_{d \rightarrow i}^\lambda \\ &= (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow i}^\lambda. \end{aligned}$$

So,

$$\begin{aligned}
\sum_{k=1}^{d-1} (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot G_1(k) &= \left(\sum_{k=1}^{d-1} (\lambda_k^{\uparrow\uparrow} - \lambda_{k+1}^{\uparrow\uparrow}) \cdot D_{k \rightarrow i}^\lambda \right) \cdot (D_{d \rightarrow j}^{\lambda+e_i} - D_{d-1 \rightarrow j}^{\lambda+e_i}) \\
&= (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow i}^\lambda \cdot (D_{d \rightarrow j}^{\lambda+e_i} - D_{d-1 \rightarrow j}^{\lambda+e_i}) \\
&= (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow i}^\lambda \cdot D_{d-1 \rightarrow j}^{\lambda+e_i} \\
&= (C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda - C_d^\lambda + 1) \cdot D_{d-1 \rightarrow ij}^\lambda \quad (\text{Theorem 13.3}) \\
&= B_{ij}. \quad (108)
\end{aligned}$$

The lemma follows by substituting Equations (107) and (108) into Equation (106). $\square$

Lemma 14.21. *Let λ be a Young diagram of length at most d , and let $i, j \in [d]$. We have*

$$\sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 = \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot B_{ji} + \frac{1}{\Delta_{ji}^2} \cdot B_{ij} + (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d)).$$

This lemma's proof is almost the same as the previous lemma's, only slightly more complicated by having to use the $s \geq t$ case of Theorem 14.10.

Proof. Write $G_2(\ell) := F(d, \ell) - F(d-1, \ell)$, and note that $|c_{d\ell \rightarrow ij}^\lambda|^2 = G_2(\ell) - G_2(\ell-1)$, where we take $G_2(0) = 0$. We have

$$\sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 = \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot (G_2(\ell) - G_2(\ell-1)) = \left(\sum_{\ell=1}^{d-1} (\lambda_\ell^{\uparrow\uparrow} - \lambda_{\ell+1}^{\uparrow\uparrow}) \cdot G_2(\ell) \right) + \lambda_d^{\uparrow\uparrow} \cdot G_2(d). \quad (109)$$

The last term can be rewritten as

$$\lambda_d^{\uparrow\uparrow} \cdot G_2(d) = (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d)). \quad (110)$$

In the remaining sum, we have $\ell \leq d-1$, so that from Theorem 14.10, we have

$$G_2(\ell) = \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{\ell \rightarrow j}^\lambda \cdot (D_{d \rightarrow i}^{\lambda+e_j} - D_{d-1 \rightarrow i}^{\lambda+e_j}) + \frac{1}{\Delta_{ji}^2} \cdot D_{\ell \rightarrow i}^\lambda \cdot (D_{d \rightarrow j}^{\lambda+e_i} - D_{d-1 \rightarrow j}^{\lambda+e_i}).$$

Then, by calculations similar to those in Equation (108) in the previous lemma's proof, we have

$$\begin{aligned}
\sum_{\ell=1}^{d-1} (\lambda_\ell^{\uparrow\uparrow} - \lambda_{\ell+1}^{\uparrow\uparrow}) \cdot G_2(\ell) &= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot \left((C_j^\lambda - C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_j^\lambda - C_d^\lambda + 1) \cdot D_{d-1 \rightarrow ij}^\lambda \right) \\
&\quad + \frac{1}{\Delta_{ji}^2} \cdot \left((C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda - C_d^\lambda + 1) \cdot D_{d-1 \rightarrow ij}^\lambda \right) \\
&= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot B_{ji} + \frac{1}{\Delta_{ji}^2} \cdot B_{ij}. \quad (111)
\end{aligned}$$

The lemma follows by substituting Equations (110) and (111) into Equation (109). $\square$

We are now ready to prove Theorem 14.13.

Lemma 14.22 (Theorem 14.13, restated). *We have the following equation:*

$$\sum_{\substack{k, \ell \\ \max(k, \ell) = d}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}^2} \lambda_d^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) \cdot D_{d-1 \rightarrow ij}^\lambda.$$

Proof. We start by splitting the left-hand side into four terms:

$$\begin{aligned} \sum_{\substack{k,\ell \\ \max(k,\ell)=d}} (\lambda_k^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} + \frac{1}{\Delta_{ji}} \lambda_d^{\uparrow\uparrow}) \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 &= \sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \lambda_d^{\uparrow\uparrow} \cdot |c_{kd \rightarrow ij}^\lambda|^2 + \sum_{\ell=1}^d \lambda_d^{\uparrow\uparrow} \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 \\ &\quad - (\lambda_d^{\uparrow\uparrow})^2 \cdot |c_{dd \rightarrow ij}^\lambda|^2 + \sum_{\substack{k,\ell \\ \max(k,\ell)=d}} \frac{1}{\Delta_{ji}} \lambda_d^{\uparrow\uparrow} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2. \end{aligned} \quad (112)$$

Note that every term on the right is proportional to $\lambda_d^{\uparrow\uparrow} = C_d^\lambda + 1$. Thus, we start by considering instead the right-hand side of this equation, divided by $\lambda_d^{\uparrow\uparrow}$:

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{kd \rightarrow ij}^\lambda|^2 + \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 - \lambda_d^{\uparrow\uparrow} \cdot |c_{dd \rightarrow ij}^\lambda|^2 + \sum_{\substack{k,\ell \\ \max(k,\ell)=d}} \frac{1}{\Delta_{ji}} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2. \quad (113)$$

The first term of Equation (113) is given by Theorem 14.20 to be:

$$\sum_{k=1}^d \lambda_k^{\uparrow\uparrow} \cdot |c_{kd \rightarrow ij}^\lambda|^2 = B_{ij} + (C_d^\lambda + 1) \cdot (F(d, d) - F(d, d-1)).$$

The second term of Equation (113), by Theorem 14.21, is:

$$\begin{aligned} \sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 &= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot B_{ji} + \frac{1}{\Delta_{ji}^2} \cdot B_{ij} + (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d)) \\ &= B_{ji} + \frac{1}{\Delta_{ji}^2} \cdot (B_{ij} - B_{ji}) + (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d)). \end{aligned}$$

Before pressing on, let us rewrite the difference $B_{ij} - B_{ji}$ as:

$$\begin{aligned} B_{ij} - B_{ji} &= (C_i^\lambda - C_j^\lambda) \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda) \\ &= (\delta_{ij} - \Delta_{ij}) \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda), \end{aligned}$$

where we have used $\Delta_{ji} = C_j^{\lambda+e_i} - C_i^\lambda = C_j^\lambda - C_i^\lambda + \delta_{ij}$. Lastly, notice that $\delta_{ij} \neq 0$ implies $i = j$, which implies $\Delta_{ji} = 1$. Thus, $\delta_{ij}/\Delta_{ji}^2 = \delta_{ij}$, and

$$\sum_{\ell=1}^d \lambda_\ell^{\uparrow\uparrow} \cdot |c_{d\ell \rightarrow ij}^\lambda|^2 = B_{ji} + \left(\delta_{ij} - \frac{1}{\Delta_{ji}}\right) \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda) + (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d)).$$

The third term of Equation (113), by Theorem 14.10, is:

$$-\lambda_d^{\uparrow\uparrow} \cdot |c_{dd \rightarrow ij}^\lambda|^2 = -(C_d^\lambda + 1) \cdot (F(d, d) - F(d, d-1) - F(d-1, d) + F(d-1, d-1)).$$

The fourth term of Equation (113), by Theorem 14.10, is:

$$\sum_{\substack{k,\ell \\ \max(k,\ell)=d}} \frac{1}{\Delta_{ji}} \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 = \frac{1}{\Delta_{ji}} \cdot \left(\sum_{k,\ell=1}^d |c_{k\ell \rightarrow ij}^\lambda|^2 - \sum_{k,\ell=1}^{d-1} |c_{k\ell \rightarrow ij}^\lambda|^2 \right) = \frac{1}{\Delta_{ji}} \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda).$$

Recombining the terms gives:

$$\begin{aligned} (113) &= (B_{ij} + B_{ji}) + \delta_{ij} \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda) + (C_d^\lambda + 1) \cdot (F(d, d) - F(d-1, d-1)) \\ &= (B_{ij} + B_{ji}) + (C_d^\lambda + 1 + \delta_{ij}) \cdot (D_{d \rightarrow ij}^\lambda - D_{d-1 \rightarrow ij}^\lambda), \end{aligned} \quad (114)$$

here using [Theorem 14.10](#) to substitute $F(s, s) = D_{s \rightarrow ij}^\lambda$. Now, we also have

$$B_{ij} + B_{ji} = (C_i^\lambda + C_j^\lambda - 2C_d^\lambda) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + C_j^\lambda - 2C_d^\lambda + 2) \cdot D_{d-1 \rightarrow ij}^\lambda,$$

so that

$$\begin{aligned} (114) &= (C_i^\lambda + C_j^\lambda - C_d^\lambda + 1 + \delta_{ij}) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + C_j^\lambda - C_d^\lambda + 3 + \delta_{ij}) \cdot D_{d-1 \rightarrow ij}^\lambda \\ &= (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 3) \cdot D_{d-1 \rightarrow ij}^\lambda. \end{aligned}$$

Thus,

$$(112) = (C_d^\lambda + 1) \cdot (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_d^\lambda + 1) \cdot (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 3) \cdot D_{d-1 \rightarrow ij}^\lambda. \quad (115)$$

Our final ingredient comes from the following identity:

$$\begin{aligned} (C_i^\lambda - C_d^\lambda) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i}) \cdot D_{d \rightarrow ij}^\lambda &= \left((C_i^\lambda - C_d^\lambda) \cdot D_{d \rightarrow i}^\lambda \right) \cdot \left((C_j^{\lambda+e_i} - C_d^{\lambda+e_i}) \cdot D_{d \rightarrow j}^{\lambda+e_i} \right) \\ &= \left((C_i^\lambda - C_d^\lambda + 1) \cdot D_{d-1 \rightarrow i}^\lambda \right) \cdot \left((C_j^{\lambda+e_i} - C_d^{\lambda+e_i} + 1) \cdot D_{d-1 \rightarrow j}^{\lambda+e_i} \right) \\ &= (C_i^\lambda - C_d^\lambda + 1) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i} + 1) \cdot D_{d-1 \rightarrow ij}^\lambda, \end{aligned}$$

using [Theorem 13.3](#). In particular, we will add zero to [Equation \(115\)](#) as:

$$0 = (C_i^\lambda - C_d^\lambda) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i}) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda - C_d^\lambda + 1) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i} + 1) \cdot D_{d-1 \rightarrow ij}^\lambda.$$

However, the coefficient of $D_{d \rightarrow ij}^\lambda$ then becomes:

$$\begin{aligned} &(C_d^\lambda + 1) \cdot (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 1) + (C_i^\lambda - C_d^\lambda) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i}) \\ &= (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) + (C_d^\lambda - C_d^{\lambda+e_i}) \cdot (C_i^\lambda - C_d^\lambda) \\ &= (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1), \end{aligned}$$

where the last step holds since the term $C_d^\lambda - C_d^{\lambda+e_i} = -\delta_{id}$ is nonzero if and only if $C_i^\lambda = C_d^\lambda$. Similarly, the coefficient of $D_{d-1 \rightarrow ij}^\lambda$ becomes:

$$\begin{aligned} &(C_d^\lambda + 1) \cdot (C_i^\lambda + C_j^{\lambda+e_i} - C_d^\lambda + 3) + (C_i^\lambda - C_d^\lambda + 1) \cdot (C_j^{\lambda+e_i} - C_d^{\lambda+e_i} + 1) \\ &= (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) + (C_d^\lambda - C_d^{\lambda+e_i}) \cdot (C_i^\lambda - C_d^\lambda) \\ &= (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2). \end{aligned}$$

Thus,

$$(115) = (C_i^\lambda + 1)(C_j^{\lambda+e_i} + 1) \cdot D_{d \rightarrow ij}^\lambda - (C_i^\lambda + 2)(C_j^{\lambda+e_i} + 2) \cdot D_{d-1 \rightarrow ij}^\lambda,$$

which completes the proof of the lemma. $\square$

14.6.2 Proof of [Theorem 14.14](#)

In this subsection, we prove the relations between the two-step Clebsch-Gordan coefficients with respect to the blocks of the Young diagram λ .

Lemma 14.23 ([Theorem 14.14](#), restated). *Let λ be a Young diagram, and $i, j \in [d]$. There exist constants $\{c_{B_1 B_2 \rightarrow ij}^\lambda\}_{B_1, B_2}$, where B_1, B_2 are blocks of λ , such that the two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$ satisfy:*

- (i) *Let B be a block of λ . Then the two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$ are equal for all pairs of distinct $k, \ell \in B$, that is:*

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{BB \rightarrow ij}^\lambda|^2.$$

In addition, when $k = \ell$,

$$|c_{kk \rightarrow ij}^\lambda|^2 = \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot |c_{BB \rightarrow ij}^\lambda|^2.$$

(ii) Let B_1, B_2 be two distinct blocks of λ . Then the two-step Clebsch-Gordan coefficients $|c_{k\ell \rightarrow ij}^\lambda|^2$ are equal for all $k \in B_1$ and $\ell \in B_2$, that is:

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{B_1 B_2 \rightarrow ij}^\lambda|^2.$$

Proof. The proof follows by combining [Theorems 14.25](#) to [14.27](#) and [14.29](#). These lemmas are proven below.

We first prove item (ii). Since blocks contain consecutive indices, and B_1, B_2 are distinct blocks of λ , one of the blocks has strictly smaller row indices than the other block. Without loss of generality, let B_1 be the block with the smaller indices. We would like to show that for all $k, k' \in B_1$ and $\ell, \ell' \in B_2$:

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{k'\ell' \rightarrow ij}^\lambda|^2.$$

Since $k, k' < \ell, \ell'$ [Theorem 14.25](#) implies the equality above for all such k, k' and ℓ, ℓ' .

Let us now prove item (i). For a block B of λ , let k, k', ℓ, ℓ' be elements of B . If $|B| = 1$, then the statement trivially holds; thus, we assume that $|B| \geq 2$, and use $m, m+1$ to denote the two smallest indices of the block.

We first consider the case of distinct $k \neq \ell$. If $k < \ell$ and $k' < \ell'$, or $k > \ell$ and $k' > \ell'$, then [Theorem 14.25](#) again implies that $|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{k'\ell' \rightarrow ij}^\lambda|^2$. Otherwise, we assume without loss of generality that $k < \ell$ and $k' > \ell'$ (the proof of the $k > \ell$ and $k' < \ell'$ case is identical to this one). [Theorem 14.25](#) implies that

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{m, m+1 \rightarrow ij}^\lambda|^2, \quad |c_{k'\ell' \rightarrow ij}^\lambda|^2 = |c_{m+1, m \rightarrow ij}^\lambda|^2.$$

We show in [Theorem 14.26](#) that the right-hand sides of the two equations above are equal, which completes the $k \neq \ell$ case of item (i).

The case of $k = \ell$ remains. If k is not the first row in the block B , then [Theorem 14.27](#) implies that $|c_{kk \rightarrow ij}^\lambda|^2 = (1 + 1/\Delta_{ji}) \cdot |c_{k-1, k \rightarrow ij}^\lambda|^2$. Since $k-1, k \in B$, we have argued earlier that $|c_{k-1, k \rightarrow ij}^\lambda|^2 = |c_{B B \rightarrow ij}^\lambda|^2$, and the desired statement holds. Note that the above argument does not work if k is the first row of block B , since $k-1$ is not in block B in that case. This case is proved separately in [Theorem 14.29](#). $\square$

Before we proceed with the proofs of the remaining lemmas, we introduce the following piece of notation that will simplify our calculations.

Notation 14.24. Let $F_1(s, t)$ and $F_2(s, t)$ denote the two expressions that arise in the partial sums of the two-step Clebsch-Gordan coefficients:

$$F_1(s, t) := D_{s \rightarrow i}^\lambda \cdot D_{t \rightarrow j}^{\lambda+e_i}, \quad F_2(s, t) := \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{t \rightarrow j}^\lambda \cdot D_{s \rightarrow i}^{\lambda+e_j} + \frac{1}{\Delta_{ji}^2} \cdot D_{t \rightarrow i}^\lambda \cdot D_{s \rightarrow j}^{\lambda+e_i}.$$

We extend the definition of $F_1(s, t)$ and $F_2(s, t)$ to all values of s, t , even though they have a more restricted domain in [Theorem 14.10](#).

In particular, the statement of [Theorem 14.10](#) simplifies to

$$F(s, t) = \sum_{k=1}^s \sum_{\ell=1}^t |c_{k\ell \rightarrow ij}^\lambda|^2 = \begin{cases} F_1(s, t) & s \leq t, \\ F_2(s, t) & s > t. \end{cases}$$

Lemma 14.25. Let λ be a Young diagram, and B_1, B_2 be two blocks of λ (not necessarily distinct), such that $k, k' \in B_1$, and $\ell, \ell' \in B_2$. If either

(i) $k < \ell$ and $k' < \ell'$, or

(ii) $k > \ell$ and $k' > \ell'$,

then

$$|c_{k\ell \rightarrow ij}^\lambda|^2 = |c_{k'\ell' \rightarrow ij}^\lambda|^2.$$

Proof. Let us consider the case when $k < \ell$ and $k' < \ell'$ first. We write

$$\begin{aligned} |c_{k\ell \rightarrow ij}^\lambda|^2 &= F(k, \ell) - F(k-1, \ell) - F(k, \ell-1) + F(k-1, \ell-1) \\ &= D_{k \rightarrow i}^\lambda D_{\ell \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow i}^\lambda D_{\ell \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow i}^\lambda D_{\ell-1 \rightarrow j}^{\lambda+e_i} + D_{k-1 \rightarrow i}^\lambda D_{\ell-1 \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 14.10}) \\ &= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{\ell \rightarrow j}^{\lambda+e_i} - D_{\ell-1 \rightarrow j}^{\lambda+e_i}). \end{aligned} \quad (116)$$

Similarly,

$$|c_{k'\ell' \rightarrow ij}^\lambda|^2 = (D_{k' \rightarrow i}^\lambda - D_{k'-1 \rightarrow i}^\lambda) \cdot (D_{\ell' \rightarrow j}^{\lambda+e_i} - D_{\ell'-1 \rightarrow j}^{\lambda+e_i}). \quad (117)$$

Since k and k' are in the same block of λ , it holds from Theorem 13.5 that

$$D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda = |c_{k \rightarrow i}^\lambda|^2 = |c_{k' \rightarrow i}^\lambda|^2 = D_{k' \rightarrow i}^\lambda - D_{k'-1 \rightarrow i}^\lambda.$$

Moreover, if ℓ and ℓ' are in the same block of $\lambda + e_i$, then it also follows that

$$D_{\ell \rightarrow j}^{\lambda+e_i} - D_{\ell-1 \rightarrow j}^{\lambda+e_i} = |c_{\ell \rightarrow j}^{\lambda+e_i}|^2 = |c_{\ell' \rightarrow j}^{\lambda+e_i}|^2 = D_{\ell' \rightarrow j}^{\lambda+e_i} - D_{\ell'-1 \rightarrow j}^{\lambda+e_i},$$

which implies that the right-hand sides of Equation (116) and Equation (117) are equal.

If ℓ and ℓ' are not in the same block of $\lambda + e_i$, then, because ℓ and ℓ' are in the same block of λ , we must have $i = \ell$ or $i = \ell'$. Assume $i = \ell$; the $i = \ell'$ case is similar. Since $k < \ell$, this means that $k < i$, and thus the respective one-step Clebsch-Gordan coefficient is equal to zero:

$$0 = |c_{k \rightarrow i}^\lambda|^2 = D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda = D_{k' \rightarrow i}^\lambda - D_{k'-1 \rightarrow i}^\lambda.$$

We conclude that in all cases the right-hand sides of Equation (116) and Equation (117) are equal, and so are the two-step Clebsch-Gordan coefficients.

We now turn our attention to the case when $k > \ell$ and $k' > \ell'$. Using Theorem 14.10 we write:

$$\begin{aligned} |c_{k\ell \rightarrow ij}^\lambda|^2 &= F(k, \ell) - F(k-1, \ell) - F(k, \ell-1) + F(k-1, \ell-1) \\ &= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (D_{\ell \rightarrow j}^\lambda - D_{\ell-1 \rightarrow j}^\lambda) \cdot (D_{k \rightarrow i}^{\lambda+e_j} - D_{k-1 \rightarrow i}^{\lambda+e_j}) + \frac{1}{\Delta_{ji}^2} \cdot (D_{\ell \rightarrow i}^\lambda - D_{\ell-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) \\ &= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot |c_{\ell k \rightarrow ji}^\lambda|^2 + \frac{1}{\Delta_{ji}^2} \cdot |c_{\ell k \rightarrow ij}^\lambda|^2, \end{aligned} \quad (118)$$

where the last equality follows from Equation (116). Similarly,

$$|c_{k'\ell' \rightarrow ij}^\lambda|^2 = \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot |c_{\ell' k' \rightarrow ji}^\lambda|^2 + \frac{1}{\Delta_{ji}^2} \cdot |c_{\ell' k' \rightarrow ij}^\lambda|^2. \quad (119)$$

Since $\ell < k$ and $\ell' < k'$, the first part of this proof implies that

$$|c_{\ell k \rightarrow ji}^\lambda|^2 = |c_{\ell' k' \rightarrow ji}^\lambda|^2, \quad |c_{\ell k \rightarrow ij}^\lambda|^2 = |c_{\ell' k' \rightarrow ij}^\lambda|^2,$$

and we conclude that the right-hand sides of Equation (118) and Equation (119) are equal, and so are the corresponding two-step Clebsch-Gordan coefficients. $\square$

Lemma 14.26. *Let λ be a Young diagram, and B be a block of λ , such that $k, k+1 \in B$. Then*

$$|c_{k, k+1 \rightarrow ij}^\lambda|^2 = |c_{k+1, k \rightarrow ij}^\lambda|^2.$$

Proof. First, we observe that since $k, k+1$ are in the same block, when $i = k+1$, both two-step coefficients are equal to zero, since $\lambda + e_{k+1}$ is an invalid Young diagram. Similarly, both coefficients are zero when $j = k+1$ and $i \neq k$. Thus, for the remainder of this proof, we will consider the case when $i = k, j = k+1$, or when both $i, j \neq k+1$.

Using Theorem 14.10 we write:

$$\begin{aligned} |c_{k, k+1 \rightarrow ij}^\lambda|^2 &= F(k, k+1) - F(k-1, k+1) - F(k, k) + F(k-1, k) \\ &= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) \\ &= (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}), \end{aligned} \quad (k, k+1 \in B)$$

and similarly

$$\begin{aligned}
|c_{k+1,k \rightarrow ij}^\lambda|^2 &= F(k+1, k) - F(k, k) - F(k+1, k-1) + F(k, k-1) \\
&= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (D_{k \rightarrow j}^\lambda - D_{k-1 \rightarrow j}^\lambda) \cdot (D_{k+1 \rightarrow i}^{\lambda+e_j} - D_{k \rightarrow i}^{\lambda+e_j}) + \frac{1}{\Delta_{ji}^2} \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) \\
&= \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (D_{k+1 \rightarrow j}^\lambda - D_{k \rightarrow j}^\lambda) \cdot (D_{k+1 \rightarrow i}^{\lambda+e_j} - D_{k \rightarrow i}^{\lambda+e_j}) + \frac{1}{\Delta_{ji}^2} \cdot (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}),
\end{aligned}$$

where the last equation follows because k and $k+1$ are in the same block of λ . Equating the two previous expressions, it suffices to show that:

$$\left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (D_{k+1 \rightarrow j}^\lambda - D_{k \rightarrow j}^\lambda) \cdot (D_{k+1 \rightarrow i}^{\lambda+e_j} - D_{k \rightarrow i}^{\lambda+e_j}) = \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}).$$

When $i = k$ and $j = k+1$, then $\Delta_{ji}^2 = 1$, and thus the equation above is trivially satisfied. Let us now restrict our attention to the case when $\Delta_{ji}^2 \neq 1$ and both i, j are not equal to $k+1$. We observe that

$$D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i} = \frac{\dim(V_{\lambda \leq k+e_i+e_j}^k)}{\dim(V_{\lambda \leq k}^k)} = D_{k \rightarrow j}^\lambda D_{k \rightarrow i}^{\lambda+e_j} = D_{k \rightarrow ij}^\lambda, \quad (120)$$

hence, what we want to show above simplifies to

$$D_{k+1 \rightarrow j}^\lambda D_{k \rightarrow i}^{\lambda+e_j} + D_{k \rightarrow j}^\lambda D_{k+1 \rightarrow i}^{\lambda+e_j} = D_{k+1 \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i} + D_{k \rightarrow i}^\lambda D_{k+1 \rightarrow j}^{\lambda+e_i}. \quad (121)$$

Towards using Equation (120) again, we would like to use Theorem 13.3 to change the “ $D_{k+1 \rightarrow \cdot}$ ” factors to “ $D_{k \rightarrow \cdot}$ ” factors. Since $i, j \neq k+1$, we can write:

$$\begin{aligned}
D_{k+1 \rightarrow j}^\lambda &= \left(1 + \frac{1}{C_j^\lambda - C_{k+1}^\lambda}\right) \cdot D_{k \rightarrow j}^\lambda, & D_{k+1 \rightarrow i}^\lambda &= \left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) \cdot D_{k \rightarrow i}^\lambda, \\
D_{k+1 \rightarrow i}^{\lambda+e_j} &= \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) \cdot D_{k \rightarrow i}^{\lambda+e_j}, & D_{k+1 \rightarrow j}^{\lambda+e_i} &= \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) \cdot D_{k \rightarrow j}^{\lambda+e_i}.
\end{aligned}$$

To show Equation (121), it suffices to show that

$$\begin{aligned}
&\iff \left(1 + \frac{1}{C_j^\lambda - C_{k+1}^\lambda} + 1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) D_{k \rightarrow ij}^\lambda = \left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda} + 1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) D_{k \rightarrow ij}^\lambda \\
&\iff \left(\frac{1}{C_j^\lambda - C_{k+1}^\lambda} + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) D_{k \rightarrow ij}^\lambda = \left(\frac{1}{C_i^\lambda - C_{k+1}^\lambda} + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) D_{k \rightarrow ij}^\lambda. \quad (122)
\end{aligned}$$

Theorem 13.3 also implies the following:

$$\begin{aligned}
D_{k+1 \rightarrow ij}^\lambda &= D_{k+1 \rightarrow i}^\lambda D_{k+1 \rightarrow j}^{\lambda+e_i} = \left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) D_{k \rightarrow i}^\lambda \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) D_{k+1 \rightarrow j}^{\lambda+e_i} \\
&= \left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) \cdot \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) \cdot D_{k \rightarrow ij}^\lambda,
\end{aligned}$$

and

$$\begin{aligned}
D_{k+1 \rightarrow ij}^\lambda &= D_{k+1 \rightarrow j}^\lambda D_{k+1 \rightarrow i}^{\lambda+e_j} = \left(1 + \frac{1}{C_j^\lambda - C_{k+1}^\lambda}\right) D_{k \rightarrow j}^\lambda \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) D_{k+1 \rightarrow i}^{\lambda+e_j} \\
&= \left(1 + \frac{1}{C_j^\lambda - C_{k+1}^\lambda}\right) \cdot \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) \cdot D_{k \rightarrow ij}^\lambda.
\end{aligned}$$

Hence

$$\left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) \cdot \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) \cdot D_{k \rightarrow ij}^\lambda = \left(1 + \frac{1}{C_j^\lambda - C_{k+1}^\lambda}\right) \cdot \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) \cdot D_{k \rightarrow ij}^\lambda. \quad (123)$$

Subtracting Equation (123) from Equation (122), we only need to show that

$$\frac{1}{C_i^\lambda - C_{k+1}^\lambda} \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}} \cdot D_{k \rightarrow ij}^\lambda = \frac{1}{C_j^\lambda - C_{k+1}^\lambda} \cdot \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}} \cdot D_{k \rightarrow ij}^\lambda.$$

We remark that $C_y^{\lambda+e_x} = C_y^\lambda + \delta_{xy}$, and thus we rewrite the above expression as

$$(C_i^\lambda - C_{k+1}^\lambda) \cdot (C_j^\lambda - C_{k+1}^\lambda + \delta_{ij} - \delta_{i(k+1)}) \cdot D_{k \rightarrow ij}^\lambda = (C_j^\lambda - C_{k+1}^\lambda) \cdot (C_i^\lambda - C_{k+1}^\lambda - \delta_{ij} - \delta_{j(k+1)}) \cdot D_{k \rightarrow ij}^\lambda.$$

Collecting terms, it suffices to show that

$$\begin{aligned} (C_j^\lambda - C_{k+1}^\lambda) \cdot (\delta_{ji} - \delta_{j(k+1)}) &= (C_i^\lambda - C_{k+1}^\lambda) \cdot (\delta_{ij} - \delta_{i(k+1)}) \\ \iff (C_i^\lambda - C_{k+1}^\lambda) \cdot \delta_{i(k+1)} - (C_j^\lambda - C_{k+1}^\lambda) \cdot \delta_{j(k+1)} + (C_j^\lambda - C_i^\lambda) \cdot \delta_{ij} &= 0. \end{aligned}$$

The last equation holds because the expression $(C_x^\lambda - C_y^\lambda) \cdot \delta_{xy}$ is zero for all x, y . $\square$

Lemma 14.27. *Let λ be a Young diagram, and B be a block of λ , such that $k, k+1 \in B$. Then*

$$|c_{k+1, k+1 \rightarrow ij}^\lambda|^2 = \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot |c_{k, k+1 \rightarrow ij}^\lambda|^2.$$

Proof. We first observe that if $i > k+1$, then the Clebsch-Gordan coefficients on both sides are equal to zero. Moreover, both coefficients vanish when $i = k+1$. This is because the first insertion of the left coefficient corresponds to the Young diagram of shape $\lambda + e_{k+1}$, and this is not a valid shape since $k, k+1$ are in the same block. Similarly, the first insertion of the right coefficient corresponds to adding k to row $i > k$, which results in an invalid SSYT. Hence, we can restrict our attention to when $i \leq k$.

Similarly, j has to be at most $k+1$ since we add a letter that is at most $k+1$ to row j . And since $k, k+1$ are in the same block, we can only have $j = k+1$ when $i = k$, in which case the new boxes are added in a vertical strip. Then $\Delta_{ji} = -1$ and the right-hand side vanishes. In that “vertical” case, the left-hand side also vanishes, because the coefficient corresponds to the SSYT $T_{k+1, k+1 \rightarrow ij}^\lambda$, which is an invalid tableau, since it contains the letter $k+1$ in the same column as the same letter. For the rest of this proof, we assume that $i, j \leq k$.

On the left-hand side, we have

$$\begin{aligned} &|c_{k+1, k+1 \rightarrow ij}^\lambda|^2 \\ &= F_1(k+1, k+1) - F_1(k, k+1) - F_2(k+1, k) + F_1(k, k) \\ &= (F_1(k+1, k+1) - F_1(k, k+1) - F_1(k+1, k) + F_1(k, k)) + (F_1(k+1, k) - F_2(k+1, k)) \\ &= (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) + (F_1(k+1, k) - F_2(k+1, k)). \end{aligned}$$

On the other hand,

$$\begin{aligned} |c_{k, k+1 \rightarrow ij}^\lambda|^2 &= F_1(k, k+1) - F_1(k-1, k+1) - F_1(k, k) + F_1(k-1, k) \\ &= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) \\ &= (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) \quad (k, k+1 \in B) \end{aligned}$$

Comparing the two expressions, it suffices to show that

$$F_1(k+1, k) - F_2(k+1, k) = \frac{1}{\Delta_{ji}} \cdot (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda)(D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}). \quad (124)$$

We prove this equation holds in Theorem 14.28 below. This proof is deferred to a separate lemma because it is also used in the proof of Theorem 14.29, and because proving Equation (124) does not require that k and $k+1$ are in the same block. $\square$

Lemma 14.28. *Let λ be a Young diagram, and let $i, j, k \in [d]$ such that $k+1 > i, j$. Then*

$$F_1(k+1, k) - F_2(k+1, k) = \frac{1}{\Delta_{ji}} \cdot (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}).$$

Proof. We use [Theorem 13.3](#) to rewrite the right-hand side of the desired statement:

$$\begin{aligned} \frac{1}{\Delta_{ji}} \cdot (D_{k+1 \rightarrow i}^\lambda - D_{k \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) &= \frac{1}{\Delta_{ji}} \cdot \frac{1}{C_i^\lambda - C_{k+1}^\lambda} \cdot D_{k \rightarrow i}^\lambda \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}} \cdot D_{k \rightarrow j}^{\lambda+e_i} \\ &= \frac{1}{\Delta_{ji}} \cdot \frac{1}{C_i^\lambda - C_{k+1}^\lambda} \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}} \cdot D_{k \rightarrow ij}^\lambda. \end{aligned} \quad (125)$$

The fact that $i, j < k+1$ ensures that the denominators of the above expressions are never zero. Similarly, we write the left-hand side of the lemma statement as:

$$\begin{aligned} &F_1(k+1, k) - F_2(k+1, k) \\ &= D_{k+1 \rightarrow i}^\lambda \cdot D_{k \rightarrow j}^{\lambda+e_i} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{k \rightarrow j}^\lambda \cdot D_{k+1 \rightarrow i}^{\lambda+e_j} - \frac{1}{\Delta_{ji}^2} \cdot D_{k \rightarrow i}^\lambda \cdot D_{k+1 \rightarrow j}^{\lambda+e_i} \\ &= \left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) \cdot D_{k \rightarrow i}^\lambda \cdot D_{k \rightarrow j}^{\lambda+e_i} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{k \rightarrow j}^\lambda \cdot \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) \cdot D_{k \rightarrow i}^{\lambda+e_j} \\ &\quad - \frac{1}{\Delta_{ji}^2} \cdot D_{k \rightarrow i}^\lambda \cdot \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right) \cdot D_{k \rightarrow j}^{\lambda+e_i} \\ &= \left[\left(1 + \frac{1}{C_i^\lambda - C_{k+1}^\lambda}\right) - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \left(1 + \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}}\right) - \frac{1}{\Delta_{ji}^2} \left(1 + \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right)\right] \cdot D_{k \rightarrow ij}^\lambda \\ &= \left[\frac{1}{C_i^\lambda - C_{k+1}^\lambda} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot \frac{1}{C_i^{\lambda+e_j} - C_{k+1}^{\lambda+e_j}} - \frac{1}{\Delta_{ji}^2} \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right] \cdot D_{k \rightarrow ij}^\lambda. \end{aligned} \quad (126)$$

We show that the two sides are equal by considering the following two cases:

Case 1: $i = j$. In this case, $\Delta_{ji} = C_j^{\lambda+e_i} - C_i^\lambda = 1$, so that we have:

$$\begin{aligned} (126) &= \left[\frac{1}{C_i^\lambda - C_{k+1}^\lambda} - \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}}\right] \cdot D_{k \rightarrow ij}^\lambda \\ &= \frac{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i} - C_i^\lambda + C_{k+1}^\lambda}{(C_i^\lambda - C_{k+1}^\lambda) \cdot (C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i})} \cdot D_{k \rightarrow ij}^\lambda \\ &= \frac{1}{(C_i^\lambda - C_{k+1}^\lambda) \cdot (C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i})} \cdot D_{k \rightarrow ij}^\lambda = (125) \quad (i \neq k+1 \implies C_{k+1}^{\lambda+e_i} = C_{k+1}^\lambda) \end{aligned}$$

Case 2: $i \neq j$. In this case, $C_i^{\lambda+e_j} = C_i^\lambda$ and $C_j^{\lambda+e_i} = C_j^\lambda$. Let us use C_i, C_j to refer to these quantities. Moreover, $\Delta_{ji} = (C_j^{\lambda+e_i} + 1) - (C_i^\lambda + 1) = C_j - C_i$. Then, we get:

$$\begin{aligned} (126) &= \left[\frac{1}{C_i - C_{k+1}} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot \frac{1}{C_i - C_{k+1}} - \frac{1}{\Delta_{ji}^2} \cdot \frac{1}{C_j - C_{k+1}}\right] \cdot D_{k \rightarrow ij}^\lambda \\ &= \frac{1}{\Delta_{ji}^2} \cdot \frac{C_j - C_i}{(C_i - C_{k+1})(C_j - C_{k+1})} \cdot D_{k \rightarrow ij}^\lambda \\ &= \frac{1}{\Delta_{ji}^2} \cdot \frac{\Delta_{ji}}{(C_i - C_{k+1})(C_j - C_{k+1})} \cdot D_{k \rightarrow ij}^\lambda = (125). \quad \square \end{aligned}$$

Lemma 14.29. *Let λ be a Young diagram, and B be a block of λ , such that $k, k+1 \in B$. Then*

$$|c_{kk \rightarrow ij}^\lambda|^2 = \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot |c_{k, k+1 \rightarrow ij}^\lambda|^2.$$

Proof. We first observe that if $i > k$, then the Clebsch-Gordan coefficients on both sides are equal to zero. Hence, we can restrict our attention to when $i \leq k$.

Similarly, if $j > k + 1$, then both coefficients are zero. The same happens when $i < k, j = k + 1$, since the left coefficient has an invalid SSYT, and the right coefficient corresponds to the invalid shape $\lambda + e_i + e_{k+1}$. When $i = k, j = k + 1$, then the left-hand coefficient is zero, and $\Delta_{ji} = -1$ (since $k, k + 1$ are in the same block), thus the right-hand side also vanishes. We conclude that we can further restrict our attention to when $i, j \leq k$.

On the left-hand side, we have

$$\begin{aligned} |c_{kk \rightarrow ij}^\lambda|^2 &= F_1(k, k) - F_1(k - 1, k) - F_2(k, k - 1) + F_1(k - 1, k - 1) \\ &= (F_1(k, k) - F_1(k - 1, k) - F_1(k, k - 1) + F_1(k - 1, k - 1)) + (F_1(k, k - 1) - F_2(k, k - 1)) \\ &= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + (F_1(k, k - 1) - F_2(k, k - 1)). \end{aligned} \quad (127)$$

The right-hand side is equal to

$$\begin{aligned} \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot |c_{k, k+1 \rightarrow ij}^\lambda|^2 &= \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot (F_1(k, k + 1) - F_1(k - 1, k + 1) - F_1(k, k) + F_1(k - 1, k)) \\ &= \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}). \end{aligned} \quad (128)$$

Let us now consider the following four cases, depending on whether i, j are equal to k or strictly less than k .

Case 1: $i, j < k$. Since $i < k, k$ and $k + 1$ are in the same block of $\lambda + e_i$, which allows us to write:

$$(128) = \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}).$$

Comparing the expressions for the left and right-hand sides, it suffices to show that

$$F_1(k, k - 1) - F_2(k, k - 1) = \frac{1}{\Delta_{ji}} \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}). \quad (129)$$

The above equation follows directly from [Theorem 14.28](#) when we use $k - 1$ in the place of k , and since $i, j < k$.

Case 2: $i = j = k$. In that case, $\Delta_{ji} = 1$, and recall that the term $D_{k-1 \rightarrow i}^\lambda$ is zero when $i > k - 1$. The left-hand side can be simplified to

$$(127) = D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i} + (F_1(k, k - 1) - F_2(k, k - 1)) = D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i},$$

where the last equality holds because both $F_1(k, k - 1)$ and $F_2(k, k - 1)$ have terms of the form $D_{k-1 \rightarrow k}^\mu$, which are zero for any diagram μ . The right-hand side can be simplified to

$$\begin{aligned} (128) &= \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot D_{k \rightarrow i}^\lambda \cdot (D_{k+1 \rightarrow j}^{\lambda+e_i} - D_{k \rightarrow j}^{\lambda+e_i}) \\ &= 2 \cdot D_{k \rightarrow i}^\lambda \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}} \cdot D_{k \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.3}) \\ &= 2 \cdot D_{k \rightarrow i}^\lambda \cdot \frac{1}{(\lambda_k + 1 - k) - (\lambda_{k+1} - k - 1)} \cdot D_{k \rightarrow j}^{\lambda+e_i} = D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i}. \quad (i = j = k \text{ and } \lambda_k = \lambda_{k+1}) \end{aligned}$$

Hence, the two sides are equal.

Case 3: $i = k, j < k$. In that case, $\Delta_{ji} = C_j^{\lambda+e_i} - C_i^\lambda = C_j^\lambda - C_k^\lambda$. The left-hand side can be simplified to

$$\begin{aligned}
(127) &= D_{k \rightarrow i}^\lambda \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + (F_1(k, k-1) - F_2(k, k-1)) \\
&= D_{k \rightarrow i}^\lambda \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + D_{k \rightarrow i}^\lambda \cdot D_{k-1 \rightarrow j}^{\lambda+e_i} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot D_{k-1 \rightarrow j}^\lambda \cdot D_{k \rightarrow i}^{\lambda+e_j} \\
&= D_{k \rightarrow i}^\lambda \cdot D_{k \rightarrow j}^{\lambda+e_i} - \left(1 - \frac{1}{\Delta_{ji}^2}\right) \cdot \left(1 - \frac{1}{C_j^\lambda - C_k^\lambda + 1}\right) \cdot D_{k \rightarrow j}^\lambda \cdot D_{k \rightarrow i}^{\lambda+e_j} \quad (\text{Theorem 13.3}) \\
&= D_{k \rightarrow ij}^\lambda - \frac{(\Delta_{ji} - 1)(\Delta_{ji} + 1)}{\Delta_{ji}^2} \cdot \frac{\Delta_{ji}}{\Delta_{ji} + 1} \cdot D_{k \rightarrow ij}^\lambda = \frac{1}{\Delta_{ji}} \cdot D_{k \rightarrow ij}^\lambda.
\end{aligned}$$

The right-hand side is equal to

$$\begin{aligned}
(128) &= \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot D_{k \rightarrow i}^\lambda \cdot \frac{1}{C_j^{\lambda+e_i} - C_{k+1}^{\lambda+e_i}} \cdot D_{k \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.3}) \\
&= \frac{\Delta_{ji} + 1}{\Delta_{ji}} \cdot \frac{1}{C_j^\lambda - C_k^\lambda + 1} \cdot D_{k \rightarrow ij}^\lambda \quad (C_{k+1}^{\lambda+e_i} = \lambda_{k+1} - k - 1 = \lambda_k - k - 1 = C_k^\lambda - 1) \\
&= \frac{\Delta_{ji} + 1}{\Delta_{ji}} \cdot \frac{1}{\Delta_{ji} + 1} D_{k \rightarrow ij}^\lambda = \frac{1}{\Delta_{ji}} \cdot D_{k \rightarrow ij}^\lambda.
\end{aligned}$$

Hence, the two sides are equal.

Case 4: $i < k, j = k$. In that case, $\Delta_{ji} = C_j^{\lambda+e_i} - C_i^\lambda = C_k^\lambda - C_i^\lambda$. The left-hand side can be simplified to

$$\begin{aligned}
(127) &= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) - \frac{1}{\Delta_{ji}^2} \cdot D_{k-1 \rightarrow i}^\lambda \cdot D_{k \rightarrow j}^{\lambda+e_i} \\
&= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) - \frac{1}{\Delta_{ji}^2} \cdot \left(1 - \frac{1}{C_i^\lambda - C_k^\lambda + 1}\right) \cdot D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.3}) \\
&= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + \frac{1}{\Delta_{ji}^2} \cdot \frac{\Delta_{ji}}{1 - \Delta_{ji}} \cdot D_{k \rightarrow ij}^\lambda.
\end{aligned}$$

Since $i < k$, k and $k+1$ are in the same block of $\lambda + e_i$, which simplifies the right-hand side to

$$\begin{aligned}
(128) &= \left(1 + \frac{1}{\Delta_{ji}}\right) \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) \\
&= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + \frac{1}{\Delta_{ji}} \cdot (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot D_{k \rightarrow j}^{\lambda+e_i} \\
&= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + \frac{1}{\Delta_{ji}} \cdot \frac{1}{C_i^\lambda - C_k^\lambda + 1} \cdot D_{k \rightarrow i}^\lambda D_{k \rightarrow j}^{\lambda+e_i} \quad (\text{Theorem 13.3}) \\
&= (D_{k \rightarrow i}^\lambda - D_{k-1 \rightarrow i}^\lambda) \cdot (D_{k \rightarrow j}^{\lambda+e_i} - D_{k-1 \rightarrow j}^{\lambda+e_i}) + \frac{1}{\Delta_{ji}} \cdot \frac{1}{1 - \Delta_{ji}} \cdot D_{k \rightarrow ij}^\lambda.
\end{aligned}$$

Hence, the two sides are equal. $\square$

15 Large- d expansion of the second moment

Recall our expression for the second moment from [Theorem 1.4](#):

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

We have previously argued that Lower_ρ can be written as a positive linear combination of terms of the form $(P \otimes P) \cdot \text{SWAP}$, for some PSD matrices P . This nice structure allowed us to drop the contribution of Lower_ρ in our applications.

In this section, we argue that Lower_ρ consists of terms which are *lower-order* in d . That is, for fixed n and $\lambda \vdash n$, we have

$$\lim_{d \rightarrow \infty} \text{Lower}_\rho = 0.$$

In particular, this allows one to obtain a clean expression for $\mathbf{E}[\widehat{\rho} \otimes \widehat{\rho}]$, by embedding the input state copies $\rho^{\otimes n}$ from a d -dimensional Hilbert space to one of dimension $D \gg d$, such that the lower-order terms vanish:

$$\mathbf{E}[\widehat{\rho} \otimes \widehat{\rho}] \approx \frac{n-1}{n} \cdot \rho \otimes \rho + \frac{1}{n} \cdot (I \otimes \rho + \rho \otimes I) \cdot \text{SWAP} + \frac{\mathbf{E}[\ell(\lambda)]}{n^2} \cdot \text{SWAP}. \quad (130)$$

We show that Lower_ρ only consists of lower-order terms in d by computing the constant-in- d terms of $M_{\text{avg}}^{(2)}$, and demonstrating that they correspond exactly to the terms in [Equation \(130\)](#). The remaining terms are subleading, and must be Lower_ρ .

Following the proof of [Theorem 14.4](#), we can write $M_{\text{avg}}^{(2)}$ as a block-diagonal matrix in the Schur basis, with each block $[M_{\text{avg}}^{(2)}]_{\{i,j\}}^S$ corresponding to an SYT and an unordered pair of indices $S, \{i, j\}$,

$$M_{\text{avg}}^{(2)} = \sum_{\lambda \vdash n} \sum_{i \leq j} |\lambda_{ij}\rangle \langle \lambda_{ij}| \otimes \sum_{S \in \lambda} |S, \{i, j\}\rangle \langle S, \{i, j\}| \otimes (M_{\text{avg}}^{(2)})_{\{i,j\}}^S \otimes I_{\dim(V_{\lambda+e_i+e_j}^d)},$$

with diagonal entry

$$[M_{\text{avg}}^{(2)}]_{ij,ij}^S = \frac{1}{n^2} \cdot \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} \sum_{k,\ell} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2.$$

We use the following formula for the dimension of the space V_λ^d ([Equation \(16\)](#)):

$$\dim(V_\lambda^d) = \prod_{\square \in \lambda} \frac{d + \text{cont}_\lambda(\square)}{h_\lambda(\square)},$$

hence

$$\begin{aligned} \frac{\dim(V_\lambda^d)}{\dim(V_{\lambda_{ij}}^d)} &= \prod_{\square \in \lambda} \frac{h_{\lambda_{ij}}(\square)}{h_\lambda(\square)} \cdot \frac{d + \text{cont}_\lambda(\square)}{d + \text{cont}_{\lambda_{ij}}(\square)} \cdot \prod_{\square \in \lambda_{ij} \setminus \lambda} \frac{h_{\lambda_{ij}}(\square)}{d + \text{cont}_{\lambda_{ij}}(\square)} \\ &= \prod_{\square \in \lambda} \frac{h_{\lambda_{ij}}(\square)}{h_\lambda(\square)} \cdot \prod_{\square \in \lambda_{ij} \setminus \lambda} \frac{h_{\lambda_{ij}}(\square)}{d + \text{cont}_{\lambda_{ij}}(\square)} \\ &= O_n(1) \cdot \frac{O(1)}{\prod_{\square \in \lambda_{ij} \setminus \lambda} (d + \text{cont}_{\lambda_{ij}}(\square))} \\ &= \frac{O_n(1)}{\prod_{\square \in \lambda_{ij} \setminus \lambda} (d + \text{cont}_{\lambda_{ij}}(\square))}. \end{aligned}$$

where $O_n(1)$ is a factor that only depends on n , and not on d . Moreover, we observe that the contents of the two new boxes on rows i and j , $\{\text{cont}_{\lambda_{ij}}(\square)\}_{\square \in \lambda_{ij} \setminus \lambda}$ have magnitude at most n , and thus the denominator is $d^2 + O(d)$. Thus, the only constant-in- d terms that survive in each $(S, \{i, j\})$ block are the $\Theta(d^2)$ terms in the numerator

$$\begin{aligned} &\sum_{k,\ell} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 \\ &= \sum_{k,\ell=1}^{\ell(\lambda)} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{k=1}^{\ell(\lambda)} \sum_{\ell=\ell(\lambda)+1}^d \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{k=\ell(\lambda)+1}^d \sum_{\ell=1}^{\ell(\lambda)} \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{k,\ell=\ell(\lambda)+1}^d \lambda_k^\uparrow \lambda_\ell^\uparrow \cdot |c_{k\ell \rightarrow ij}^\lambda|^2. \end{aligned}$$

In the above equation, we broke down the summation into the non-zero and zero rows of λ , since $\lambda_i = 0$ if and only if $i > \ell(\lambda)$. Since the Clebsch-Gordan coefficients capture overlaps between unit vectors, they have magnitude at most 1. On the other hand, we can write

$$\lambda_i^\uparrow = \begin{cases} \lambda_i - \sum_{\lambda_j > \lambda_i} 1 + \sum_{\lambda_i > \lambda_j > 0} 1 + d - \ell(\lambda), & \lambda_i > 0 \\ -\ell(\lambda), & \lambda_i = 0 \end{cases}.$$

In the first case, only the fourth term depends on d , whereas the remaining terms have magnitude at most n . Moreover, there are at most n rows with $\lambda_i > 0$. On the other hand, there are $d - \ell(\lambda)$ rows with zero length. Hence, the $\Theta(d^2)$ terms of the numerator are included in the following expression:

$$d^2 \cdot \left(\sum_{k,\ell=1}^{\ell(\lambda)} |c_{k\ell \rightarrow ij}^\lambda|^2 \right) - d\ell(\lambda) \cdot \left(\sum_{k=1}^{\ell(\lambda)} \sum_{\ell=\ell(\lambda)+1}^d |c_{k\ell \rightarrow ij}^\lambda|^2 + \sum_{k=\ell(\lambda)+1}^d \sum_{\ell=1}^{\ell(\lambda)} |c_{k\ell \rightarrow ij}^\lambda|^2 \right) + \ell(\lambda)^2 \cdot \left(\sum_{k,\ell=\ell(\lambda)+1}^d |c_{k\ell \rightarrow ij}^\lambda|^2 \right).$$

Using the formulas for the partial sums in [Theorem 14.10](#), we can verify that the expression above is equal to $d^2(X_{n+1}X_{n+2} + \ell(\lambda)/\Delta_{ji}) + O(d)$, which gives rise to the terms in [Equation \(130\)](#).

References

- [Aar18] Scott Aaronson. Shadow tomography of quantum states. In *Proceedings of the 50th Annual ACM Symposium on Theory of Computing*, pages 325–338, 2018.
- [ARS88] Robert Alicki, Sławomir Rudnicki, and Sławomir Sadowski. Symmetry properties of product states for the system of N n -level atoms. *Journal of mathematical physics*, 29(5):1158–1162, 1988.
- [ATD19] Francesco Albarelli, Mankei Tsang, and Animesh Datta. Upper bounds on the holevo cramér-rao bound for multiparameter quantum parametric and semiparametric estimation. Technical report, arXiv:1911.11036, 2019.
- [BC94] Samuel Braunstein and Carlton Caves. Statistical distance and the geometry of quantum states. *Physical Review Letters*, 72(22):3439, 1994.
- [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.
- [Ber12] Sonya Berg. *A quantum algorithm for the quantum Schur-Weyl transform*. PhD thesis, University of California, Davis, 2012.
- [BLM⁺23] Harry Buhrman, Noah Linden, Laura Mančinska, Ashley Montanaro, and Maris Ozols. Quantum majority vote. In *Proceedings of the 14th Innovations in Theoretical Computer Science*, 2023.
- [BO24] Costin Bădescu and Ryan O’Donnell. Improved quantum data analysis. *TheoretCS*, 3, 2024.
- [BOW19] Costin Bădescu, Ryan O’Donnell, and John Wright. Quantum state certification. In *Proceedings of the 51st Annual ACM Symposium on Theory of Computing*, pages 503–514, 2019.
- [BWB24] Adam Bene Watts and John Bostanci. Quantum event learning and gentle random measurements. In *Proceedings of the 15th Innovations in Theoretical Computer Science*, pages 97–1, 2024.
- [Can20] Clément Canonne. A short note on learning discrete distributions, 2020.
- [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.
- [CL24] Sitan Chen and Jerry Li. Personal communication, 2024.
- [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.

- [CM06] Matthias Christandl and Graeme Mitchison. The spectra of quantum states and the Kronecker coefficients of the symmetric group. *Communications in mathematical physics*, 261(3):789–797, 2006.
- [CSDV19] Angelo Carollo, Bernardo Spagnolo, Alexander Dubkov, and Davide Valenti. On quantumness in multi-parameter quantum estimation. *Journal of Statistical Mechanics: Theory and Experiment*, 2019(9):094010, 2019.
- [CYZ25] Kean Chen, Nengkun Yu, and Zhicheng Zhang. Tight bound for quantum unitary time-reversal. Technical report, arXiv:2507.05736, 2025.
- [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.
- [Fis05] Steve Fisk. A very short proof of Cauchy’s interlace theorem for eigenvalues of Hermitian matrices. *Technical report*, arXiv:0502408, 2005.
- [FO24] Steven Flammia and Ryan O’Donnell. Quantum chi-squared tomography and mutual information testing. *Quantum*, 8:1381, 2024.
- [GBO24] Dmitry Grinko, Adam Burchardt, and Maris Ozols. Gelfand-Tsetlin basis for partially transposed permutations, with applications to quantum information. *27th Conference on Quantum Information Processing*, 2024.
- [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.
- [GW09] Roe Goodman and Nolan Wallach. *Symmetry, representations, and invariants*. Springer, 2009.
- [Har05] Aram Harrow. *Applications of coherent classical communication and the Schur transform to quantum information theory*. PhD thesis, Massachusetts Institute of Technology, 2005.
- [Har16] Aram Harrow. Personal communication, 2016.
- [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.
- [HJW18] Yanjun Han, Jiantao Jiao, and Tsachy Weissman. Local moment matching: a unified methodology for symmetric functional estimation and distribution estimation under Wasserstein distance. In *Proceedings of the 31st Annual Conference on Learning Theory*, pages 3189–3221, 2018.
- [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.

- [HM02] Masahito Hayashi and Keiji Matsumoto. Quantum universal variable-length source coding. *Physical Review A*, 66(2):022311, 2002.
- [Hol11] Alexander Holevo. *Probabilistic and statistical aspects of quantum theory*. Springer Science & Business Media, 2011.
- [HW23] Jonas Helsen and Michael Walter. Thrifty shadow estimation: Reusing quantum circuits and bounding tails. *Physical Review Letters*, 131(24):240602, 2023.
- [Key06] Michael Keyl. Quantum state estimation and large deviations. *Reviews in Mathematical Physics*, 18(01):19–60, 2006.
- [KG09] Jonas Kahn and Mădălin Guță. Local asymptotic normality for finite dimensional quantum systems. *Communications in Mathematical Physics*, 289(2):597–652, 2009.
- [KW01] Michael Keyl and Reinhard Werner. Estimating the spectrum of a density operator. *Physical Review A*, 64(5):052311, 2001.
- [LFIC24] Zhaoyi Li, Honghao Fu, Takuya Isogawa, and Isaac Chuang. Optimal quantum purity amplification. Technical report, arXiv:2409.18167, 2024.
- [Mat02] Keiji Matsumoto. A new approach to the Cramér-Rao-type bound of the pure-state model. *Journal of Physics A: Mathematical and General*, 35(13):3111, 2002.
- [NC10] Michael Nielsen and Isaac Chuang. *Quantum computation and quantum information*. Cambridge University Press, 2010.
- [Ngu24] Quynh Nguyen. The mixed Schur transform: efficient quantum circuit and applications. In *27th Conference on Quantum Information Processing*, 2024.
- [OW15a] Ryan O’Donnell and John Wright. A note on the Haah et al. tomography algorithm, 2015. <https://people.eecs.berkeley.edu/~jswright/papers/tomography-note.pdf>.
- [OW15b] 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.
- [OW25] Ryan O’Donnell and Chirag Wadhwa. Instance-optimal quantum state certification with entangled measurements. Technical report, arXiv:2507.06010, 2025.
- [Pil90] Shai Pilpel. Descending subsequences of random permutations. *Journal of Combinatorial Theory, Series A*, 53(1):96–116, 1990.
- [PTTW25] Angelos Pelecanos, Xinyu Tan, Ewin Tang, and John Wright. Beating full state tomography for unentangled spectrum estimation. *Technical report, arXiv:2504.02785*, 2025.
- [Sag01] Bruce E Sagan. *The symmetric group: representations, combinatorial algorithms, and symmetric functions*. Springer, 2001.
- [SSW25] Thilo Scharnhorst, Jack Spilecki, and John Wright. Optimal lower bounds for quantum state tomography. Manuscript, 2025.
- [Sta99] Richard P Stanley. *Enumerative combinatorics Volume 2*. Cambridge University Press, Cambridge, 1999.

- [VK85] Anatoly Vershik and Sergei Kerov. Asymptotic of the largest and the typical dimensions of irreducible representations of a symmetric group. *Functional Analysis and its Applications*, 19(1):21–31, 1985.
- [VK92] Naum Vilenkin and Anatoli Klimy. *Representation of Lie Groups and Special Functions. Volume 3: Classical and Quantum Groups and Special Functions*. Springer, 1992.
- [VN96] Vasilii Voinov and Mikhail Nikulin. *Unbiased estimators and their applications: volume 2: multivariate case*. Springer Science & Business Media, 1996.
- [VN12] Vasilii Voinov and Mikhail Nikulin. *Unbiased estimators and their applications: volume 1: univariate case*. Springer Science & Business Media, 2012.
- [VO05] Anatoli Vershik and Andrei Okounkov. A new approach to the representation theory of the symmetric groups. II. *Journal of Mathematical Sciences*, 131:5471–5494, 2005.
- [VV11] Gregory Valiant and Paul Valiant. Estimating the unseen: an $n/\log(n)$ -sample estimator for entropy and support size, shown optimal via new CLTs. In *Proceedings of the 43rd Annual ACM Symposium on Theory of Computing*, pages 685–694, 2011.
- [VV17] Gregory Valiant and Paul Valiant. Estimating the unseen: improved estimators for entropy and other properties. *Journal of the ACM*, 64(6):1–41, 2017.
- [Wat18] John Watrous. *The theory of quantum information*. Cambridge University Press, 2018.
- [Win02] Andreas Winter. Coding theorem and strong converse for quantum channels. *IEEE Transactions on information theory*, 45(7):2481–2485, 2002.
- [Wri16] John Wright. *How to learn a quantum state*. PhD thesis, Carnegie Mellon University, 2016.
- [Wri24a] John Wright. Lecture 10 from CS 294: Quantum learning theory. Found at <https://people.eecs.berkeley.edu/~jswright/quantumlearningtheory24/subscribe%20notes/lecture10.pdf>, 2024.
- [Wri24b] John Wright. Lecture 12 from CS 294: Quantum learning theory. Found at <https://people.eecs.berkeley.edu/~jswright/quantumlearningtheory24/subscribe%20notes/lecture12.pdf>, 2024.
- [YCH19] Yuxiang Yang, Giulio Chiribella, and Masahito Hayashi. Attaining the ultimate precision limit in quantum state estimation. *Communications in Mathematical Physics*, 368(1):223–293, 2019.
- [YFG13] Koichi Yamagata, Akio Fujiwara, and Richard Gill. Quantum local asymptotic normality based on a new quantum likelihood ratio. *The Annals of Statistics*, 41(4):2197–2217, 2013.
- [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.

Part V

Appendix

A Obtaining Young’s orthogonal basis

The purpose of this section is to show the following theorem.

Theorem A.1 (Theorem 11.3, restated). *The unitary in Theorem 11.2 is a Schur transform. That is, it satisfies*

$$\mathcal{U}_{\text{SW}}^{(n)} \cdot \left(\mathcal{P}^{(n)}(\pi) \mathcal{Q}^{(n)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} = \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\pi) \otimes \nu_\lambda(U), \quad (131)$$

for all $\pi \in S_n$ and $U \in U(d)$.

We also restate the recursive construction of the Schur transform, for the reader's convenience.

Definition A.2 (Schur transform; Theorem 11.2, restated). We define the *Schur transform*, $\mathcal{U}_{\text{SW}}^{(n)} : (\mathbb{C}^d)^{\otimes n} \rightarrow \bigoplus_{\lambda \vdash (n,d)} \text{Sp}_\lambda \otimes V_\lambda^d$, by the following recursive construction. We explicitly define $\mathcal{U}_{\text{SW}}^{(1)}$ as the unitary such that

$$\mathcal{U}_{\text{SW}}^{(1)} |i\rangle := |\square\rangle \otimes |\mathbb{1}\rangle \otimes |i\rangle,$$

for all $i \in [d]$. Then, for $n > 1$,

$$\mathcal{U}_{\text{SW}}^{(n)} := \mathcal{U}_{\text{R}}^{(n)} \cdot \mathcal{U}_{\text{CG}}^{(n)} \cdot (\mathcal{U}_{\text{SW}}^{(n-1)} \otimes I).$$

We start by proving the base case, $n = 1$.

Lemma A.3. *Theorem A.1 holds for $n = 1$.*

Proof. Note that for $n = 1$, $S_n = \{e\}$ so that we just need to check that unitaries are correctly transformed. Moreover, there is only one Young diagram of size 1, $\lambda = \square$, and one SYT of that shape, $S = \mathbb{1}$. We have:

$$\mathcal{U}_{\text{SW}}^{(1)} \cdot U \cdot \mathcal{U}_{\text{SW}}^{(1)\dagger} = \sum_{i,j \in [d]} U_{ij} \cdot \left(\mathcal{U}_{\text{SW}}^{(1)} \cdot |i\rangle\langle j| \cdot \mathcal{U}_{\text{SW}}^{(1)\dagger} \right) = |\square\rangle\langle\square| \otimes |\mathbb{1}\rangle\langle\mathbb{1}| \otimes \sum_{i,j \in [d]} U_{ij} \cdot |i\rangle\langle j| = |\square\rangle\langle\square| \otimes |\mathbb{1}\rangle\langle\mathbb{1}| \otimes U.$$

Since $\nu_{(1)}(U) = U$, Theorem 11.3 holds for $n = 1$. $\square$

We now prove the following partial result, which nevertheless will be a useful stepping stone towards the complete result.

Lemma A.4. *Suppose Theorem A.1 holds for some $n \in \mathbb{N}$. Then $\mathcal{U}_{\text{SW}}^{(n+1)}$ transforms correctly, up to phase. More precisely:*

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(\sigma) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \tilde{\kappa}_\mu(\pi) \otimes \nu_\mu(U),$$

where $\tilde{\kappa}_\mu$ is a representation of S_{n+1} isomorphic to κ_μ , such that $\tilde{\kappa}_\mu \downarrow_{S_n} = \kappa_\mu \downarrow_{S_n}$, and $\tilde{\kappa}_\mu = W_\mu \cdot \kappa_\mu \cdot W_\mu^\dagger$, for some diagonal unitary W_μ .

Proof. First, for all $\sigma \in S_n$ and $U \in U(d)$

$$\begin{aligned} (\mathcal{U}_{\text{SW}}^{(n)} \otimes I) \cdot \left(\mathcal{P}^{(n+1)}(\sigma) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot (\mathcal{U}_{\text{SW}}^{(n)\dagger} \otimes I) &= \left(\mathcal{U}_{\text{SW}}^{(n)} \cdot \left(\mathcal{P}^{(n)}(\sigma) \cdot \mathcal{Q}^{(n,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \right) \otimes U \\ &= \left(\sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma) \otimes \nu_\lambda(U) \right) \otimes U. \end{aligned}$$

Here, we have used that $\mathcal{P}^{(n+1)}(\sigma) = \mathcal{P}^{(n)}(\sigma) \otimes I$, and $\mathcal{Q}^{(n+1)}(U) = \mathcal{Q}^{(n)}(U) \otimes U$. We can then apply the Clebsch-Gordan transform, obtaining the following, where we have defined $\mathcal{V} := \mathcal{U}_{\text{CG}}^{(n+1)} \cdot (\mathcal{U}_{\text{SW}}^{(n)} \otimes I)$:

$$\begin{aligned} \mathcal{V} \cdot \left(\mathcal{P}^{(n+1)}(\sigma) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{V}^\dagger &= \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma) \otimes \left(\mathcal{U}_{\text{CG}}^{(\lambda)} \cdot (\nu_\lambda(U) \otimes U) \cdot \mathcal{U}_{\text{CG}}^{(\lambda)\dagger} \right) \\ &= \sum_{\substack{\lambda \vdash n \\ \ell(\lambda) \leq d}} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma) \otimes \left(\sum_{\substack{\lambda \succ \mu \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \nu_\mu(U) \right). \end{aligned} \quad (132)$$

Finally, applying the rearranging unitary $\mathcal{U}_R^{(n+1)}$, we have:

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(\sigma) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \left(\sum_{\lambda \nearrow \mu} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma) \right) \otimes \nu_\mu(U). \quad (133)$$

Recall that for $\sigma \in S_n$, we have $\kappa_\mu(\sigma) = \sum_{\lambda \nearrow \mu} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma)$, the symmetric group branching rule and the construction of the Young-Yamanouchi basis (see [Section 2.4.3](#)). Thus, $\mathcal{U}_{\text{SW}}^{(n+1)}$ correctly implements the Schur transform for $\sigma \in S_n$. What about $\pi \in S_{n+1}$, more generally? Firstly, we claim

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(\pi) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes X_\mu(\pi) \otimes \nu_\mu(U), \quad (134)$$

for some matrices $\{X_\mu(\pi)\}$. To see this, note that because $\mathcal{P}^{(n+1)}(\pi)$ and $\mathcal{Q}^{(n+1,d)}(U)$ commute for all $U \in U(d)$, the matrix $\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \mathcal{P}^{(n+1)}(\pi) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger}$ commutes with all matrices of the form

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \mathcal{Q}^{(n+1,d)}(U) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \mathcal{U}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(e) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes I_{\dim(\mu)} \otimes \nu_\mu(U). \quad (135)$$

Since the $\{\nu_\mu\}$ are non-isomorphic, irreducible representations, we must have that [[Sag01](#), Theorem 1.7.8, Item (2)]

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \mathcal{P}^{(n+1)}(\pi) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes X_\mu(\pi) \otimes I_{\dim(V_\mu^d)}, \quad (136)$$

for some matrices $\{X_\mu(\pi)\}$. However, these matrices form a representation of S_{n+1} , and by the uniqueness of the decomposition in Schur-Weyl duality, X_μ is isomorphic as a representation to κ_μ . We relabel $X_\mu \rightarrow \tilde{\kappa}_\mu$. Combining [Equations \(135\) and \(136\)](#) gives us [Equation \(134\)](#).

By [Theorem 2.17](#), for each μ there exists a fixed unitary W_μ such that $W_\mu \cdot \tilde{\kappa}_\mu(\pi) \cdot W_\mu^\dagger = \kappa_\mu(\pi)$ for all $\pi \in S_{n+1}$. Restricting to $\sigma \in S_n$, we obtain

$$W_\mu \cdot \left(\sum_{\lambda \nearrow \mu} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma) \right) \cdot W_\mu^\dagger = \sum_{\lambda \nearrow \mu} |\lambda\rangle\langle\lambda| \otimes \kappa_\lambda(\sigma).$$

Since the $\{\kappa_\lambda\}$ are non-isomorphic and irreducible, W_μ , lying in the commutant of the sum, is necessarily constant on each λ -subspace (as in [Equation \(136\)](#), except now the multiplicity space is one-dimensional). Each constant must be a phase, since each λ -subspace is orthogonal, and W_μ is unitary. That is, W_μ is of the form

$$W_\mu = \sum_{\lambda \nearrow \mu} e^{i\theta_{\lambda\mu}} \cdot |\lambda\rangle\langle\lambda| \otimes I_{\dim(\lambda)},$$

for some angles $\{\theta_{\lambda\mu}\}_{\lambda \nearrow \mu}$. Therefore,

$$\mathcal{U}_{\text{SW}}^{(n+1)} \cdot \left(\mathcal{P}^{(n+1)}(\pi) \cdot \mathcal{Q}^{(n+1,d)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(n+1)\dagger} = \sum_{\substack{\mu \vdash n+1 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \tilde{\kappa}_\mu(\pi) \otimes \nu_\mu(U),$$

with $\tilde{\kappa}_\mu$ isomorphic to κ_μ , such that $\tilde{\kappa}_\mu \downarrow_{S_n} = \kappa_\mu \downarrow_{S_n}$, and $\tilde{\kappa}_\mu = W_\mu \cdot \kappa_\mu \cdot W_\mu^\dagger$, for some diagonal unitary W_μ . $\square$

This partial result implies the $n = 2$ case as well.

Corollary A.5. *[Theorem A.1](#) holds for $n = 2$.*

Proof. By [Theorem A.3](#) and [Theorem A.4](#), we know that for all $\pi \in S_2$ and $U \in U(d)$:

$$\mathcal{U}_{\text{SW}}^{(2)} \cdot \left(\mathcal{P}^{(2)}(\pi) \cdot \mathcal{Q}^{(2)}(U) \right) \cdot \mathcal{U}_{\text{SW}}^{(2)\dagger} = \sum_{\substack{\mu \vdash 2 \\ \ell(\mu) \leq d}} |\mu\rangle\langle\mu| \otimes \tilde{\kappa}(\pi) \otimes \nu_\mu(U),$$

where $\tilde{\kappa}_\mu(\pi) = W_\mu \cdot \kappa_\mu(\pi) \cdot W_\mu^\dagger$, for some unitaries W_μ . However, for $|\mu| = 2$, we either have $\mu = (2)$ or $\mu = (1, 1)$. In either case, $\dim(\mu) = 1$, as there is a unique SSYT for any Young diagram with a single row, or with a single column. Since all matrices of size 1 commute:

$$\tilde{\kappa}_\mu(\pi) = W_\mu \cdot \kappa_\mu(\pi) \cdot W_\mu^\dagger = \kappa_\mu(\pi). \quad \square$$

We are now ready to prove [Theorem 11.3](#).

Proof of Theorem 11.3. The proof will follow by induction on the number of qudits n . In particular, we show that if [Theorem A.1](#) holds for n and $n + 1$ qudits, then it also holds for $n + 2$. Given the base cases of $n = 1$ ([Theorem A.3](#)) and $n = 2$ ([Theorem A.5](#)), this will show the result for all $n \geq 1$.

Given the inductive hypothesis, [Theorem A.4](#), states that for $\mu \vdash n + 2$,

$$\tilde{\kappa}_\mu \downarrow_{S_{n+1}} = \kappa_\mu \downarrow_{S_{n+1}}.$$

This implies that for any permutation $\pi \in S_{n+2}$ that satisfies $\pi(n + 2) = n + 2$, it holds that $\tilde{\kappa}_\mu(\pi) = \kappa_\mu(\pi)$. Thus, since the set of permutations $\{(n + 1, n + 2)\} \cup \{\sigma \in S_{n+2} \mid \sigma(n + 2) = n + 2\}$ generates the entirety of S_{n+2} , it suffices to prove that

$$\tilde{\kappa}_\mu(n + 1, n + 2) = \kappa_\mu(n + 1, n + 2),$$

to conclude that $\tilde{\kappa}_\mu(\pi) = \kappa_\mu(\pi)$ for all π .

We restrict our attention to the permutation $(n + 1, n + 2)$. Let $\lambda \vdash n$ be a Young diagram. We know that both $\kappa_{\lambda_{ij}}(n + 1, n + 2)$ and $\tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2)$ have the same block-diagonal structure in Young's orthogonal basis. In particular, each matrix has a 2×2 submatrix whenever (S_{ij}, S_{ji}) are both valid, and a 1×1 submatrix whenever only one is valid. Here, we use S_{ij} to denote the SYT we obtain when we start from a valid SYT S with n boxes, and add an $n + 1$ at the end of row i , and an $n + 2$ at the end of row j .

[Theorem A.4](#) states that

$$\tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2) = W_{\lambda_{ij}} \cdot \kappa_{\lambda_{ij}}(n + 1, n + 2) \cdot W_{\lambda_{ij}}^\dagger,$$

for $W_{\lambda_{ij}}^\dagger$ a diagonal unitary. Since matrices of size 1 commute, we can deduce that $\tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2)$ and $\kappa_{\lambda_{ij}}(n + 1, n + 2)$ have the same 1×1 submatrices in that basis.

We now focus on a 2×2 submatrix indexed by the SYTs S_{ij} and S_{ji} . Let $e^{i\theta_1}, e^{i\theta_2}$ be the entries of $W_{\lambda_{ij}}^\dagger$ on $|S_{ij}\rangle\langle S_{ij}|$ and $|S_{ji}\rangle\langle S_{ji}|$ respectively. Then $\tilde{\kappa}_{\lambda_{ij}}$ and $\kappa_{\lambda_{ij}}$ agree on the diagonal entries of the submatrix, since

$$\langle S_{ij} | \tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2) | S_{ij} \rangle = e^{i\theta_1} \cdot \langle S_{ij} | \kappa_{\lambda_{ij}}(n + 1, n + 2) | S_{ij} \rangle \cdot e^{-i\theta_1} = \langle S_{ij} | \kappa_{\lambda_{ij}}(n + 1, n + 2) | S_{ij} \rangle,$$

and similarly for S_{ji} . On the other hand, the off-diagonal entries satisfy

$$\begin{aligned} \langle S_{ij} | \tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2) | S_{ji} \rangle &= e^{i\theta_1} \cdot \langle S_{ij} | \kappa_{\lambda_{ij}}(n + 1, n + 2) | S_{ji} \rangle \cdot e^{-i\theta_2} \\ &= e^{i(\theta_1 - \theta_2)} \cdot \langle S_{ij} | \kappa_{\lambda_{ij}}(n + 1, n + 2) | S_{ji} \rangle. \end{aligned}$$

Since the off-diagonal entries of $\kappa_{\lambda_{ij}}(n + 1, n + 2)$ are given by the expression³

$$\langle S_{ij} | \kappa_{\lambda_{ij}}(n + 1, n + 2) | S_{ji} \rangle = \sqrt{1 - \frac{1}{\Delta_{ji}^2}},$$

it suffices to show that $\langle S_{ij} | \tilde{\kappa}_{\lambda_{ij}}(n + 1, n + 2) | S_{ji} \rangle$ is real and positive for all pairs of SYTs of the form (S_{ij}, S_{ji}) .

³Recall that Δ_{ji} is the difference in contents of the boxes with $n + 2$ and $n + 1$ in S_{ij} .

The expression $\langle S_{ij} | \tilde{\kappa}_{\lambda_{ij}}(n+1, n+2) | S_{ji} \rangle$ can be written to involve the SWAP = $\mathcal{P}(n+1, n+2)$ operator on qudits $n+1$ and $n+2$, by noting that

$$\langle S_{ij} | \tilde{\kappa}_{\lambda_{ij}}(n+1, n+2) | S_{ji} \rangle = \langle \lambda_{ij}, S_{ij}, T | \cdot \mathcal{U}_{\text{SW}}^{(n+2)} \cdot \mathcal{P}(n+1, n+2) \cdot \mathcal{U}_{\text{SW}}^{(n+2)\dagger} \cdot | \lambda_{ij}, S_{ji}, T \rangle, \quad (137)$$

for any $T \in \text{SSYT}(\lambda_{ij})$. The vectors on the left and right of $\mathcal{P}(n+1, n+2)$ are written in “our” Schur basis, which contains a Young-Yamanouchi basis in the permutation register that is a priori related by phases to Young’s orthogonal basis. We will now “undo” our Schur transform to obtain the Schur basis on n qudits, which, by induction, we will assume is the true Schur basis. In particular, we follow the original steps from [Section 11](#), in reverse. As each step is unitary, we can undo each step by applying the conjugate transpose of the relevant unitary. We have

$$\begin{aligned} | \lambda_{ij}, S_{ji}, T \rangle &= \mathcal{U}_{\text{R}}^{(n+2)} \cdot | \lambda + e_j, S_j, \lambda_{ij}, T \rangle \\ &= \mathcal{U}_{\text{R}}^{(n+2)} \cdot \mathcal{U}_{\text{CG}}^{(n+2)} \cdot \left(\sum_k \sum_{T' \in \lambda + e_j} \langle T', k | T \rangle \cdot | \lambda + e_j, S_j, T' \rangle \otimes | k \rangle \right) \\ &= \mathcal{U}_{\text{R}}^{(n+2)} \cdot \mathcal{U}_{\text{CG}}^{(n+2)} \cdot \mathcal{U}_{\text{R}}^{(n+1)} \cdot \left(\sum_k \sum_{T' \in \lambda + e_j} \langle T', k | T \rangle \cdot | \lambda, S, \lambda + e_j, T' \rangle \otimes | k \rangle \right) \\ &= \mathcal{U}_{\text{R}}^{(n+2)} \cdot \mathcal{U}_{\text{CG}}^{(n+2)} \cdot \mathcal{U}_{\text{R}}^{(n+1)} \cdot \mathcal{U}_{\text{CG}}^{(n+1)} \cdot \left(\sum_{k, \ell} \sum_{T' \in \lambda + e_j} \sum_{T'' \in \lambda} \langle T', k | T \rangle \cdot \langle T'', \ell | T' \rangle \cdot | \lambda, S, T'' \rangle \otimes | \ell \rangle \otimes | k \rangle \right) \\ &= \mathcal{U}_{\text{SW}}^{(n+2)} \cdot (\mathcal{U}_{\text{SW}}^{(n)\dagger} \otimes I \otimes I) \cdot \left(\sum_{k, \ell} \sum_{T' \in \lambda + e_j} \sum_{T'' \in \lambda} \langle T', k | T \rangle \cdot \langle T'', \ell | T' \rangle \cdot | \lambda, S, T'' \rangle \otimes | \ell \rangle \otimes | k \rangle \right). \end{aligned}$$

Throughout this section, we use $T \in \lambda$ to denote in a succinct way that the SSYT T has shape λ . Similarly, we can write

$$| \lambda_{ij}, S_{ij}, T \rangle = \mathcal{U}_{\text{SW}}^{(n+2)} \cdot (\mathcal{U}_{\text{SW}}^{(n)\dagger} \otimes I \otimes I) \cdot \left(\sum_{k', \ell'} \sum_{Y' \in \lambda + e_i} \sum_{Y'' \in \lambda} \langle Y', \ell' | T \rangle \cdot \langle Y'', k' | Y' \rangle \cdot | \lambda, S, Y'' \rangle \otimes | k' \rangle \otimes | \ell' \rangle \right).$$

Then the off-diagonal entry that we are trying to compute is equal to

$$\begin{aligned} &\sum_{k, \ell, k', \ell'} \sum_{\substack{T'' \in \lambda \\ T' \in \lambda + e_j}} \sum_{\substack{Y'' \in \lambda \\ Y' \in \lambda + e_i}} \langle Y'', k' | Y' \rangle \langle Y', \ell' | T \rangle \langle T'', \ell | T' \rangle \langle T', k | T \rangle \cdot \langle \lambda, S, T'', \ell, k | \cdot \mathcal{U}_{\text{SW}}^{(n)} \cdot \text{SWAP} \cdot \mathcal{U}_{\text{SW}}^{(n)\dagger} \cdot | \lambda, S, Y'', k', \ell' \rangle \\ &= \sum_{k, \ell, k', \ell'} \sum_{\substack{T'' \in \lambda \\ T' \in \lambda + e_j}} \sum_{\substack{Y'' \in \lambda \\ Y' \in \lambda + e_i}} \langle Y'', k' | Y' \rangle \langle Y', \ell' | T \rangle \langle T'', \ell | T' \rangle \langle T', k | T \rangle \cdot \langle \lambda, S, T'', \ell, k | \lambda, S, Y'', \ell', k' \rangle \\ &= \sum_{k, \ell} \sum_{\substack{T'' \in \lambda \\ T' \in \lambda + e_j}} \sum_{\substack{Y'' \in \lambda \\ Y' \in \lambda + e_i}} \langle Y'', k | Y' \rangle \langle Y', \ell | T \rangle \langle T'', \ell | T' \rangle \langle T', k | T \rangle \cdot \langle T'' | Y'' \rangle, \end{aligned} \quad (138)$$

where the last equality follows because each term is zero unless $(\ell, k) = (\ell', k')$. Since [Equation \(137\)](#) holds for any SSYT T , we will choose T to make the above summation easier to compute (its sign, at least). In particular, the one-step and two-step Clebsch-Gordan coefficients that we have developed so far gave us a solid understanding of how a highest-weight SSYT changes when we tensor onto it one or two elements of the fundamental representation $V_{(1)}^d$, like $|k\rangle$ and $|\ell\rangle$. Here, we are instead starting from an SSYT T , and want to understand which SSYT T'' and Y'' can take us to T when we tensor on $|k\rangle$ ’s and $|\ell\rangle$ ’s. Pictorially, the SSYT in [Equation \(138\)](#) satisfy the following:

$$\underbrace{T''}_{\lambda} \xrightarrow{\otimes |\ell\rangle} \underbrace{T'}_{\lambda + e_j} \xrightarrow{\otimes |k\rangle} \underbrace{T}_{\lambda_{ij}} \xleftarrow{\otimes |\ell\rangle} \underbrace{Y'}_{\lambda + e_i} \xleftarrow{\otimes |k\rangle} \underbrace{Y''}_{\lambda}$$

Since T'' and Y'' are not necessarily highest-weight SSYT, it feels that our work from the previous section is not helpful here. Fortunately, we can relate the Clebsch-Gordan coefficients that arise in [Equation \(137\)](#) to the one-step and two-step Clebsch-Gordan coefficients, by setting T to be the *lowest-weight* SSYT of shape λ_{ij} , and then considering the *complement* tableaux of T, T', T'' , and Y', Y'' . We introduce these notions below.

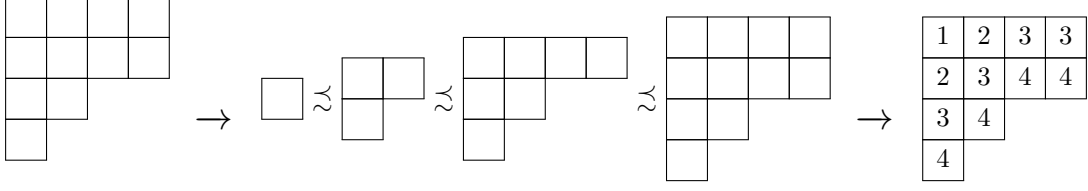


Figure 24: An example illustrating lowest-weight SSYTs. Take $d = 4$. On the left, we have the Young diagram $\lambda = (4, 4, 2, 1)$. In the middle, we have the sequence $T_\lambda^{\leq p}$, for $p \in [d]$. Note that $T_\lambda^{\leq p}$ consists of the lowest p rows of λ , shifted upwards by $(d - p)$. Thus, $T_\lambda^{\leq p}$ contains the cells of λ directly above at least $(d - p)$ boxes. On the right is T_λ , the lowest-weight SSYT, corresponding to the chain $T_\lambda^{\leq 1} \preceq \cdots \preceq T_\lambda^{\leq d}$. The boxes filled with symbol p are the boxes in λ that are above exactly $(d - p)$ boxes. This SSYT can also be obtained by filling the lowest box in each column with d , then the lowest remaining box in each column with $(d - 1)$, and so on.

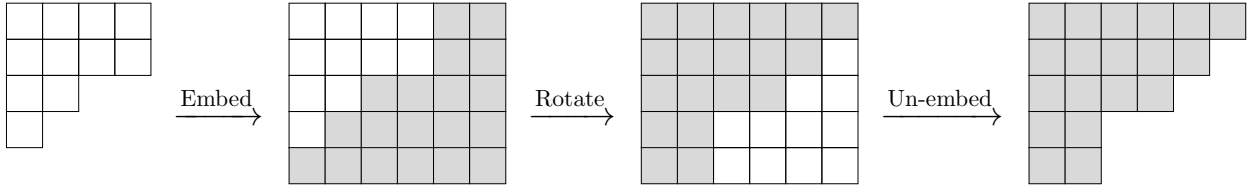


Figure 25: An example illustrating complementary Young diagrams. Take $d = 5$, and $m = 6$. On the left, we have the Young diagram $\lambda = (4, 4, 2, 1, 0)$. On the right, we have the Young diagram $\mu = \text{compl}_{(d,m)}(\lambda) = (6, 5, 4, 2, 2)$. In between, we illustrate how we can obtain μ from λ by first “embedding” λ into a $d \times m$ grid of boxes, then rotating the diagram 180° , and finally “un-embedding” the shaded boxes, which are the boxes that were not originally in λ .

Definition A.6 (Lowest-weight SSYTs). Let $\lambda = (\lambda_1, \dots, \lambda_d) \vdash n$ be a Young diagram. The *lowest weight SSYT* of shape λ and alphabet $[d]$ is denoted by T_λ and satisfies:

$$T_\lambda^{\leq p} = (\lambda_{d-p+1}, \dots, \lambda_d).$$

The *lowest-weight vector* in V_λ^d is the vector $|T_\lambda\rangle$.

This formal definition can be used to show straightforwardly that T_λ is, in fact, an SSYT, using the interlacing property described in Section 2.4.4. Indeed, note that $(T_\lambda)_i^{\leq p} = \lambda_{d-p+i}$. Then $T_\lambda^{\leq p} \preceq T_\lambda^{\leq (p+1)}$ since $(T_\lambda)_{i+1}^{\leq (p+1)} \leq (T_\lambda)_i^{\leq p} \leq (T_\lambda)_i^{\leq (p+1)}$, which holds because $\lambda_{d-(p+1)+(i+1)} = \lambda_{d-p+i} \leq \lambda_{d-(p+1)+i}$, as the rows of λ are weakly decreasing. However, we can also give a more intuitive picture.

Note that the rows of $T_\lambda^{\leq p}$ are the same as the bottom p rows of λ . As illustrated in Figure 24, this means that a box in T_λ is filled with the symbol p when it is above exactly $(d - p)$ boxes. Thus, the lowest-weight SSYT has a simple, alternate characterization: T_λ is the SSYT of shape λ obtained by filling the lowest box in each column with the value d , and then repeating this procedure with the unfilled boxes using the values $d - 1, \dots, 1$. Intuitively, T_λ is the SSYT of shape λ with the maximum number of d 's, then the maximum number of $(d - 1)$'s, and so on.

Definition A.7 (Complementary Young diagrams). Let $\lambda \in \mathbb{Z}^d$ be a Young diagram, and let m be an integer at least as large as the number of columns in λ . The Young diagram which *complements* λ with respect to m columns, $\mu \in \mathbb{Z}^d$, is the Young diagram such that

$$\mu_i = m - \lambda_{d+1-i}.$$

We write $\mu =: \text{compl}_{(d,m)}(\lambda)$, though we will drop (d, m) when these parameters are clear. We also say μ is the *complement* of λ , or μ is *complementary* to λ .

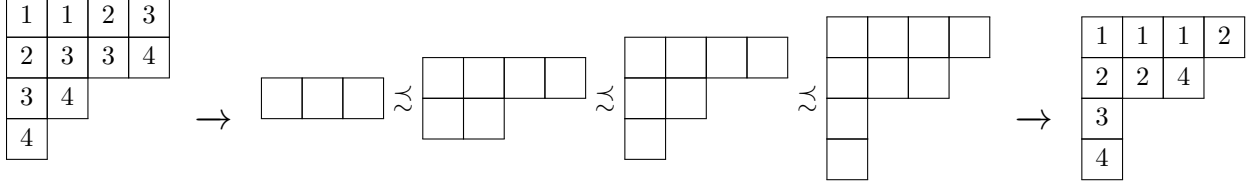


Figure 26: An example illustrating complementing SSYT. Take $d = 4$ and $m = 5$. On the left, we have an SSYT S of shape $\lambda = (4, 4, 2, 1)$ and alphabet $[d]$. In the middle, we have the sequence $\text{compl}_{(p,m)}(S^{\leq p})$, for $p \in [d]$. On the right, we have the SSYT $T = \text{compl}_{(d,m)}(S)$ corresponding to this chain.

We note that $\text{compl}(\lambda)$ is, indeed, a Young diagram, since $\mu_{i+1} = m - \lambda_{p-i} \leq m - \lambda_{p+1-i} = \mu_i$. The complementary Young diagram can be obtained by a three-step geometric process illustrated in Figure 25. It is also clear from the figure that $\text{compl}_{(d,m)}(\text{compl}_{(d,m)}(\lambda)) = \lambda$, so that complementation with fixed parameters is bijective.

For S an SSYT, we can consider the complements of the Young diagrams $S^{\leq p} \in \mathbb{Z}^p$, with respect to a common value of m , for $p = 1, \dots, d$.

Lemma A.8. *Let S be an SSYT with alphabet $[d]$. Let m be an integer at least as large as the number of columns of S . For each $p \in [d]$, construct the complementary tableau $\text{compl}_{(p,m)}(S^{\leq p}) \in \mathbb{Z}^p$. Then, for each $p \in [d-1]$, $\text{compl}_{(p,m)}(S^{\leq p}) \preceq \text{compl}_{(p+1,m)}(S^{\leq p+1})$.*

Proof. Let $A, B \in \mathbb{Z}^d$ be Young diagrams. Recall that $A \preceq B$ iff the row-lengths of B interlace the row-lengths of A : $A_{i+1} \leq B_{i+1} \leq A_i$ for all $i \in [d-1]$.

Since S is a SSYT, we have $S^{\leq p} \preceq S^{\leq (p+1)}$ for each $p \in [d-1]$. Thus, for each $p \in [d-1]$ and $j \in [p-1]$, we have $(S^{\leq p})_{j+1} \leq (S^{\leq (p+1)})_{j+1} \leq (S^{\leq p})_j$. Now, from $(S^{\leq p})_{j+1} \leq (S^{\leq (p+1)})_{j+1}$ with $j = p-i$, we have

$$\text{compl}(S^{\leq (p+1)})_{i+1} = m - (S^{\leq (p+1)})_{p+1-i} \leq m - (S^{\leq p})_{p+1-i} = \text{compl}(S^{\leq p})_i. \quad (139)$$

Moreover, from $(S^{\leq (p+1)})_{j+1} \leq (S^{\leq p})_j$ with $j = p+1-i$, we have

$$\text{compl}(S^{\leq p})_i = m - (S^{\leq p})_{p+1-i} \leq m - (S^{\leq (p+1)})_{p+2-i} = \text{compl}(S^{\leq (p+1)})_i. \quad (140)$$

Combining Equations (139) and (140) shows the interlacing condition:

$$\text{compl}(S^{\leq p})_{i+1} \leq \text{compl}(S^{\leq p})_i \leq \text{compl}(S^{\leq (p+1)})_i.$$

We conclude that $\text{compl}(S^{\leq p}) \preceq \text{compl}(S^{\leq (p+1)})$. □

Thus, $\text{compl}_{(1,m)}(S^{\leq 1}) \preceq \dots \preceq \text{compl}_{(d,m)}(S^{\leq d})$. This chain is associated with an SSYT with alphabet $[d]$, by the correspondence described in Section 2.4.4.

Corollary A.9. *Let S be an SSYT with alphabet $[d]$, and let m be an integer at least as large as the number of columns of S . Then, there exists a unique SSYT with alphabet $[d]$, T , such that, for each $p \in [d]$, $T^{\leq p} = \text{compl}_{(p,m)}(S^{\leq p})$.*

Definition A.10 (Complementary SSYT). Let S be an SSYT, and let m be an integer larger than the number of columns in S . The SSYT which *complements* S to be the SSYT T that satisfies

$$T^{\leq p} = \text{compl}_{(p,m)}(S^{\leq p}).$$

Thus, $T_i^{\leq p} = m - S_{p+1-i}^{\leq p}$. We write $T = \text{compl}_{(d,m)}(S)$, dropping (d, m) when these parameters are clear from context. We also T is the *complement* of S , or T is *complementary* to S .

Note that if $T = \text{compl}_{(d,m)}(S)$, then $\text{compl}_{(p,m)}(T^{\leq p}) = S^{\leq p}$ for all $p \in [d]$, so that $S = \text{compl}_{(d,m)}(T)$ as well. That is, $S = \text{compl}(\text{compl}(S))$. See Figure 26 for an example of complementary SSYT.

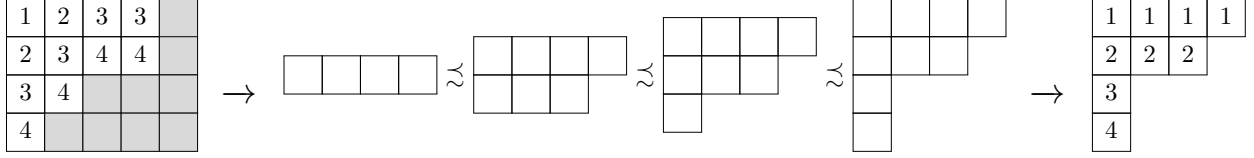


Figure 27: Another example of complementing SSYT, here illustrating how the complements of lowest-weight SSYT are highest-weight SSYT. As in the previous figure, take $d = 4$, $m = 5$, and $\lambda = (4, 4, 2, 1)$. On the left, we have the lowest-weight SSYT S_λ of shape λ and alphabet $[d]$, which we have depicted as embedded into a $d \times m$ grid. The gray cells are those that do not belong to λ . In the middle, we have the sequence $\text{compl}_{(p,m)}(S_\lambda^{\leq p})$, for $p \in [d]$. Note that $\text{compl}_{(p,m)}(S_\lambda^{\leq p})$ is the complement of the bottom p rows, regarded as a Young diagram. Thus, $\text{compl}_{(p,m)}(S_\lambda^{\leq p})$ consists of the gray cells in the bottom p rows of λ , rotated 180° . Therefore, the p -th element in the chain adds boxes only in the p -th row. This means that the SSYT corresponding to this chain, $T^\lambda = \text{compl}_{(d,m)}(S)$ is highest-weight. On the right is T^λ , and indeed it is highest weight.

Lemma A.11. *Let $\lambda \in \mathbb{Z}^d$ be a Young diagram, and let m be an integer at least as large as the number of columns of λ . Let $\mu := \text{compl}_{(d,m)}(\lambda)$. Let S_λ be the lowest-weight SSYT of shape λ , and let T^μ be the highest-weight SSYT of shape μ , both with alphabet $[d]$. Then S_λ and T^μ are complementary SSYT.*

It is perhaps easiest to understand this lemma pictorially; see Figure 27. We now give a formal argument.

Proof. Recall $(S_\lambda^{\leq p})_i = \lambda_{d-p+i}$. Then

$$\text{compl}_{(p,m)}(S_\lambda^{\leq p})_i = m - (S_\lambda^{\leq p})_{p+1-i} = m - \lambda_{d-p+(p+1-i)} = m - \lambda_{d+1-i},$$

so that $\text{compl}_{(p,m)}(S_\lambda^{\leq p}) = (m - \lambda_d, \dots, m - \lambda_{d+1-p})$. Since $\text{compl}_{(p+1,m)}(S_\lambda^{\leq (p+1)})_i = \text{compl}_{(p,m)}(S_\lambda^{\leq p})_i$ for $i \in [p]$, the only boxes in $\text{compl}_{(p+1,m)}(S_\lambda^{\leq (p+1)}) \setminus \text{compl}_{(p,m)}(S_\lambda^{\leq p})$ are those in the $(p+1)$ -th row. Therefore, all boxes labelled p in $\text{compl}_{(d,m)}(S_\lambda)$ occur in the p -th row, and therefore this SSYT is highest-weight. Since $\text{compl}_{(d,m)}(S_\lambda)$ is an SSYT of shape $\text{compl}_{(d,m)}(S_\lambda^{\leq d}) = \text{compl}_{(d,m)}(\lambda) = \mu$, we have $\text{compl}_{(d,m)}(S_\lambda) = T^\mu$. $\square$

The reason why we set T to be the lowest-weight SSYT, is because its complement is a highest-weight SSYT, and the following lemma allows us to relate the Clebsch-Gordan coefficients in Equation (138) to the one- and two-step Clebsch-Gordan coefficients.

Lemma A.12. *Let S, S' be SSYT with shapes λ and $\lambda + e_i$ respectively, and m be an integer larger than the number of columns in both S and S' . Moreover, let $T = \text{compl}(S')$ and $T' = \text{compl}(S)$. Then, it holds that*

$$\langle S, k | S' \rangle = (-1)^{d-k} \cdot \left(\frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_\lambda^d)} \right)^{1/2} \cdot \langle T, k | T' \rangle$$

The proof of this lemma is deferred to the end of this section. We now use this lemma to rewrite the coefficients from Equation (138) in the following way:

$$\begin{aligned} \langle Y'', k | Y' \rangle &= (-1)^{d-k} \left(\frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_\lambda^d)} \right)^{1/2} \langle \widehat{Y}', k | \widehat{Y}'' \rangle, & \langle Y', \ell | T \rangle &= (-1)^{d-\ell} \left(\frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_\lambda^d)} \right)^{1/2} \langle \widehat{T}, \ell | \widehat{Y}' \rangle, \\ \langle T'', \ell | T' \rangle &= (-1)^{d-\ell} \left(\frac{\dim(V_{\lambda+e_j}^d)}{\dim(V_\lambda^d)} \right)^{1/2} \langle \widehat{T}', \ell | \widehat{T}'' \rangle, & \langle T', k | T \rangle &= (-1)^{d-k} \left(\frac{\dim(V_{\lambda+e_j}^d)}{\dim(V_\lambda^d)} \right)^{1/2} \langle \widehat{T}, k | \widehat{T}' \rangle. \end{aligned}$$

For brevity, we use $\widehat{T}$ to mean $\text{compl}(T)$, and similarly for the SSYT T', T'', Y', Y'' . Moreover, since

complementation is injective (as shown above), $\langle T'' | Y'' \rangle = \langle \widehat{T}'' | \widehat{Y}'' \rangle$. Hence Equation (138) becomes

$$\begin{aligned}
& \sum_{k,\ell} \sum_{\substack{T'' \in \lambda \\ T' \in \lambda + e_j}} \sum_{\substack{Y'' \in \lambda \\ Y' \in \lambda + e_i}} \langle Y'', k | Y' \rangle \langle Y', \ell | T \rangle \langle T'', \ell | T' \rangle \langle T', k | T \rangle \cdot \langle T'' | Y'' \rangle \\
&= \sum_{k,\ell} (-1)^{d-k} (-1)^{d-k} (-1)^{d-\ell} (-1)^{d-\ell} \cdot \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_\lambda^d)} \\
&\cdot \sum_{\substack{\widehat{T}'' \in \text{compl}(\lambda) \\ \widehat{T}' \in \text{compl}(\lambda + e_j)}} \sum_{\substack{\widehat{Y}'' \in \text{compl}(\lambda) \\ \widehat{Y}' \in \text{compl}(\lambda + e_i)}} \langle \widehat{Y}'', k | \widehat{Y}' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle \langle \widehat{T}', \ell | \widehat{T}'' \rangle \langle \widehat{T}, k | \widehat{T}' \rangle \cdot \langle \widehat{T}'' | \widehat{Y}'' \rangle \\
&= \frac{\dim(V_{\lambda_{ij}}^d)}{\dim(V_\lambda^d)} \sum_{k,\ell} \sum_{\substack{\widehat{T}'' \in \text{compl}(\lambda) \\ \widehat{T}' \in \text{compl}(\lambda + e_j)}} \sum_{\substack{\widehat{Y}'' \in \text{compl}(\lambda) \\ \widehat{Y}' \in \text{compl}(\lambda + e_i)}} \langle \widehat{Y}'', k | \widehat{Y}' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle \langle \widehat{T}', \ell | \widehat{T}'' \rangle \langle \widehat{T}, k | \widehat{T}' \rangle \cdot \langle \widehat{T}'' | \widehat{Y}'' \rangle. \quad (141)
\end{aligned}$$

Now the nonzero terms in the summation above correspond to SSYT's that look as follows:

$$\begin{array}{ccccccc}
\widehat{T}'' & \xleftarrow{\otimes|\ell|} & \widehat{T}' & \xleftarrow{\otimes|k|} & \widehat{T} & \xrightarrow{\otimes|\ell|} & \widehat{Y}' & \xrightarrow{\otimes|k|} & \widehat{Y}'' \\
\mu + e_{i'} + e_{j'} & & \mu + e_{i'} & & \mu & & \mu + e_{j'} & & \mu + e_{i'} + e_{j'}
\end{array}$$

where $\widehat{T}$ is a highest-weight SSYT of shape $\mu = \text{compl}(\lambda_{ij})$. We set $i' = d + 1 - i$, $j' = d + 1 - j$, and we used the fact that $\text{compl}(\mu + e_{i'}) = \text{compl}(\mu) - e_i = \lambda + e_j$. Thus, now the terms that appear correspond to the coefficients we get when we add boxes to a highest-weight SSYT, something that we understand much better.

It suffices to show that Equation (141) is positive. The factor before the summation is always positive, and can be thus ignored. Then, we observe that the terms that appear in this expression are exactly the two-step Clebsch Gordan coefficients:

$$\langle \widehat{T}', \ell | \widehat{T}'' \rangle \langle \widehat{T}, k | \widehat{T}' \rangle = \begin{cases} a_{k\ell \rightarrow i'j'}^\mu, & \text{or} \\ b_{k\ell \rightarrow i'j'}^\mu \end{cases} \quad \text{and} \quad \langle \widehat{Y}', k | \widehat{Y}'' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle = \begin{cases} a_{\ell k \rightarrow j'i'}^\mu, & \text{or} \\ b_{\ell k \rightarrow j'i'}^\mu \end{cases}$$

From Theorem 10.18, we know that the a coefficients are always nonnegative. Thus, any negative terms in Equation (141) must come when at least one of the terms is a b coefficient.

We consider the two cases. First, let $\langle \widehat{T}', \ell | \widehat{T}'' \rangle \langle \widehat{T}, k | \widehat{T}' \rangle = b_{k\ell \rightarrow i'j'}^\mu$ have value strictly less than zero. From Theorem 10.16, this must mean that $k > \ell$, and that $\widehat{T}'' = \widehat{T}_{k\ell \rightarrow j'i'}$, that is, $\widehat{T}''$ has a k at the end of the j' -th row, and an ℓ at the end of the i' -th row. Since $k > \ell$, $b_{k\ell \rightarrow j'i'}^\mu = 0$, and thus the other term $\langle \widehat{Y}', k | \widehat{Y}'' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle$ is either 0, or $a_{\ell k \rightarrow j'i'}^\mu$. When $\langle \widehat{Y}', k | \widehat{Y}'' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle = a_{\ell k \rightarrow j'i'}^\mu$, this must mean that $\widehat{Y}'' = \widehat{T}_{\ell k \rightarrow j'i'}$, and thus $\widehat{Y}''$ has an ℓ at the end of the j' -th row, and a k at the end of the i' -th row. Since $i' \neq j'$, and $k > \ell$, the SSYT's $\widehat{Y}''$ and $\widehat{T}''$ are different, and thus $\langle \widehat{T}'' | \widehat{Y}'' \rangle = 0$. Hence, we can only have positive or zero contributions to Equation (141) in that case.

We proceed in the same fashion in the second case, when we let $\langle \widehat{Y}', k | \widehat{Y}'' \rangle \langle \widehat{T}, \ell | \widehat{Y}' \rangle = b_{\ell k \rightarrow j'i'}^\mu$ have value strictly less than zero. From Theorem 10.16, this must mean that $\ell > k$, and that $\widehat{Y}'' = \widehat{Y}_{\ell k \rightarrow i'j'}$, that is, $\widehat{Y}''$ has an ℓ at the end of the i' -th row, and a k at the end of the j' -th row. Since $\ell > k$, $b_{\ell k \rightarrow i'j'}^\mu = 0$, and thus the other term $\langle \widehat{T}', k | \widehat{T}'' \rangle \langle \widehat{T}, \ell | \widehat{T}' \rangle$ is either 0, or $a_{k\ell \rightarrow i'j'}^\mu$. When $\langle \widehat{T}', k | \widehat{T}'' \rangle \langle \widehat{T}, \ell | \widehat{T}' \rangle = a_{k\ell \rightarrow i'j'}^\mu$, this must mean that $\widehat{T}'' = \widehat{T}_{k\ell \rightarrow i'j'}$, and thus $\widehat{T}''$ has a k at the end of the i' -th row, and an ℓ at the end of the j' -th row. Since $i' \neq j'$, and $\ell > k$, the SSYT's $\widehat{T}''$ and $\widehat{Y}''$ are different, and thus $\langle \widehat{T}'' | \widehat{Y}'' \rangle = 0$. Hence, we can only have positive or zero contributions to Equation (141) in that case as well.

We conclude that this off-diagonal term must be real and positive, and thus equal to $\sqrt{1 - \frac{1}{\Delta_{ji}^2}}$. $\square$

Proof of Theorem A.12. As per Equation (58), we decompose the left-hand side as a product of scalar factors

$$\langle S, k | S' \rangle = \prod_{p=1}^d (S^{=p}, \overline{k}^{=p} | (S')^{=p}) = (S^{=k}, e_{10}^{=k} | (S')^{=k}) \cdot \prod_{p=k+1}^d (S^{=p}, e_{11}^{=p} | (S')^{=p}). \quad (142)$$

As in Equation (63), the only way for the Clebsch-Gordan coefficient $\langle S, k | S' \rangle$ to be nonzero, is if there exist indices $i_k, \dots, i_d$, such that $i_d = i$, and

$$(S')^{\leq p} = \begin{cases} S^{\leq p} & p < k \\ S_{+i_p}^{\leq p} & p = k \\ S_{+i_p, -i_{p-1}}^{\leq p} & p > k \end{cases} \implies (S')^{\leq p} = \begin{cases} S^{\leq p} & p < k \\ S^{\leq p} + e_{i_p} & p \geq k \end{cases}.$$

Taking the complements of the tableaux in the right-hand side expression, we deduce that

$$(T')^{\leq p} = \begin{cases} T^{\leq p} & p < k \\ T^{\leq p} + e_{j_p} & p \geq k \end{cases} \implies (T')^{\leq p} = \begin{cases} T^{\leq p} & p < k \\ T_{+j_p}^{\leq p} & p = k \\ T_{+j_p, -j_{p-1}}^{\leq p} & p > k \end{cases},$$

where we define $j_p := p + 1 - i_p$. This, in particular, means that the tableaux $S^{\leq p}$ and $T^{\leq p}$ are related using the following two expressions:

$$S_x^{\leq p} = m - (T')_{p+1-x}^{\leq p} = \begin{cases} m - T_{p+1-x}^{\leq p} & p < k \\ m - T_{p+1-x}^{\leq p} - \delta_{x(i_p)} & p \geq k \end{cases},$$

$$T_x^{\leq p} = m - (S')_{p+1-x}^{\leq p} = \begin{cases} m - S_{p+1-x}^{\leq p} & p < k \\ m - S_{p+1-x}^{\leq p} - \delta_{x(j_p)} & p \geq k \end{cases}.$$

We use the formula in Equation (60) to write the first factor in Equation (142):

$$\begin{aligned} (S^{=k}, e_{10}^{=k} | S_{+i_k}^{=k}) &= \left| \frac{\prod_{x=1}^{k-1} (c_x(S^{\leq k-1}) - c_{i_k}(S^{\leq k}) - 1)}{\prod_{x=1, x \neq i_k}^k (c_x(S^{\leq k}) - c_{i_k}(S^{\leq k}))} \right|^{1/2} \\ &= \left| \frac{\prod_{x=1}^{k-1} (S_x^{\leq k-1} - S_{i_k}^{\leq k} - x + i_k - 1)}{\prod_{x=1, x \neq i_k}^k (S_x^{\leq k} - S_{i_k}^{\leq k} - x + i_k)} \right|^{1/2} \\ &= \left| \frac{\prod_{x=1}^{k-1} (m - T_{k-x}^{\leq k-1} - m + T_{k+1-i_k}^{\leq k} + 1 - x + i_k - 1)}{\prod_{x=1, x \neq i_k}^k (m - T_{k+1-x}^{\leq k} - m + T_{k+1-i_k}^{\leq k} + 1 - x + i_k)} \right|^{1/2} \\ &= \left| \frac{\prod_{x=1}^{k-1} - (T_{k-x}^{\leq k-1} - T_{k+1-i_k}^{\leq k} - (k-x) + (k+1-i_k) - 1)}{\prod_{x=1, x \neq i_k}^k - (T_{k+1-x}^{\leq k} - T_{k+1-i_k}^{\leq k} - (k+1-x) + (k+1-i_k) - 1)} \right|^{1/2}. \end{aligned}$$

We will now change the iterator of the numerator to $x' \leftarrow k - x$, and the iterator of the denominator to $x'' \leftarrow k + 1 - x$. Moreover, recall that $j_k = k + 1 - i_k$. The expression thus becomes

$$\begin{aligned} (S^{=k}, e_{10}^{=k} | S_{+i_k}^{=k}) &= \left| \frac{\prod_{x'=1}^{k-1} (T_{x'}^{\leq k-1} - T_{j_k}^{\leq k} - x' + j_k - 1)}{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k - 1)} \right|^{1/2} \\ &= \left| \frac{\prod_{x'=1}^{k-1} (T_{x'}^{\leq k-1} - T_{j_k}^{\leq k} - x' + j_k - 1)}{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k - 1)} \right|^{1/2} \cdot \left| \frac{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k)}{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k - 1)} \right|^{1/2} \\ &= (T^{=k}, e_{10}^{=k} | T_{+j_k}^{=k}) \cdot \left| \frac{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k)}{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k - 1)} \right|^{1/2}. \end{aligned}$$

We can now undo the operations above to rewrite the multiplicative correction term in terms of the original tableau S :

$$\begin{aligned}
(S^{=k}, e_{10}^{=k} \mid S_{+i_k}^{=k}) &= (T^{=k}, e_{10}^{=k} \mid T_{+j_k}^{=k}) \cdot \left| \frac{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k)}{\prod_{x''=1, x'' \neq j_k}^k (T_{x''}^{\leq k} - T_{j_k}^{\leq k} - x'' + j_k - 1)} \right|^{1/2} \\
&= (T^{=k}, e_{10}^{=k} \mid T_{+j_k}^{=k}) \cdot \left| \frac{\prod_{x=1, x \neq i_k}^k (S_x^{\leq k} - S_{i_k}^{\leq k} - x + i_k - 1)}{\prod_{x=1, x \neq i_k}^k (S_x^{\leq k} - S_{i_k}^{\leq k} - x + i_k)} \right|^{1/2} \tag{143}
\end{aligned}$$

We proceed to deal with the remaining scalar factors, which correspond to $p = k + 1$ up to $p = d$. We use the formula in [Equation \(61\)](#) to write:

$$\begin{aligned}
&(S^{=p}, e_{11}^{=p} \mid S_{+i_p, -i_{p-1}}^{=p}) \\
&= S(i_p, i_{p-1}) \left| \frac{\prod_{x \neq i_{p-1}}^{p-1} (c_x(S^{\leq p-1}) - c_{i_p}(S^{\leq p}) - 1) \prod_{x \neq i_p}^p (c_x(S^{\leq p}) - c_{i_{p-1}}(S^{\leq p-1}))}{\prod_{x \neq i_p}^p (c_x(S^{\leq p}) - c_{i_p}(S^{\leq p})) \prod_{x \neq i_{p-1}}^{p-1} (c_x(S^{\leq p-1}) - c_{i_{p-1}}(S^{\leq p-1}) - 1)} \right|^{1/2} \\
&= S(i_p, i_{p-1}) \left| \frac{\prod_{x \neq i_{p-1}}^{p-1} (S_x^{\leq p-1} - S_{i_p}^{\leq p} - x + i_p - 1) \prod_{x \neq i_p}^p (S_x^{\leq p} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1})}{\prod_{x \neq i_p}^p (S_x^{\leq p} - S_{i_p}^{\leq p} - x + i_p) \prod_{x \neq i_{p-1}}^{p-1} (S_x^{\leq p-1} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1} - 1)} \right|^{1/2} \\
&= S(i_p, i_{p-1}) \left| \frac{\prod_{x \neq i_{p-1}}^{p-1} (m - T_{p-x}^{\leq p-1} - m + T_{p+1-i_p}^{\leq p} + 1 - x + i_p - 1)}{\prod_{x \neq i_p}^p (m - T_{p+1-x}^{\leq p} - m + T_{p+1-i_p}^{\leq p} + 1 - x + i_p)} \right|^{1/2} \\
&\quad \cdot \left| \frac{\prod_{x \neq i_p}^p (m - T_{p+1-x}^{\leq p} - m + T_{p-i_{p-1}}^{\leq p-1} + 1 - x + i_{p-1})}{\prod_{x \neq i_{p-1}}^{p-1} (m - T_{p-x}^{\leq p-1} - m + T_{p-i_{p-1}}^{\leq p-1} + 1 - x + i_{p-1} - 1)} \right|^{1/2} \\
&= S(i_p, i_{p-1}) \left| \frac{\prod_{x \neq i_{p-1}}^{p-1} - (T_{p-x}^{\leq p-1} - T_{p+1-i_p}^{\leq p} - (p-x) + (p+1-i_p) - 1)}{\prod_{x \neq i_p}^p - (T_{p+1-x}^{\leq p} - T_{p+1-i_p}^{\leq p} - (p+1-x) + (p+1-i_p) - 1)} \right|^{1/2} \\
&\quad \cdot \left| \frac{\prod_{x \neq i_p}^p - (T_{p+1-x}^{\leq p} - T_{p-i_{p-1}}^{\leq p-1} - (p+1-x) + (p-i_{p-1}))}{\prod_{x \neq i_{p-1}}^{p-1} - (T_{p-x}^{\leq p-1} - T_{p-i_{p-1}}^{\leq p-1} - (p-x) + (p-i_{p-1}))} \right|^{1/2}.
\end{aligned}$$

We will now change the iterators of the numerators to $x' \leftarrow p - x$, $y' \leftarrow p + 1 - x$, and the iterators of the denominators to $x'' \leftarrow p + 1 - x$, $y'' \leftarrow p - x$. Moreover, recall that $j_p = p + 1 - i_p$ and $j_{p-1} = p - i_{p-1}$. The expression thus becomes

$$\begin{aligned}
&(S^{=p}, e_{11}^{=p} \mid S_{+i_p, -i_{p-1}}^{=p}) \\
&= S(i_p, i_{p-1}) \left| \frac{\prod_{x' \neq j_{p-1}}^{p-1} (T_{x'}^{\leq p-1} - T_{j_p}^{\leq p} - x' + j_p - 1)}{\prod_{y' \neq j_p}^p (T_{y'}^{\leq p} - T_{j_p}^{\leq p} - y' + j_p - 1)} \cdot \frac{\prod_{x'' \neq j_p}^p (T_{x''}^{\leq p} - T_{j_{p-1}}^{\leq p-1} - x'' + j_{p-1})}{\prod_{y'' \neq j_{p-1}}^{p-1} (T_{y''}^{\leq p-1} - T_{j_{p-1}}^{\leq p-1} - y'' + j_{p-1})} \right|^{1/2} \\
&= \frac{S(i_p, i_{p-1})}{S(j_p, j_{p-1})} (T^{=p}, e_{11}^{=p} \mid T_{+j_p, -j_{p-1}}^{=p}) \\
&\quad \cdot \left| \prod_{y' \neq j_p}^p \frac{T_{y'}^{\leq p} - T_{j_p}^{\leq p} - y' + j_p}{T_{y'}^{\leq p} - T_{j_p}^{\leq p} - y' + j_p - 1} \cdot \prod_{y'' \neq j_{p-1}}^{p-1} \frac{T_{y''}^{\leq p-1} - T_{j_{p-1}}^{\leq p-1} - y'' + j_{p-1} - 1}{T_{y''}^{\leq p-1} - T_{j_{p-1}}^{\leq p-1} - y'' + j_{p-1}} \right|^{1/2}.
\end{aligned}$$

We note here that

$$\frac{S(i_p, i_{p-1})}{S(j_p, j_{p-1})} = \frac{S(i_p, i_{p-1})}{S(p+1-i_p, p-i_{p-1})} = -1,$$

since whenever $i_p > i_{p-1}$, then $p+1-i_p \leq p-i_{p-1}$, and whenever $i_p \leq i_{p-1}$, then $p+1-i_p > p-i_{p-1}$. We will now “undo” the operations above to rewrite the multiplicative correction term in terms of the original tableau S :

$$\begin{aligned}
& \left(S^=p, e_{11}^=p \mid S_{+i_p, -i_{p-1}}^=p \right) \\
&= - \left(T^=p, e_{11}^=p \mid T_{+j_p, -j_{p-1}}^=p \right) \cdot \left| \prod_{y' \neq j_p}^p \frac{m - S_{p+1-y'}^{\leq p} - m + S_{p+1-j_p}^{\leq p} + 1 - y' + j_p}{m - S_{p+1-y'}^{\leq p} - m + S_{p+1-j_p}^{\leq p} + 1 - y' + j_p - 1} \right|^{1/2} \\
& \quad \cdot \left| \prod_{y'' \neq j_{p-1}}^{p-1} \frac{m - S_{p-y''}^{\leq p-1} - m + S_{p-j_{p-1}}^{\leq p-1} + 1 - y'' + j_{p-1} - 1}{m - S_{p-y''}^{\leq p-1} - m + S_{p-j_{p-1}}^{\leq p-1} + 1 - y'' + j_{p-1} - 1} \right|^{1/2} \\
&= - \left(T^=p, e_{11}^=p \mid T_{+j_p, -j_{p-1}}^=p \right) \cdot \left| \prod_{x \neq i_p}^p \frac{S_x^{\leq p} - S_{i_p}^{\leq p} - x + i_p - 1}{S_x^{\leq p} - S_{i_p}^{\leq p} - x + i_p} \right|^{1/2} \\
& \quad \cdot \left| \prod_{x \neq i_{p-1}}^{p-1} \frac{S_x^{\leq p-1} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1} - 1}{S_x^{\leq p-1} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1} - 1} \right|^{1/2}. \tag{144}
\end{aligned}$$

We can thus write the original Clebsch-Gordan coefficient as

$$\begin{aligned}
& \langle S, k \mid S' \rangle \\
&= (S^=k, e_{10}^=k \mid S_{+i_k}^=k) \cdot \prod_{p=k+1}^d \left(S^=p, e_{11}^=p \mid S_{+i_p, -i_{p-1}}^=p \right) \\
&= (T^=k, e_{10}^=k \mid T_{+j_k}^=k) \cdot \left| \frac{\prod_{x=1, x \neq i_k}^k (S_x^{\leq k} - S_{i_k}^{\leq k} - x + i_k - 1)}{\prod_{x=1, x \neq i_k}^k (S_x^{\leq k} - S_{i_k}^{\leq k} - x + i_k)} \right|^{1/2} \cdot \prod_{p=k+1}^d - \left(T^=p, e_{11}^=p \mid T_{+j_p, -j_{p-1}}^=p \right) \\
& \quad \cdot \left| \prod_{x \neq i_p}^p \frac{S_x^{\leq p} - S_{i_p}^{\leq p} - x + i_p - 1}{S_x^{\leq p} - S_{i_p}^{\leq p} - x + i_p} \cdot \prod_{x \neq i_{p-1}}^{p-1} \frac{S_x^{\leq p-1} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1} - 1}{S_x^{\leq p-1} - S_{i_{p-1}}^{\leq p-1} - x + i_{p-1} - 1} \right|^{1/2}. \quad (\text{Eqs. (143) and (144)})
\end{aligned}$$

If we telescope the above expression, all the products cancel except the one with the iterator $x \neq i_d$. Moreover, we have defined $i_d = i$, and $S^{\leq d} = \text{sh}(S) = \lambda$. Thus,

$$\begin{aligned}
& \langle S, k \mid S' \rangle \\
&= (-1)^{d-k} (T^=k, e_{10}^=k \mid T_{+j_k}^=k) \cdot \prod_{p=k+1}^d \left(T^=p, e_{11}^=p \mid T_{+j_p, -j_{p-1}}^=p \right) \cdot \left| \prod_{x \neq i_d}^d \frac{S_x^{\leq d} - S_{i_d}^{\leq d} - x + i_d - 1}{S_x^{\leq d} - S_{i_d}^{\leq d} - x + i_d} \right|^{1/2} \\
&= (-1)^{d-k} \langle T, T^k \mid T' \rangle \cdot \left| \prod_{x \neq i}^d \frac{\lambda_x - \lambda_i - x + i - 1}{\lambda_x - \lambda_i - x + i} \right|^{1/2} \\
&= (-1)^{d-k} \langle T, T^k \mid T' \rangle \cdot \left(\frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_\lambda^d)} \right)^{1/2}.
\end{aligned}$$

The last equality follows from the Weyl dimension formula in [Equation \(15\)](#). First, we write

$$\begin{aligned}
\dim(V_\lambda^d) &= \prod_{1 \leq x < y \leq d} \frac{(\lambda_x - \lambda_y) + (y - x)}{y - x}, \\
\dim(V_{\lambda+e_i}^d) &= \prod_{\substack{1 \leq x < y \leq d \\ x, y \neq i}} \frac{(\lambda_x - \lambda_y) + (y - x)}{y - x} \cdot \prod_{x < i} \frac{(\lambda_x - \lambda_i) + (i - x) - 1}{i - x} \cdot \prod_{i < x} \frac{(\lambda_i - \lambda_x) + (x - i) + 1}{x - i}.
\end{aligned}$$

Thus

$$\begin{aligned}
\frac{\dim(V_{\lambda+e_i}^d)}{\dim(V_\lambda^d)} &= \prod_{x < i} \frac{(\lambda_x - \lambda_i) + (i - x) - 1}{(\lambda_x - \lambda_i) + (i - x)} \cdot \prod_{i < x} \frac{(\lambda_i - \lambda_x) + (x - i) + 1}{(\lambda_i - \lambda_x) + (x - i)} \\
&= \prod_{x < i} \frac{\lambda_x - \lambda_i - x + i - 1}{\lambda_x - \lambda_i - x + i} \cdot \prod_{i < x} \left| \frac{\lambda_x - \lambda_i - x + i - 1}{\lambda_x - \lambda_i - x + i} \right| \\
&= \left| \prod_{x \neq i} \frac{\lambda_x - \lambda_i - x + i - 1}{\lambda_x - \lambda_i - x + i} \right|,
\end{aligned}$$

Where the second equality holds because the left-hand side is always nonnegative. □