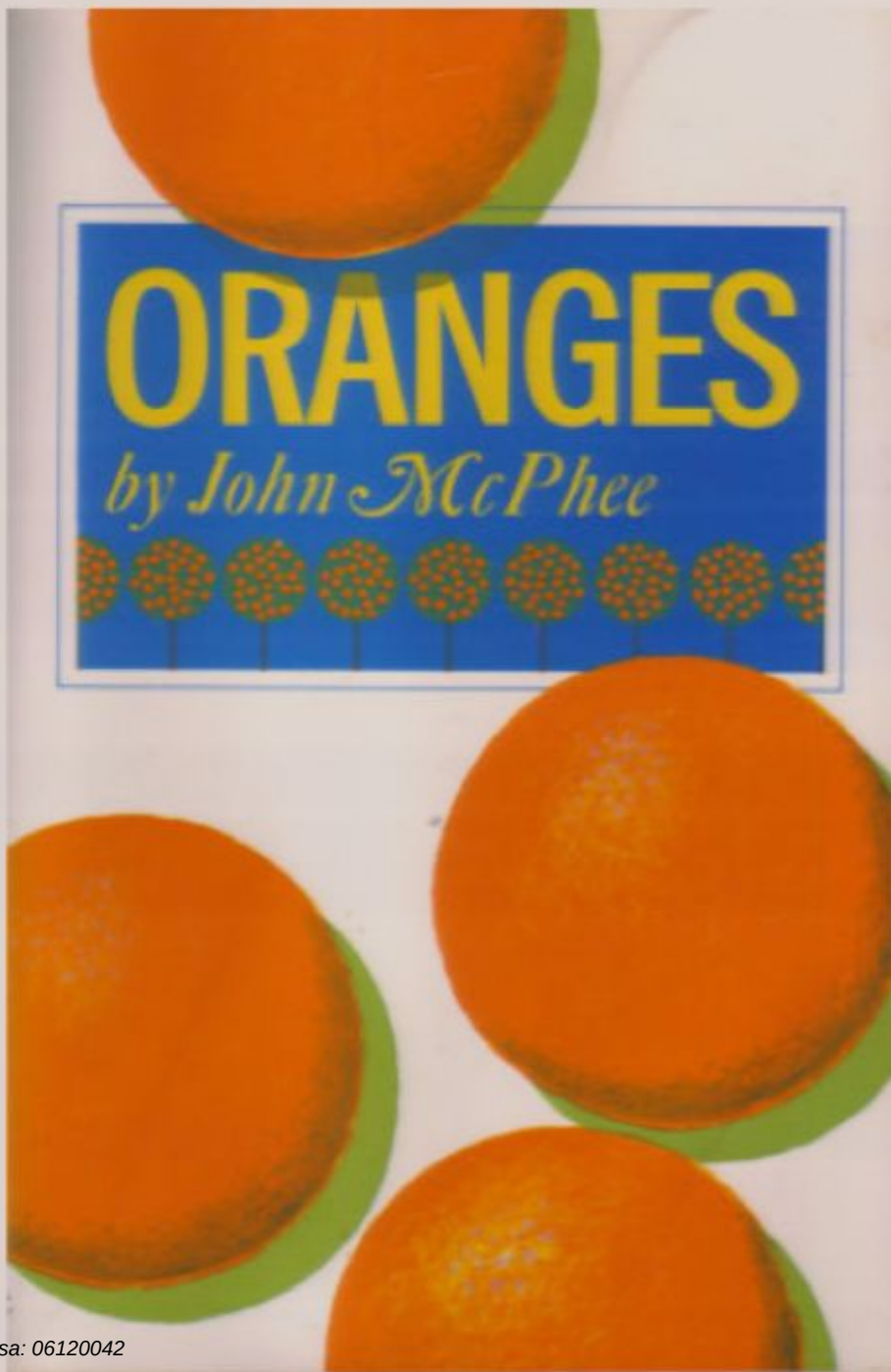


Title: Reliable Quantum State Estimation from Quantum Scoring Rules

Date: Dec 08, 2006 11:05 AM

URL: <http://pirsa.org/06120042>

Abstract: Inferring a quantum system's state, from repeated measurements, is critical for verifying theories and designing quantum hardware. It's also surprisingly easy to do wrong, as illustrated by maximum likelihood estimation (MLE), the current state of the art. I'll explain why MLE yields unreliable and rank-deficient estimates, why you shouldn't be a quantum frequentist, and why we need a different approach. I'll show how operational divergences -- well-motivated metrics designed to evaluate estimates -- follow from quantum strictly proper scoring rules. This motivates Bayesian Mean Estimation (BME), and I'll show how it fixes most of the problems with MLE. I'll conclude with a couple of speculations about the future of quantum state and process estimation.



Oranges (Paperback)

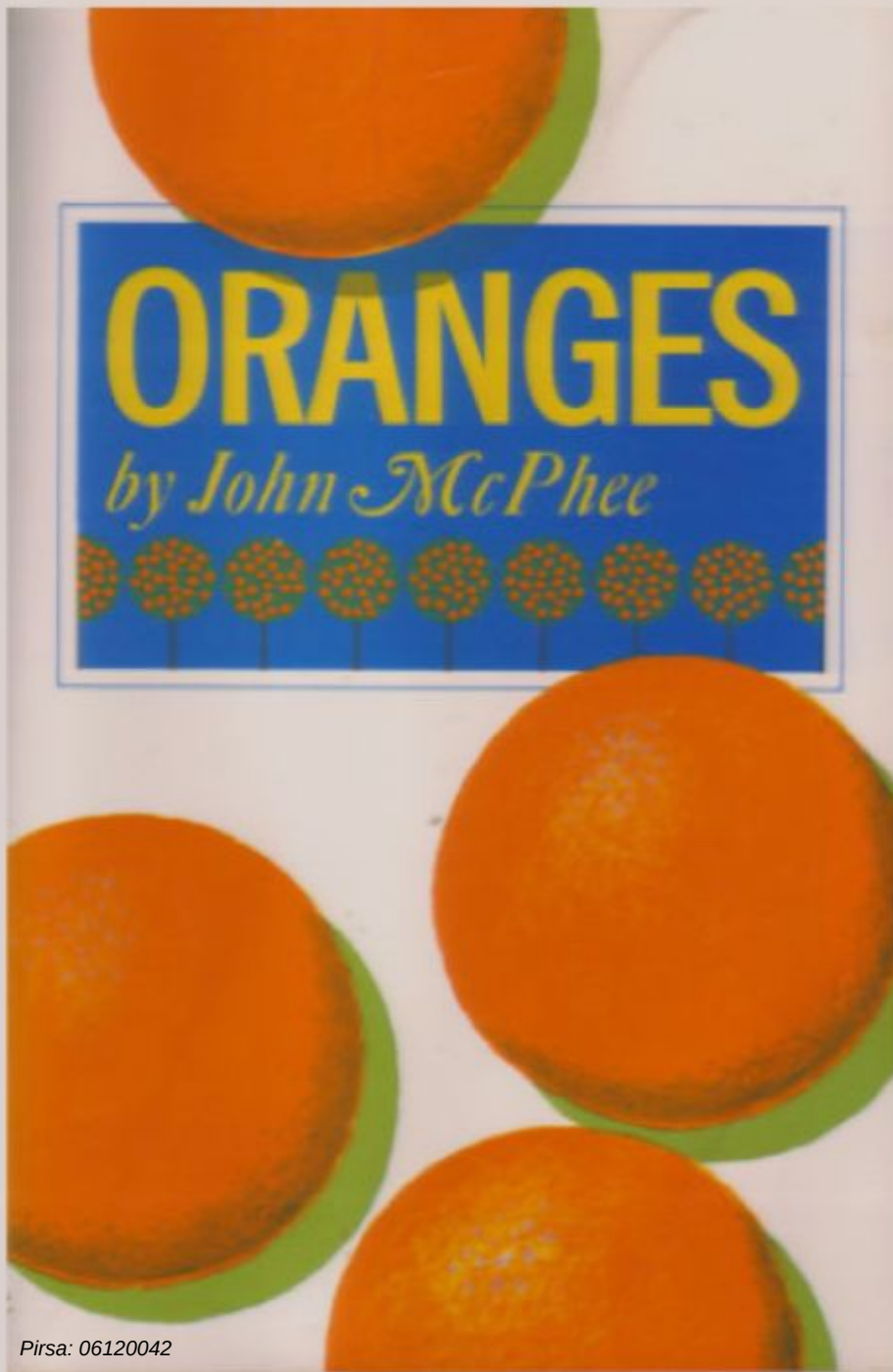
by [John McPhee](#)

★★★★★ [\(18 customer reviews\)](#)

List Price: ~~\$13.00~~

Price: **\$10.40** & eligible for **FREE**
Super Saver Shipping on
orders over \$25. [Details](#)

You Save: **\$2.60 (20%)**



Pirsa: 06120042



Oranges (Paperback)

by [John McPhee](#)

★★★★★ ([18 customer reviews](#))

List Price: ~~\$13.00~~

Price: **\$10.40** & eligible for **FREE**
Super Saver Shipping on
orders over \$25. [Details](#)

You Save: **\$2.60 (20%)**



Reliable Quantum State Estimation

quant-ph/0611080

from

Quantum Scoring Rules

quant-ph/0603116

Robin Blume-Kohout (Caltech-IQI)

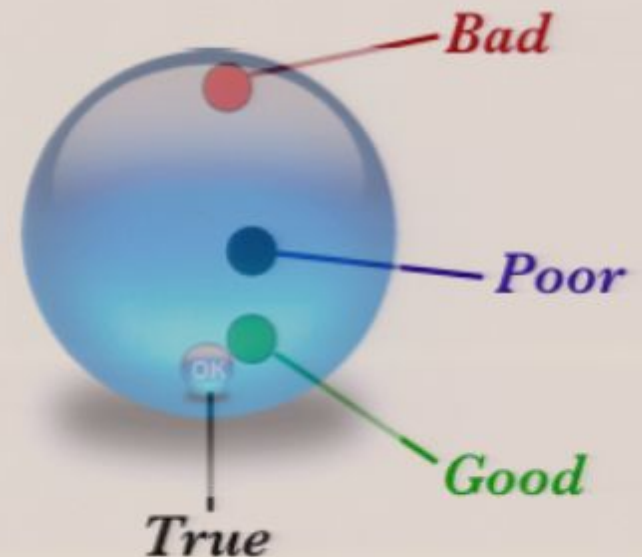
with Patrick Hayden (McGill)
and Karan Malhotra (IIT-Kanpur)

What is State Estimation?

- **Goal:** Characterize a *source* of quantum systems.
- You measure N identical copies.
- Measurement record:

$$\mathcal{M} = \{ \hat{E}_1, \hat{E}_2, \dots, \hat{E}_N \}$$

- Then report your best guess for ρ .
- **Key Application:** Verifying quantum hardware for fault tolerant quantum computation.



A few points up front:

A few points up front:

- Estimation procedures consist of:
 - a) a rule for making measurements,
 - b) a rule for analyzing the data.

A few points up front:

- Estimation procedures consist of:
 - a) a rule for making measurements,
 - b) a rule for analyzing the data.
- I'm only talking about the *data analysis* part.

A few points up front:

- Estimation procedures consist of:
 - a) a rule for making measurements,
 - b) a rule for analyzing the data.
- I'm only talking about the *data analysis* part.
- Every [sane] procedure works roughly the same as $N \Rightarrow \infty$. I'm interested in $1 < N < \infty$.

What I'm a-gonna say

- 1. Motivation:** problems with MLE*.
- 2. Foundation:** quantum scoring rules.
- 3. Solution:** the BME algorithm (& its features)
- 4. Testing:** dueling procedures! (numerics)

What I'm a-gonna say

1. Motivation: problems with MLE*.

2. Foundation: quantum scoring rules.

3. Solution: the BME algorithm (& its features)

4. Testing: dueling procedures! (numerics)

What's the Problem?

- Current technology: **Maximum Likelihood Estimation**

$$\mathcal{M} = \{ \hat{E}_1, \hat{E}_2, \dots, \hat{E}_N \} \longrightarrow \mathcal{L}(\rho) \equiv p(\mathcal{M}|\rho) \longrightarrow \rho_{\text{MLE}}$$

- ρ_{MLE} does not honestly represent the experimentalist's knowledge about the observed ensemble.
- **Why?** ρ_{MLE} typically has zero eigenvalues:

$$\lambda_i = \langle \phi_i | \rho_{\text{MLE}} | \phi_i \rangle = 0$$

- Zero eigenvalue = zero probability

What's the Problem?

- Current technology: **Maximum Likelihood Estimation**

$$\mathcal{M} = \{\hat{E}_1, \hat{E}_2, \dots, \hat{E}_N\} \longrightarrow \mathcal{L}(\rho) \equiv p(\mathcal{M}|\rho) \longrightarrow \rho_{\text{MLE}}$$

- ρ_{MLE} does not honestly represent the experimentalist's knowledge about the observed ensemble.
- **Why?** ρ_{MLE} typically has zero eigenvalues:

$$\lambda_i = \langle \phi_i | \rho_{\text{MLE}} | \phi_i \rangle = 0$$

- Zero eigenvalue = zero probability

↳ absolute certainty

↳ doesn't admit error bars!

Yes, it's a problem

PHYSICAL REVIEW A, VOLUME 64, 052312

Measurement of qubits

Daniel F. V. James,^{1,*} Paul G. Kwiat,^{2,3} William J. Munro,^{4,5} and Andrew G. White^{2,4}

¹Theoretical Division T-4, Los Alamos National Laboratory, Los Alamos, New Mexico 87545

²Physics Division P-21, Los Alamos National Laboratory, Los Alamos, New Mexico 87545

³Department of Physics, University of Illinois, Urbana-Champaign, Illinois 61801

⁴Department of Physics, University of Queensland, Brisbane, Queensland 4072, Australia

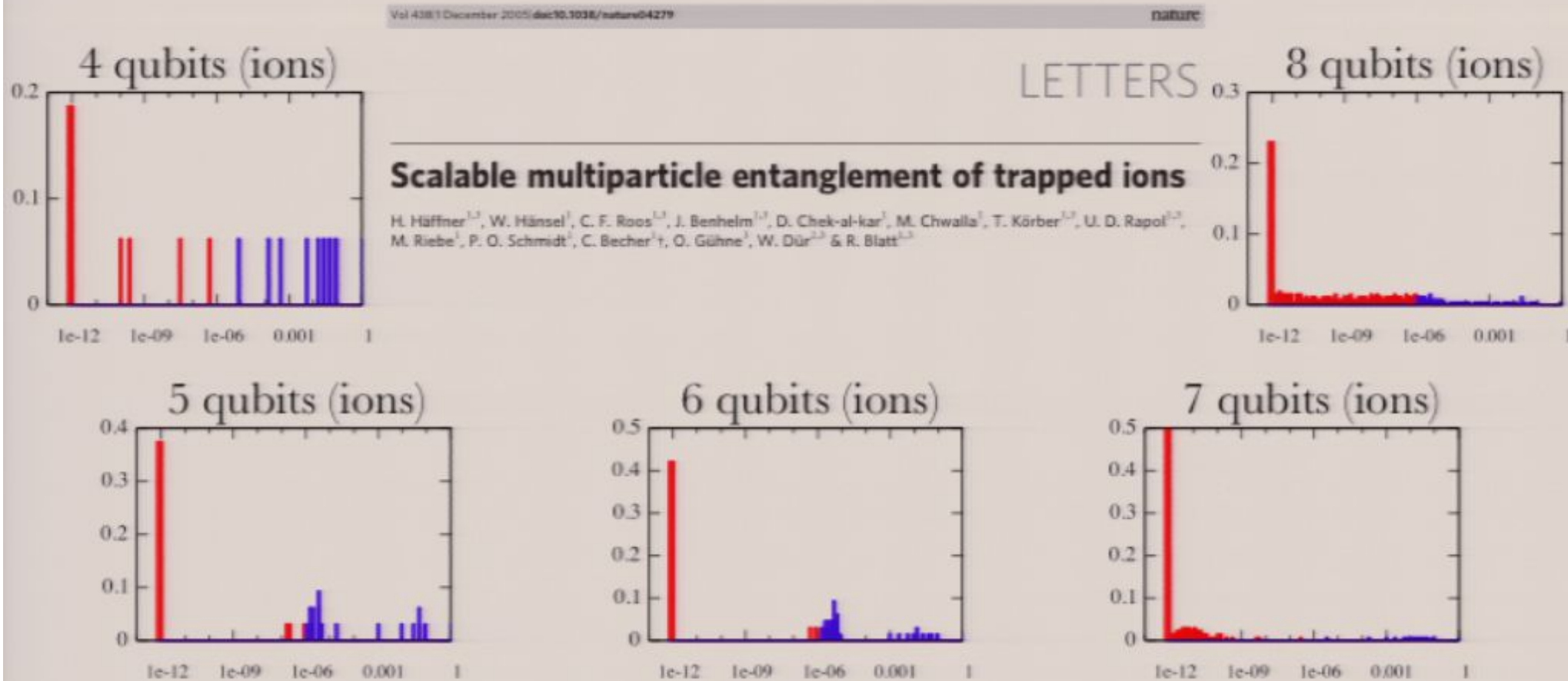
⁵Hewlett-Packard Laboratories, Filton Road, Stoke Gifford, Bristol BS34 8QZ, United Kingdom

(Received 20 March 2001; published 16 October 2001)

$$\hat{\rho} = \begin{pmatrix} 0.5069 & -0.0239 + i0.0106 & -0.0412 - i0.0221 & 0.4833 + i0.0329 \\ -0.0239 - i0.0106 & 0.0048 & 0.0023 + i0.0019 & -0.0296 - i0.0077 \\ -0.0412 + i0.0221 & 0.0023 - i0.0019 & 0.0045 & -0.0425 + i0.0192 \\ 0.4833 - i0.0329 & -0.0296 + i0.0077 & -0.0425 - i0.0192 & 0.4839 \end{pmatrix}$$

This matrix is illustrated in Fig. 3 (right). In this case, the matrix has eigenvalues 0.986 022, 0.013 977 7, 0, and 0; and $\text{Tr}\{\hat{\rho}^2\} = 0.972 435$, indicating that, while the linear reconstruction gave a nonphysical density matrix, the maximum likelihood reconstruction gives a legitimate density matrix.

Yes, it's a problem



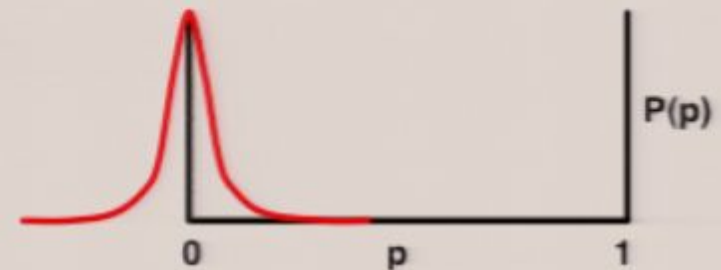
1. D. F. V. James et. al., "Measurement of Qubits," *Phys. Rev. A*, 64:052312 (2001)

Pirsa: 06120042

2. Häffner et. al., "Scalable multiparticle entanglement of trapped ions," *Nature*, 438:643-6 (2005)

Error Bars?

- “But aren’t there error bars in there?”
- Not always.
- What does $p = 0 \pm 0.1$ mean?



Error Bars?

- “But aren’t there error bars in there?”
- Not always.
- What does ~~$p = 0 \pm 0.1$~~ mean?
 $\Rightarrow p = 0.05 \pm$



Error Bars?

- “But aren’t there error bars in there?”
- Not always.
- What does ~~$p = 0 \pm 0.1$~~ mean?
 $\Rightarrow p = 0.05 \pm 0.05$



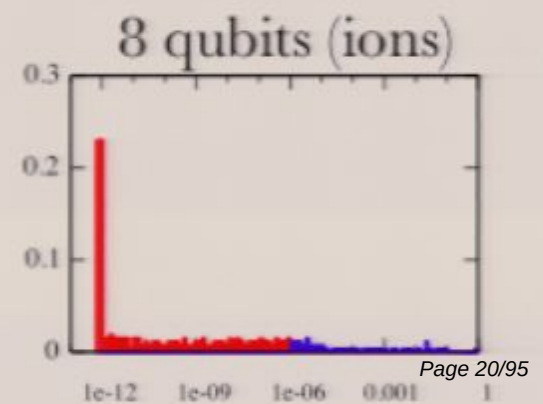
Error Bars?

- “But aren’t there error bars in there?”
- Not always.
- What does ~~$p = 0 \pm 0.1$~~ mean?
 $\Rightarrow p = 0.05 \pm 0.05$
- Increasing *small* probabilities
 $\Rightarrow$ decreasing *large* ones



Error Bars?

- “But aren’t there error bars in there?”
- Not always.
- What does ~~$p = 0 \pm 0.1$~~ mean?
 $\Rightarrow p = 0.05 \pm 0.05$
- Increasing *small* probabilities
 $\Rightarrow$ decreasing *large* ones
- Effect of “fixing”
~218 zero eigenvalues?



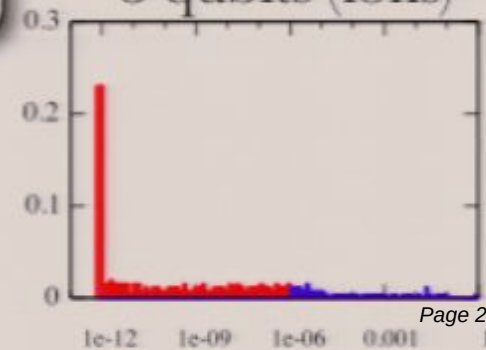
Error Bars?

$$\rho = \begin{pmatrix} .75 & & & \\ & .06 & & \\ & & \ddots & \\ & & & 0 \\ & & & & \ddots \\ & & & & & 0 \end{pmatrix}$$

ere?"



8 qubits (ions)



- Effect of “fixing”
~218 zero eigenvalues?

Error Bars?

$$\rho = \begin{pmatrix} .75 & & & \\ & .06 & & \\ & & \ddots & \\ & & & \epsilon & & \\ & & & & \ddots & \\ & & & & & \epsilon \end{pmatrix}$$

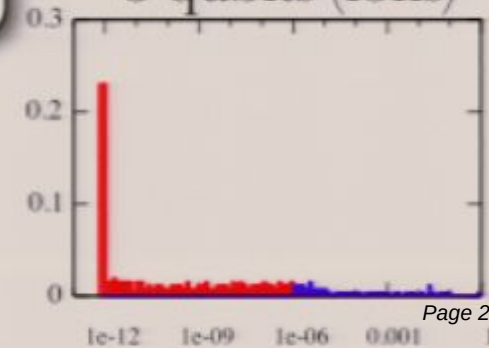
7×10^5 measurements
 65,536 indep. components
 6,561 different meas.
 256 eigenvalues

$\epsilon = \epsilon(N_{\text{measurements}}) \dots$ what is N ?

ere?"



8 qubits (ions)



- Effect of “fixing”
~218 zero eigenvalues?

Error Bars?

$$\rho = \begin{pmatrix} .75 & & & \\ & .06 & & \\ & & \ddots & \\ & & & 10^{-6} & & \\ & & & & \ddots & \\ & & & & & 10^{-6} \end{pmatrix}$$

7×10^5 measurements
 65,536 indep. components
 6,561 different meas.
 256 eigenvalues

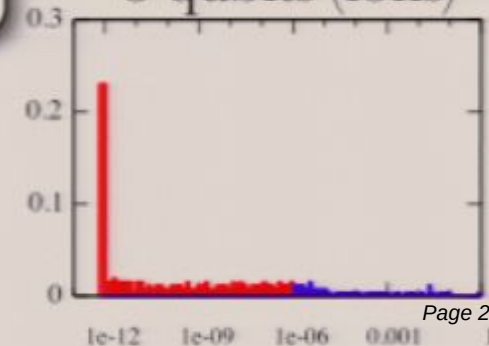
ere?"

$\varepsilon = \varepsilon(N_{\text{measurements}}) \dots$ what is N ?

$N = 7 \times 10^5?$



8 qubits (ions)



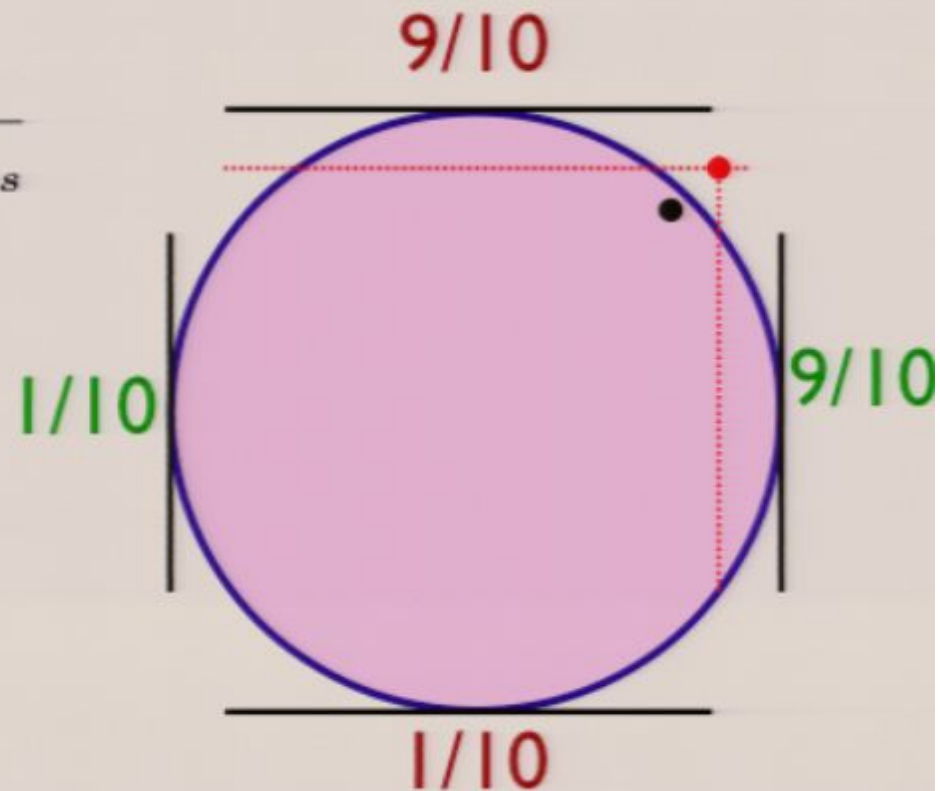
- Effect of “fixing”
~218 zero eigenvalues?

A man wearing a blue jacket, a wide-brimmed hat, and shorts is standing on a rocky mountain ridge. He is pointing his right hand towards a vast mountain range in the background. The mountains are rugged and brownish-grey, with some green patches. The sky is blue with scattered white clouds. The foreground is a rocky, uneven surface.

So... how did
we get here?

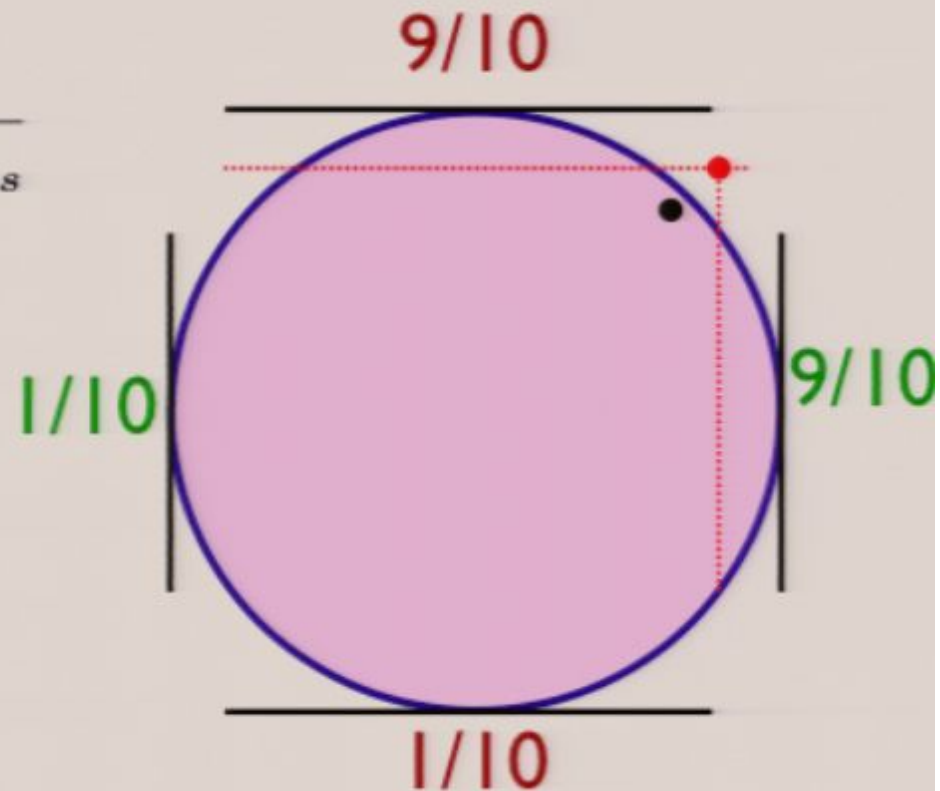
Tomography

- Assume $p(\hat{M}_i) = \frac{n_i}{N_{meas}}$
- Probabilities give expectation values.
- Expectation values uniquely identify the density matrix.
- Negativity: $\rho_{\text{tomo}} \not\geq 0$
 Why? (a) equation of frequencies with probabilities,
 (b) observables aren't truly independent.



Tomography

- Assume $p(\hat{M}_i) = \frac{n_i}{N_{meas}}$
- Probabilities give expectation values.
- Expectation values uniquely identify the density matrix.
- Negativity: $\rho_{\text{tomo}} \not\geq 0$
 Why? (a) equation of frequencies with probabilities,
 (b) observables aren't truly independent.



Maximum Likelihood (MLE)

- Report the state that *makes the data most likely* -- i.e., maximize the likelihood of data.

$$\mathcal{L}(\rho) = p(\mathcal{M}|\rho) = \text{Tr}(\hat{M}_1\rho) \text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)$$

- Involves maximization over valid states (parameterized by their square roots)
- Z. Hradil, “Quantum-state Estimation,” *Phys. Rev. A*, 55:R1561 (1997)

How MLE Works

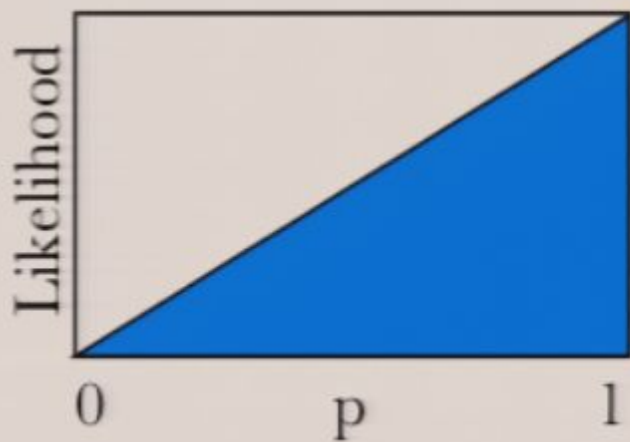
How MLE Works

Example I: Bias of a Coin

How MLE Works

Example I: Bias of a Coin

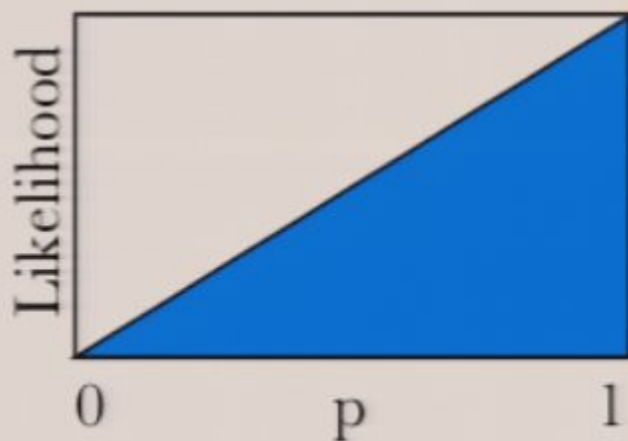
(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



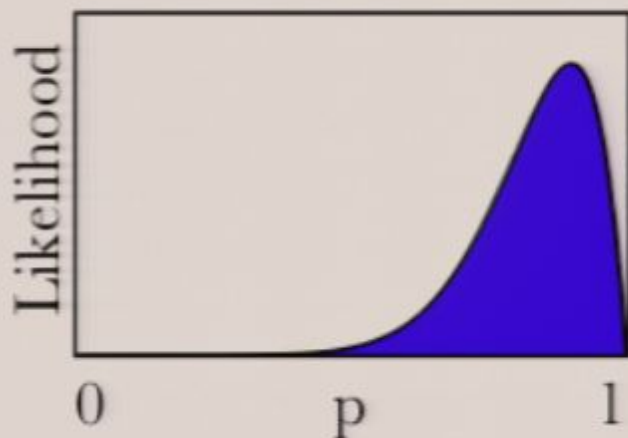
How MLE Works

Example I: Bias of a Coin

(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



(b)
 $9 \times \text{heads}$
 $1 \times \text{tails}$

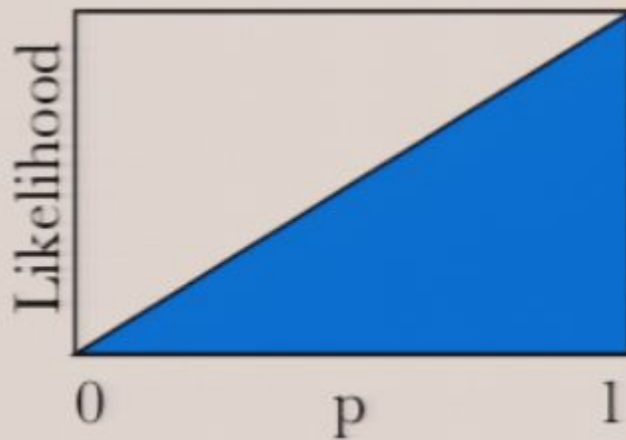


How MLE Works

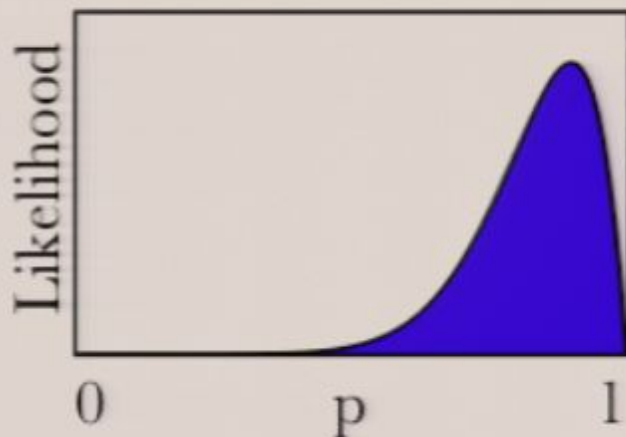
Example 1: Bias of a Coin

Example 2: Qubit

(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



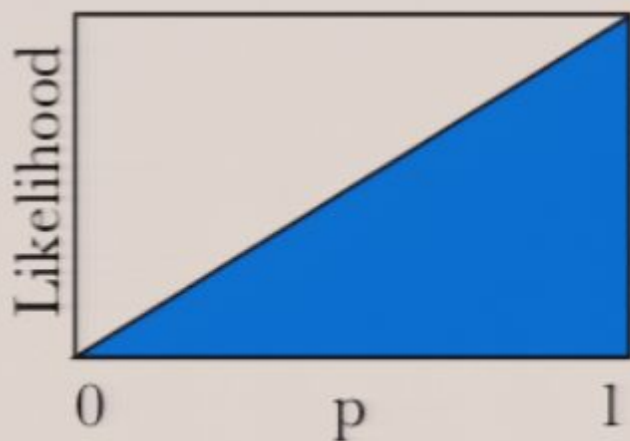
(b)
 $9 \times \text{heads}$
 $1 \times \text{tails}$



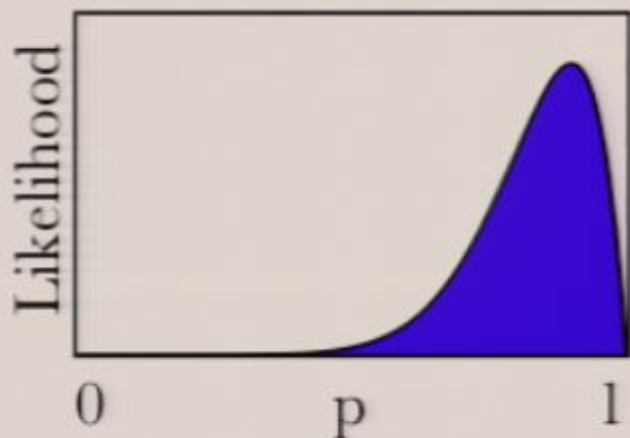
How MLE Works

Example 1: Bias of a Coin

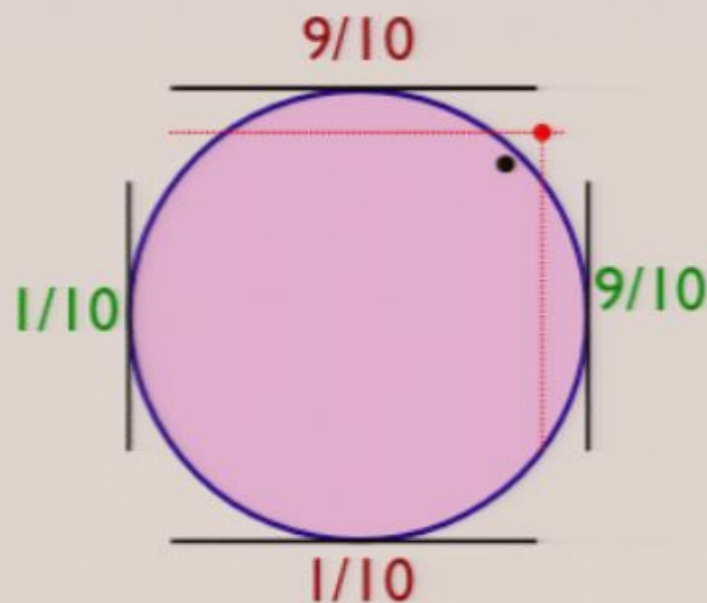
(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



(b)
 $9 \times \text{heads}$
 $1 \times \text{tails}$



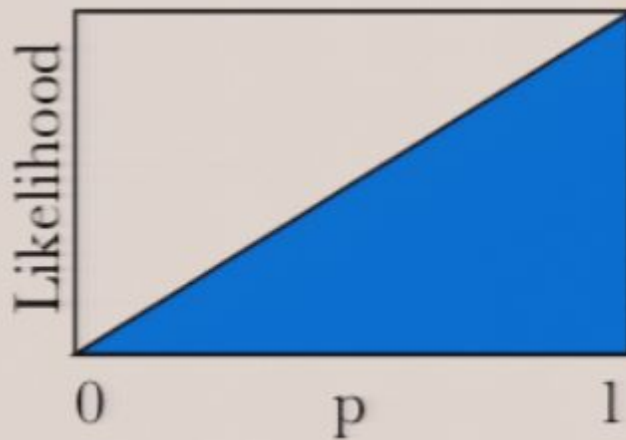
Example 2: Qubit



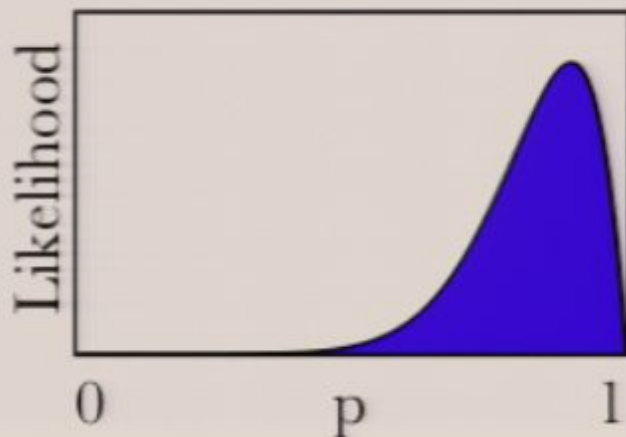
How MLE Works

Example 1: Bias of a Coin

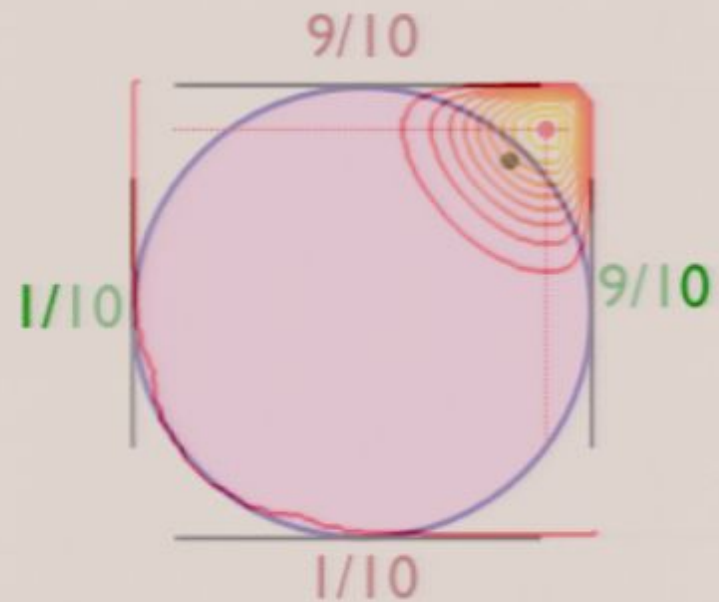
(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



(b)
 $9 \times \text{heads}$
 $1 \times \text{tails}$



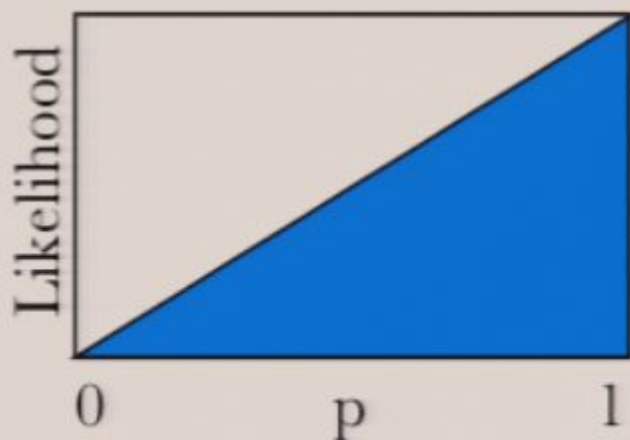
Example 2: Qubit



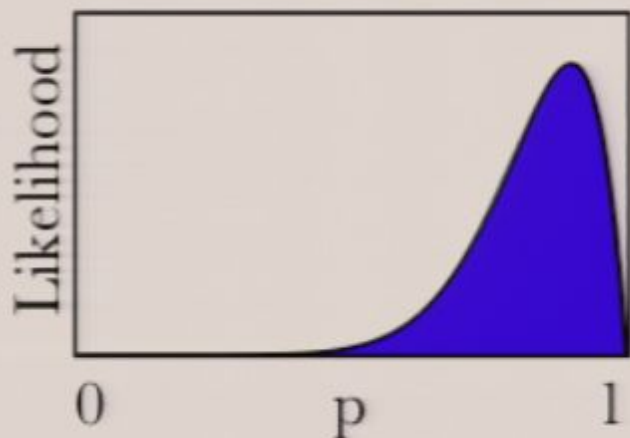
How MLE Works

Example 1: Bias of a Coin

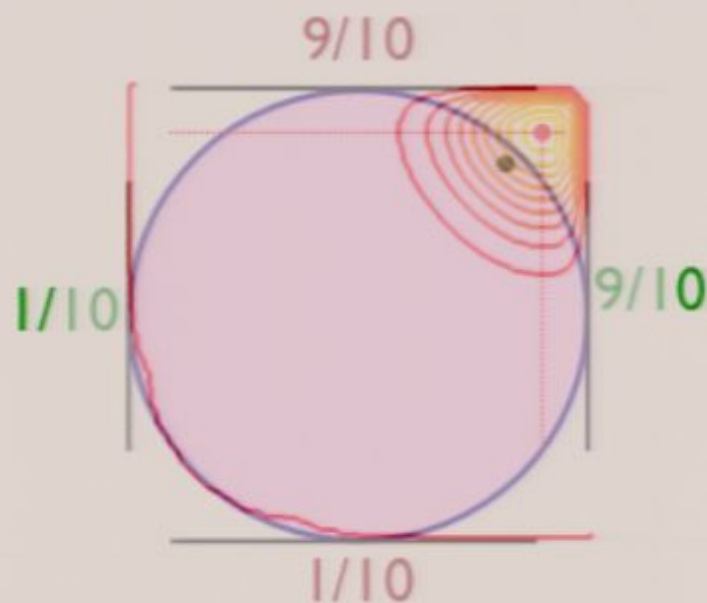
(a)
 $1 \times \text{head}$
 $0 \times \text{tails}$



(b)
 $9 \times \text{heads}$
 $1 \times \text{tails}$



Example 2: Qubit



MLE replaces *negative* eigenvalues with zeros.

A minimal fix for tomography.

What's going wrong?

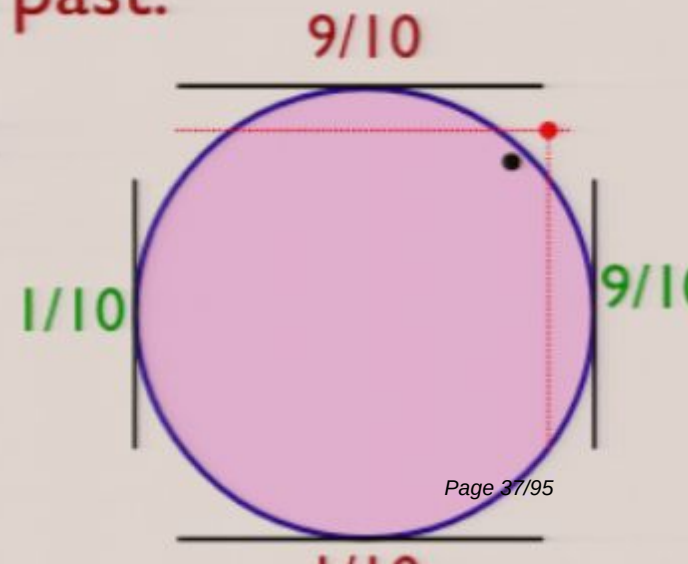
(the really, *really* short version!)

- MLE is a great way of *fitting the data*.
- Best estimate: *predicts* what we've seen.
- MLE takes observed *frequencies* literally --
“The future will look like the past.”
- But... in QM the future might contain measurements that we haven't made yet!

What's going wrong?

(the really, really short version!)

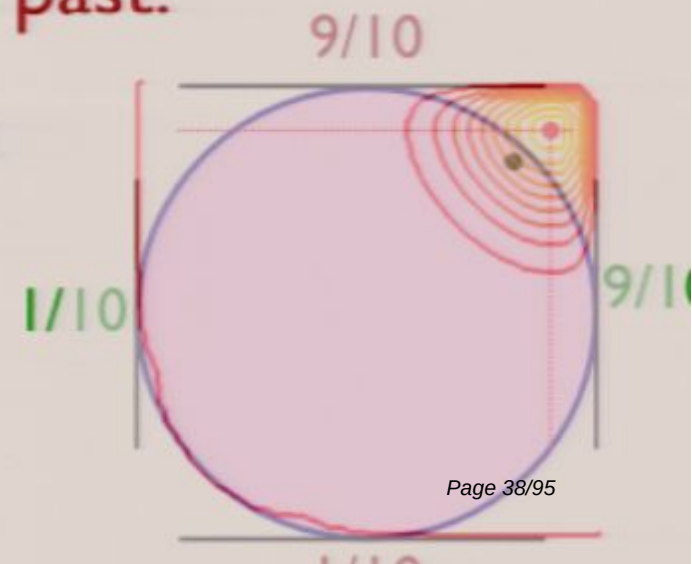
- MLE is a great way of *fitting the data*.
- Best estimate: *predicts* what we've seen.
- MLE takes observed *frequencies* literally --
“The future will look like the past.”
- But... in QM the future might contain measurements that we haven't made yet!



What's going wrong?

(the really, *really* short version!)

- MLE is a great way of *fitting the data*.
- Best estimate: *predicts* what we've seen.
- MLE takes observed *frequencies* literally --
“The future will look like the past.”
- But... in QM the future might contain measurements that we haven't made yet!



What's going wrong?

(the really, really short version!)

- MLE is a great way of *fitting the data*.
- Best estimate: *predicts* what we've seen.
- MLE takes observed *frequencies* literally --
“The future will look like the past.”
- But... in QM the future might contain measurements that we haven't made yet!



...so MLE assigns $p=0$ to events that were never observed not to happen...

Hmm... what to do?

- 1. Motivation:** problems with MLE.
- 2. Foundation:** quantum scoring rules.
- 3. Solution:** the BME algorithm (& its features)
- 4. Testing:** dueling procedures! (numerics)

Hmm... what to do?

1. Motivation: problems with MLE.

2. Foundation: quantum scoring rules.

3. Solution: the BME algorithm (& its features)

4. Testing: dueling procedures! (numerics)

States as Predictions

- Quantum states are like probability distributions: they predict the outcome of future measurements.
- Lets define a metric $f(\rho : \sigma)$ that measures how well σ predicts measurements on ρ .
 1. **The best estimate of ρ is ρ itself.**
 - i.e., $f(\rho : \rho) > f(\rho : \sigma)$ for all $\sigma \neq \rho$.
 2. **$f(\rho : \sigma)$ should correspond to an *operational* test**
 - i.e., someone's utility (reward or cost) for some practical procedure.

States as Predictions

- Quantum states are like probability distributions: they predict the outcome of future measurements.
- Lets define a metric $f(\rho : \sigma)$ that measures how well σ predicts measurements on ρ .
 1. **The best estimate of ρ is ρ itself.**
 - i.e., $f(\rho : \rho) > f(\rho : \sigma)$ for all $\sigma \neq \rho$.
 2. **$f(\rho : \sigma)$ should correspond to an *operational* test**
 - i.e., someone's utility (reward or cost) for some practical procedure.

$\implies$ **Quantum Strictly Proper Scoring Rules**

Quantum Strictly Proper Scoring Rules

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Victor's goal: to motivate you to tell the truth
if you know ρ exactly.

Your goal: to maximize your expected reward
(you greedy Scrooge, you!)

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Define $f(\rho : \sigma)$ as your expected reward, from Victor, for reporting σ when the true state is ρ .

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Define $f(\rho : \sigma)$ as your expected reward, from Victor, for reporting σ when the true state is ρ .

$$f(\rho : \sigma) = \sum_i p(i) R_i(\sigma)$$

$$\text{Measurement} = \{E_i\} : \sum E_i = \mathbb{1}.$$

$$\begin{aligned} f(\rho : \sigma) &= \sum_i \text{Tr}[\rho E_i] R_i(\sigma) \\ &= \text{Tr}[\rho \mathcal{R}(\sigma)] \end{aligned}$$

where $\mathcal{R}(\sigma) = \sum_i E_i R_i(\sigma)$.

Define “value”: $G(\rho) \equiv f(\rho : \rho) = \text{Tr}[\rho \mathcal{R}(\rho)]$.

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

$$1) f(\rho : \sigma) = \sum_i p(i) R_i(\sigma)$$

$$\text{Measurement} = \{E_i\} : \sum E_i = \mathbb{1}.$$

$$\begin{aligned} f(\rho : \sigma) &= \sum_i \text{Tr}[\rho E_i] R_i(\sigma) \\ &= \text{Tr}[\rho \mathcal{R}(\sigma)] \end{aligned}$$

$$\text{where } \mathcal{R}(\sigma) = \sum_i E_i R_i(\sigma).$$

$$\text{Define "value": } G(\rho) \equiv f(\rho : \rho) = \text{Tr}[\rho \mathcal{R}(\rho)].$$

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

1) $f(\rho)$
Meas

Now, demand that the reward for truth
be greater than the reward for *any* lie.

where
Defin

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

1) $f(\rho : \rho)$
Meas

Now, demand that the reward for truth be greater than the reward for *any* lie.

2) $f(\rho : \rho) > f(\rho : \sigma)$ if $\sigma \neq \rho$.

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma)\mathcal{R}(\sigma)]$$

where
Defin

Quantum Strictly Proper Scoring Rules

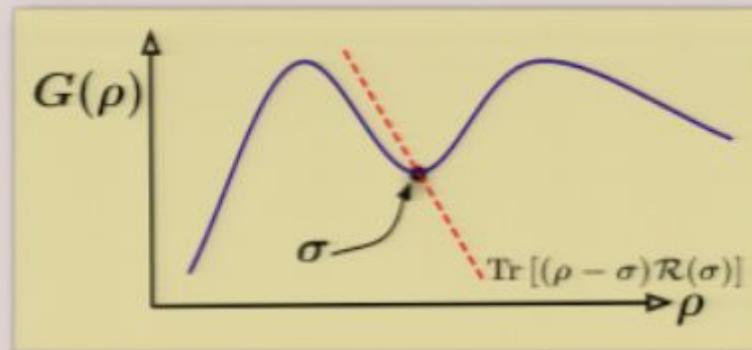
Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Now, demand that the reward for truth be greater than the reward for *any* lie.

$$2) f(\rho : \rho) > f(\rho : \sigma) \text{ if } \sigma \neq \rho.$$

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma)\mathcal{R}(\sigma)]$$



Quantum Strictly Proper Scoring Rules

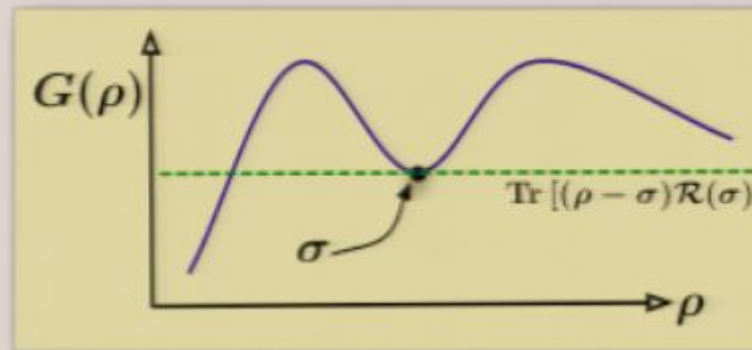
Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Now, demand that the reward for truth be greater than the reward for *any* lie.

$$2) f(\rho : \rho) > f(\rho : \sigma) \text{ if } \sigma \neq \rho.$$

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma)\mathcal{R}(\sigma)]$$



- $f(\rho : \sigma)$ is a subgradient to $G(\rho)$

Quantum Strictly Proper Scoring Rules

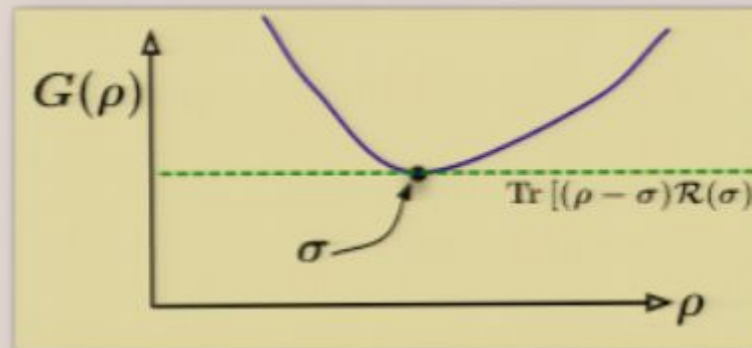
Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

Now, demand that the reward for truth be greater than the reward for *any* lie.

$$2) f(\rho : \rho) > f(\rho : \sigma) \text{ if } \sigma \neq \rho.$$

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma)\mathcal{R}(\sigma)]$$



- $f(\rho : \sigma)$ is a subgradient to $G(\rho)$
- $G(\rho)$ must be strictly convex.

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

$$1) f(\rho : \sigma) = \sum_i p(i) R_i(\sigma)$$

$$\text{Measurement} = \{E_i\} : \sum E_i = \mathbb{1}.$$

$$\begin{aligned} f(\rho : \sigma) &= \sum_i \text{Tr}[\rho E_i] R_i(\sigma) \\ &= \text{Tr}[\rho \mathcal{R}(\sigma)] \end{aligned}$$

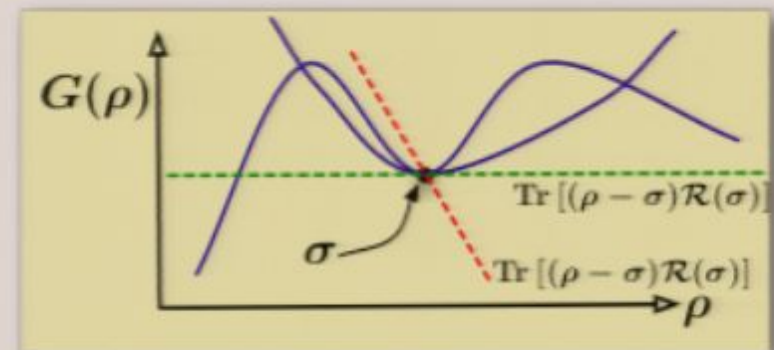
$$\text{where } \mathcal{R}(\sigma) = \sum_i E_i R_i(\sigma).$$

$$\text{Define "value": } G(\rho) \equiv f(\rho : \rho) = \text{Tr}[\rho \mathcal{R}(\rho)].$$

$$2) f(\rho : \rho) > f(\rho : \sigma) \text{ if } \sigma \neq \rho.$$

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma) \mathcal{R}(\sigma)]$$



- $f(\rho : \sigma)$ is a subgradient to $G(\rho)$
- $G(\rho)$ must be strictly convex.

Quantum Strictly Proper Scoring Rules

Victor the Verifier measures a copy of ρ , and pays you $R_i(\sigma)$ if he gets outcome i .

$$1) f(\rho : \sigma) = \sum_i p(i) R_i(\sigma)$$

$$\text{Measurement} = \{E_i\} : \sum E_i = \mathbb{1}.$$

$$\begin{aligned} f(\rho : \sigma) &= \sum_i \text{Tr}[\rho E_i] R_i(\sigma) \\ &= \text{Tr}[\rho \mathcal{R}(\sigma)] \end{aligned}$$

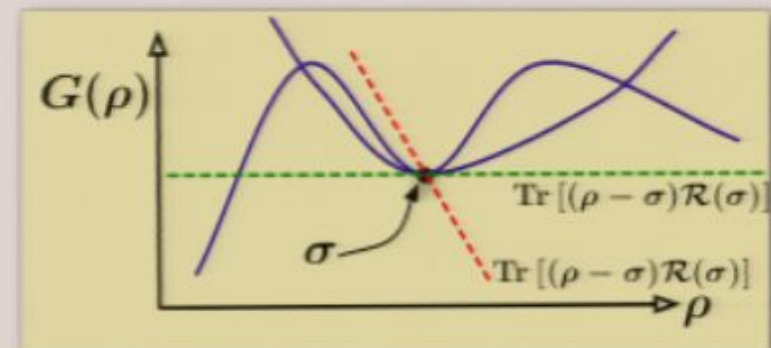
$$\text{where } \mathcal{R}(\sigma) = \sum_i E_i R_i(\sigma).$$

$$\text{Define "value": } G(\rho) \equiv f(\rho : \rho) = \text{Tr}[\rho \mathcal{R}(\rho)].$$

$$2) f(\rho : \rho) > f(\rho : \sigma) \text{ if } \sigma \neq \rho.$$

$$f(\rho : \rho) > f(\sigma : \sigma) + f(\rho : \sigma) - f(\sigma : \sigma)$$

$$G(\rho) > G(\sigma) + \text{Tr}[(\rho - \sigma) \mathcal{R}(\sigma)]$$



- $f(\rho : \sigma)$ is a subgradient to $G(\rho)$
- $G(\rho)$ must be strictly convex.

- The estimator's *expected loss* for lying is $\Delta(\rho : \sigma) \equiv G(\rho) - f(\rho : \sigma)$.

Pirsa: 0612.0042 • $\Delta(\rho : \sigma)$ is an *operational divergence* = good measure of σ 's predictive accuracy. Page 56/95

- Operational divergences are $\mathbf{1} : \mathbf{1}$ with strictly convex "entropies" $G(\rho)$.

A survey of metrics

- GOOD METRICS (OPERATIONAL DIVERGENCES)

1. L_2 distance: $\text{Tr}[(\rho - \sigma)^2] \Leftrightarrow R_i = 2 \langle i | \sigma | i \rangle - \text{Tr}(\sigma^2) - 1$

2. Relative entropy: $S(\rho || \sigma) = \text{Tr} \left[\rho \log \frac{\sigma}{\rho} \right] \Leftrightarrow R_i = \ln(\langle i | \sigma | i \rangle)$

- BAD METRICS

1. Overlap: $1 - \text{Tr}[\rho\sigma]$ *Improper (motivates lying)*

2. Trace-norm: $\|\rho - \sigma\|_1$

3. Fidelity: $1 - \left[\text{Tr} \sqrt{\sqrt{\sigma} \rho \sqrt{\sigma}} \right]^2$

Non-operational

A survey of metrics

- GOOD METRICS (OPERATIONAL DIVERGENCES)

1. L_2 distance: $\text{Tr}[(\rho - \sigma)^2] \Leftrightarrow R_i = 2 \langle i | \sigma | i \rangle - \text{Tr}(\sigma^2) - 1$

2. Relative entropy: $S(\rho || \sigma) = \text{Tr} \left[\rho \log \frac{\sigma}{\rho} \right] \Leftrightarrow R_i = \ln(\langle i | \sigma | i \rangle)$

- BAD METRICS

1. Overlap: $1 - \text{Tr}[\rho\sigma]$ *Improper (motivates lying)*

2. Trace-norm: $\|\rho - \sigma\|_1$

3. Fidelity: $1 - \left[\text{Tr} \sqrt{\sqrt{\sigma} \rho \sqrt{\sigma}} \right]^2$ *Non-operational*

Operational divergences compare
physical states to classical descriptions.

So... what do we do?

- 1. Motivation:** problems with MLE.
- 2. Foundation:** quantum scoring rules.
- 3. Solution:** the BME algorithm (& its features)
- 4. Testing:** dueling procedures! (numerics)

So... what do we do?

1. Motivation: problems with MLE.

2. Foundation: quantum scoring rules.

3. Solution: the BME algorithm (& its features)

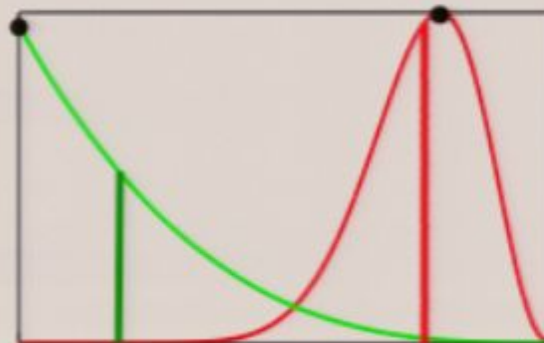
4. Testing: dueling procedures! (numerics)

Bayesian Mean Estimation

- Procedure: report the mean of a probability distribution obtained via Bayes' Rule.
- Begin with a *prior distribution*, $\pi_0(\rho)d\rho$.
- Update to $\pi(\rho) = \mathcal{L}(\rho)\pi_0(\rho)$.
- Compute the mean value: $\hat{\sigma} = \int \hat{\rho} \pi(\rho)d\rho$

MLE reports the **mode**

BME reports the **mean**



Bayesian Inference

1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Bayesian Inference

1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle +|, |-\rangle\langle -|, |+i\rangle\langle +i|, |-i\rangle\langle -i|\}$$

Bayesian Inference

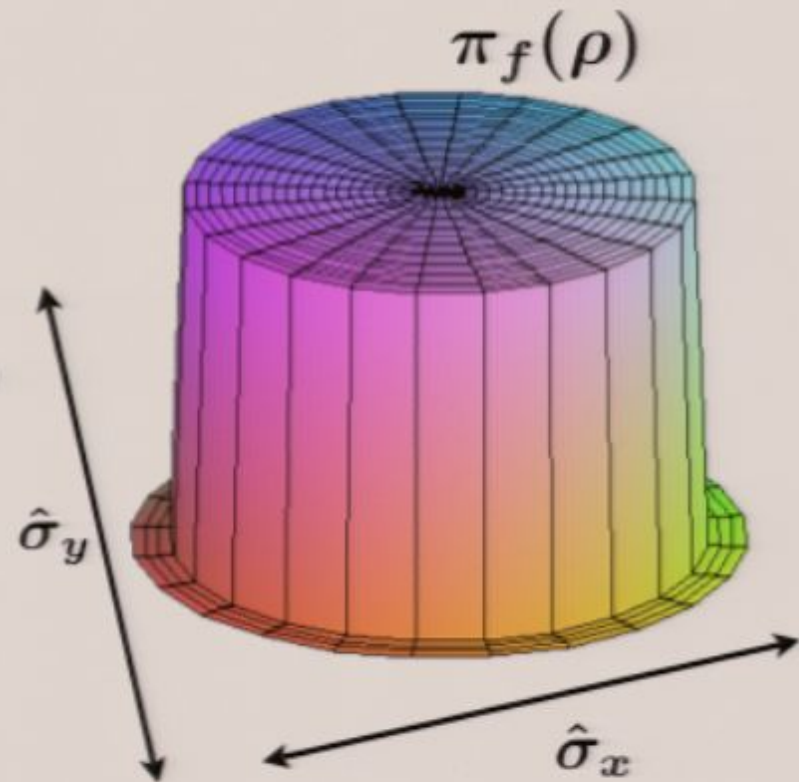
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $M = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

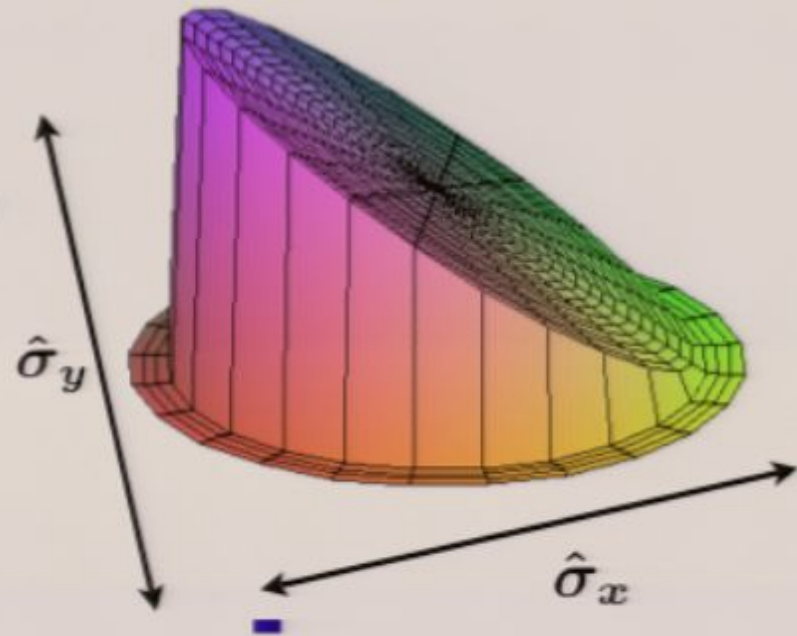
Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$

$$\pi_f(\rho)$$



Bayesian Inference

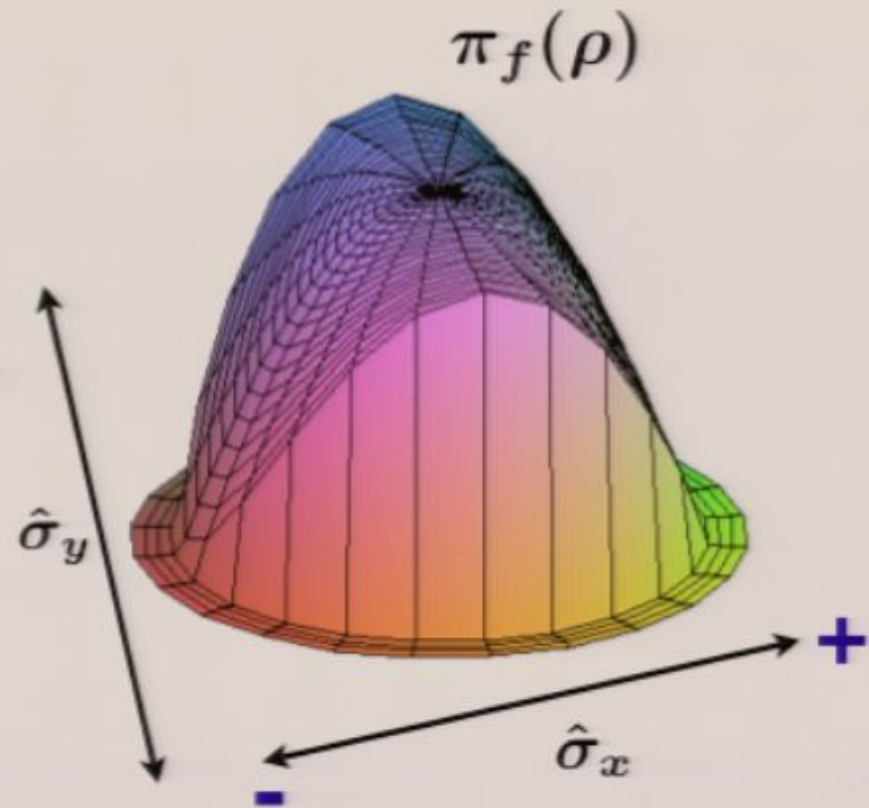
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+i\rangle\langle+i|, |-i\rangle\langle-i|\}$$



Bayesian Inference

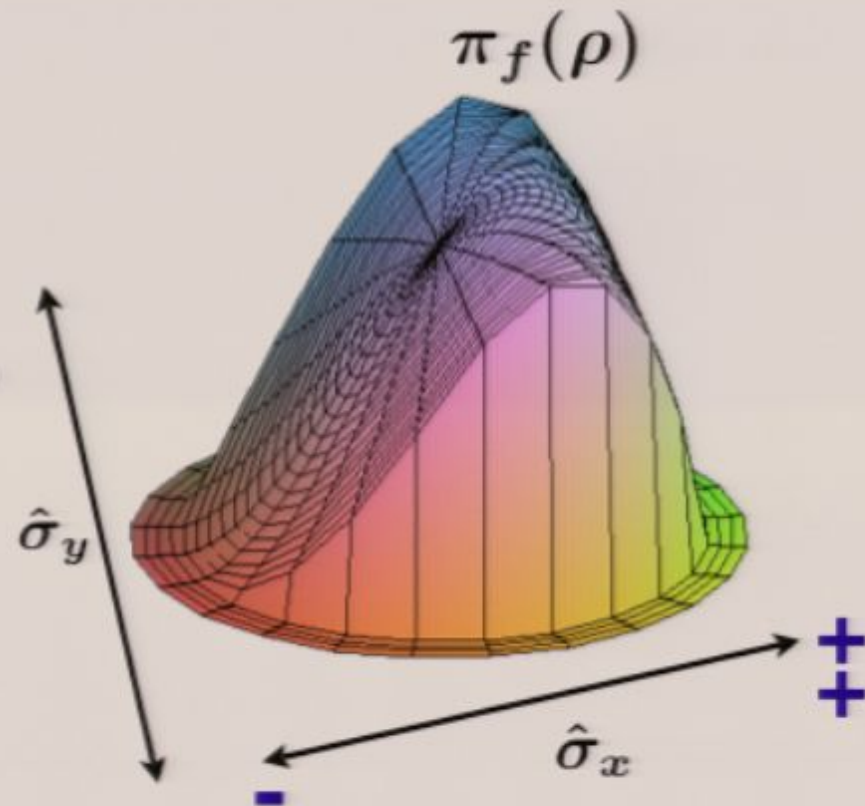
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+i\rangle\langle+i|, |-i\rangle\langle-i|\}$$



Bayesian Inference

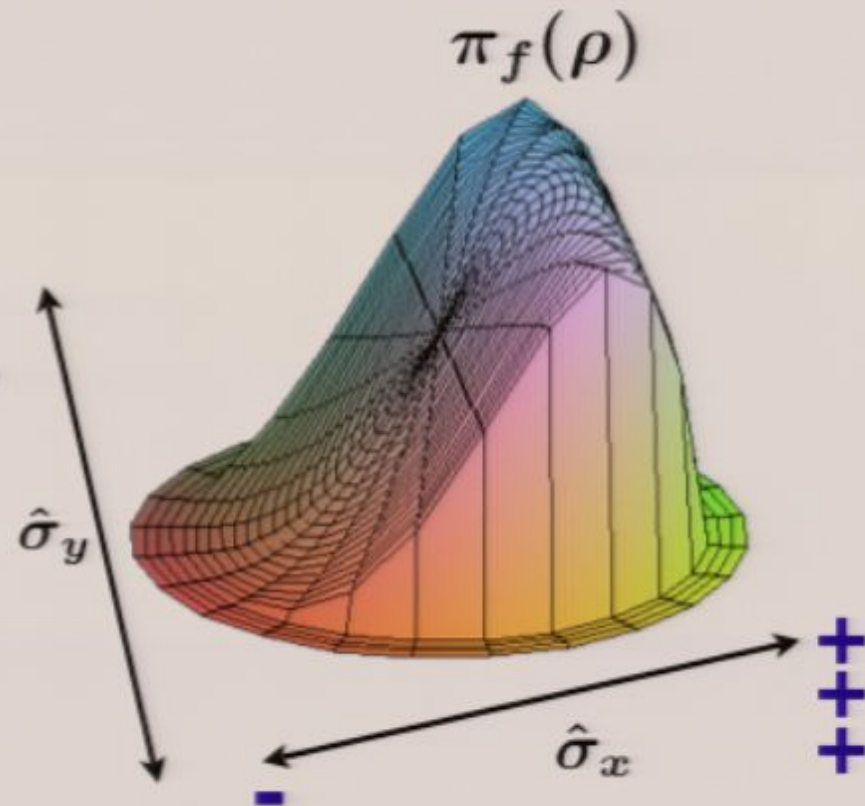
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $M = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

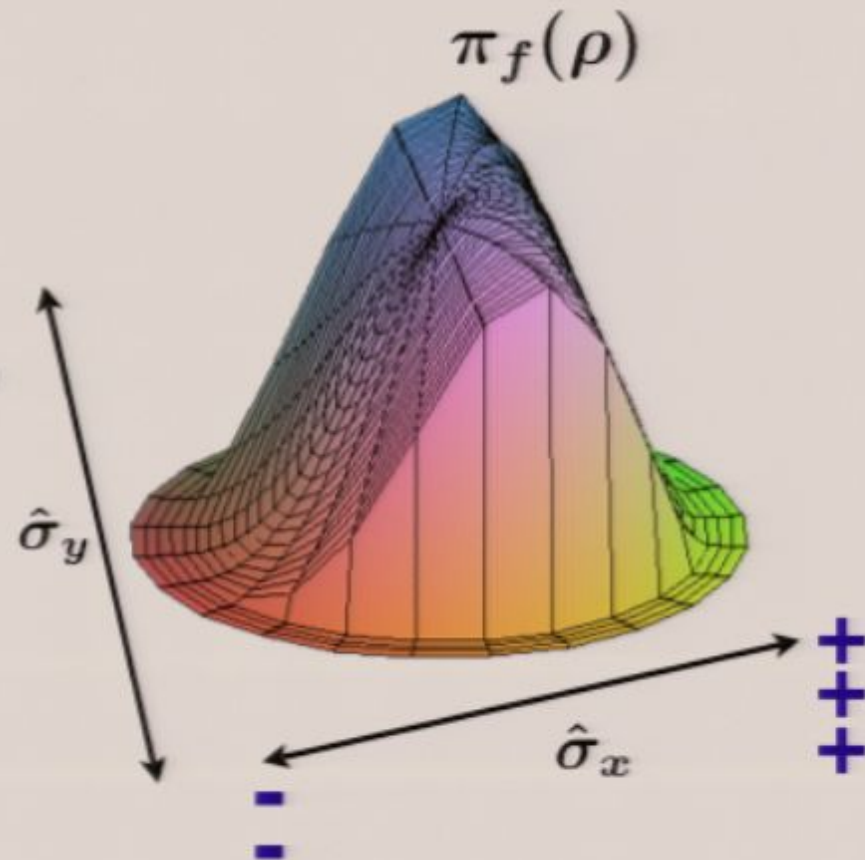
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $M = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

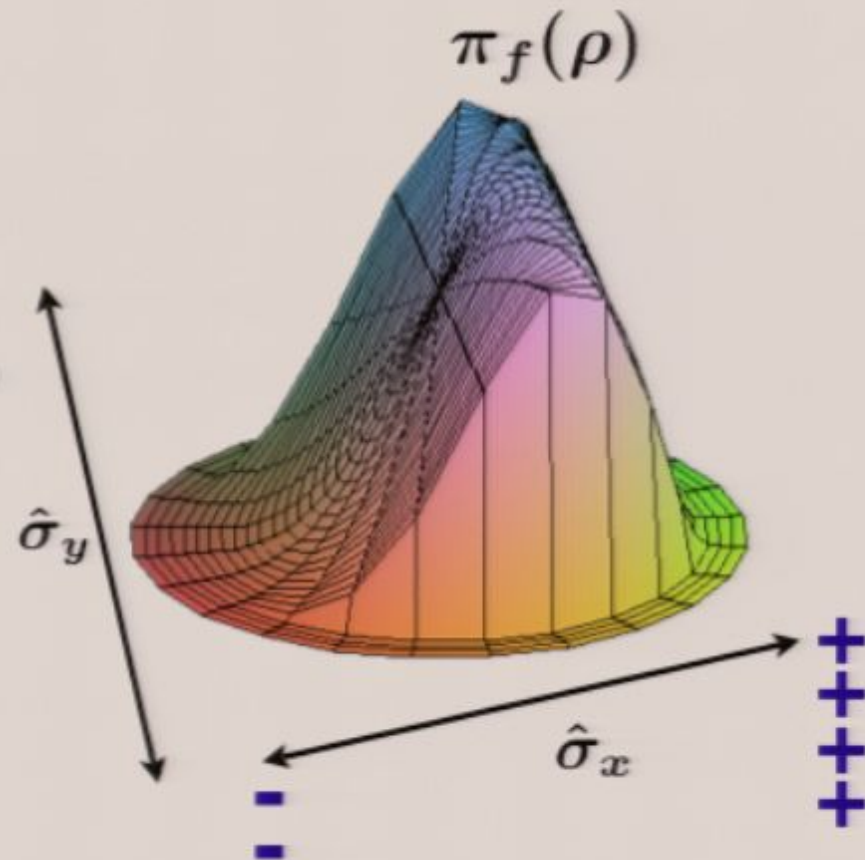
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

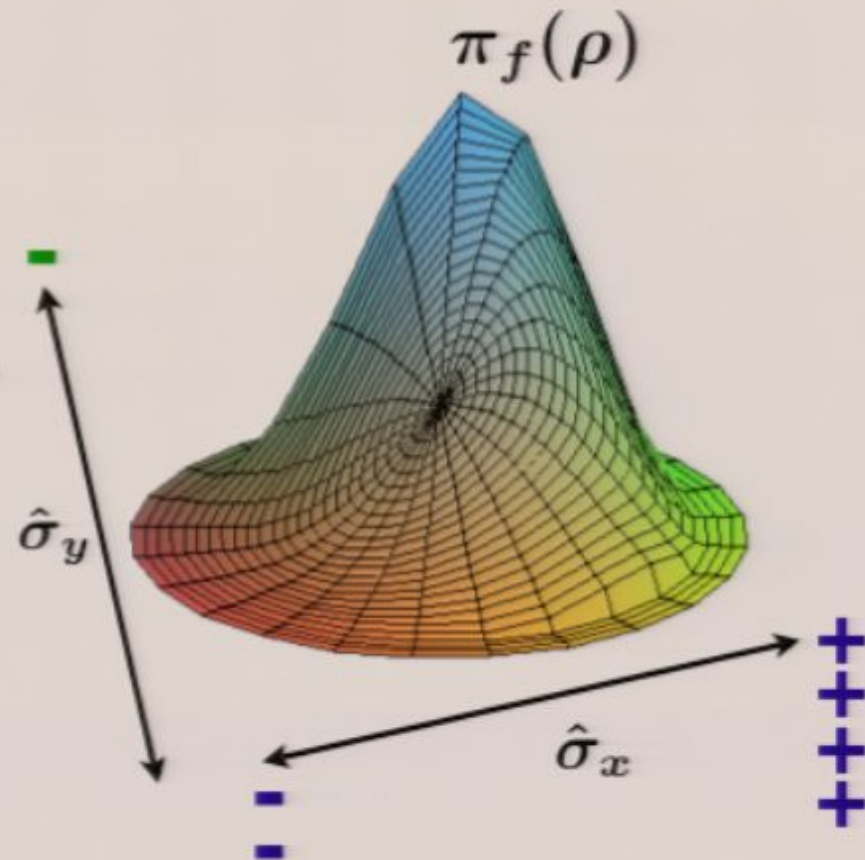
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

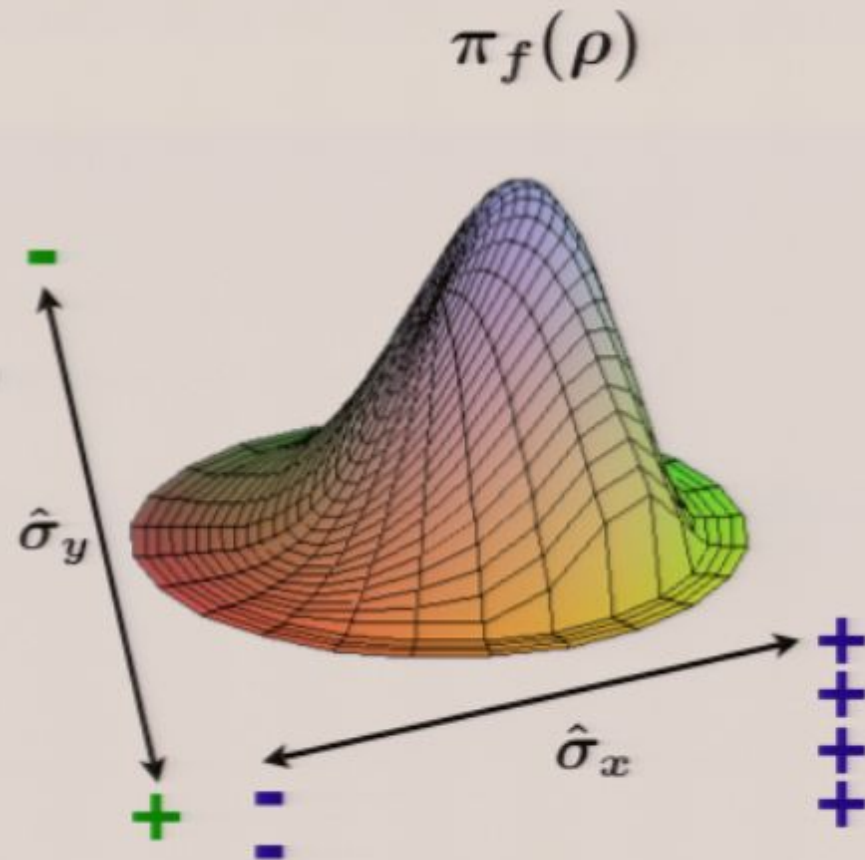
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+i\rangle\langle+i|, |-i\rangle\langle-i|\}$$



Bayesian Inference

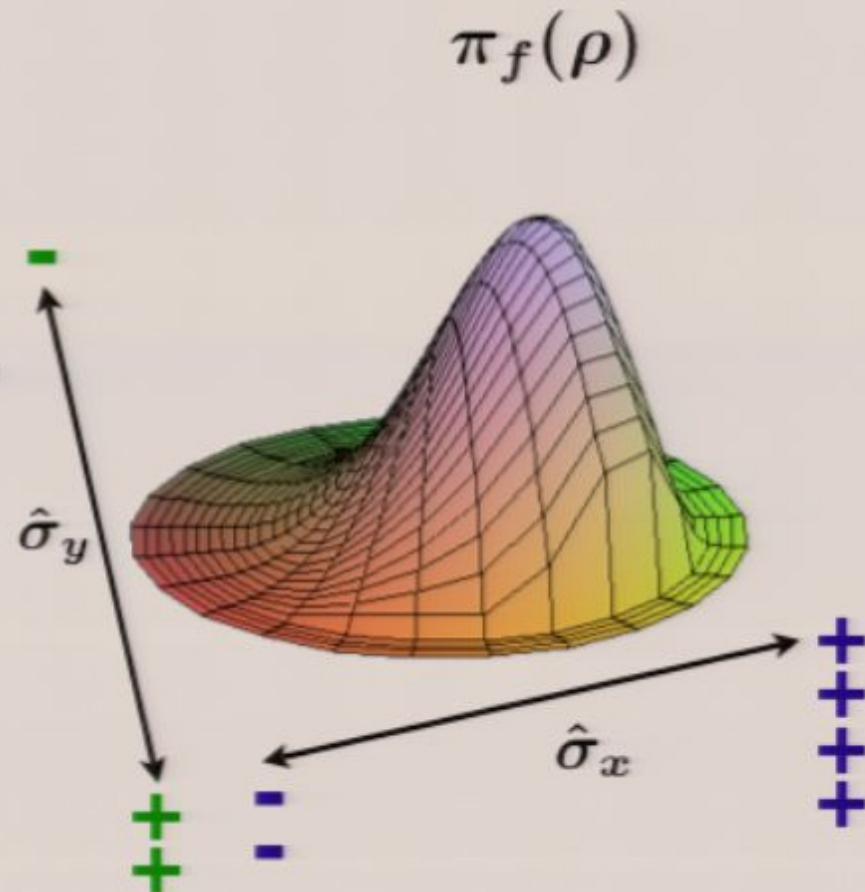
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+i\rangle\langle+i|, |-i\rangle\langle-i|\}$$



Bayesian Inference

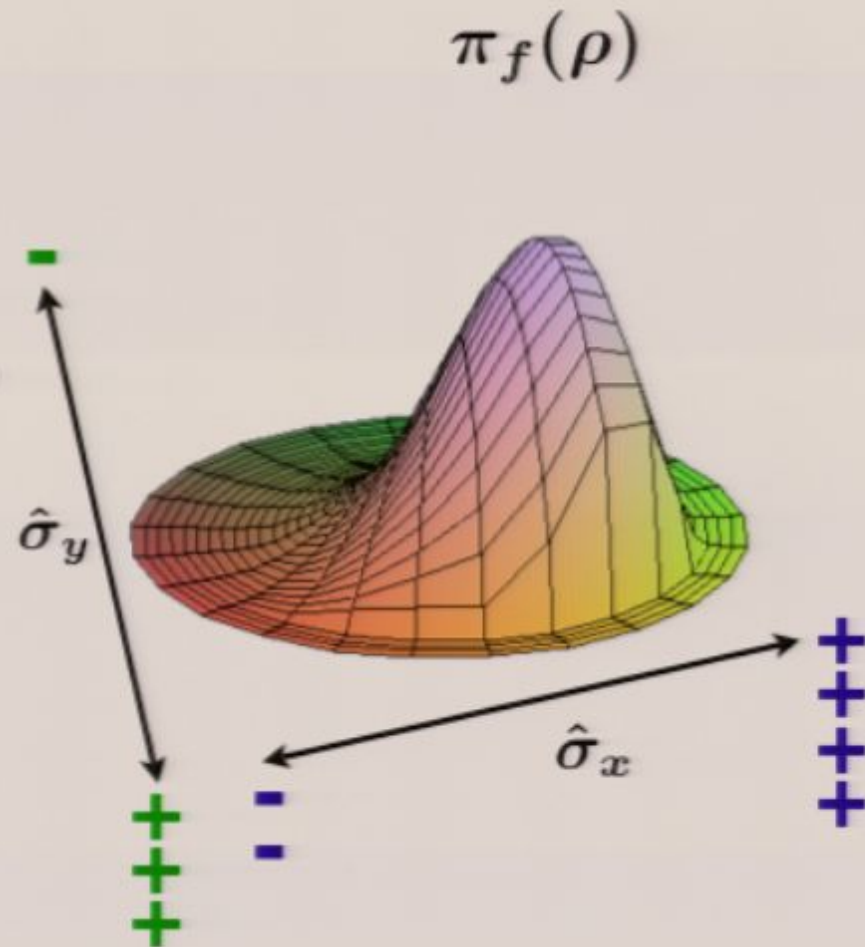
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+i\rangle\langle+i|, |-i\rangle\langle-i|\}$$



Bayesian Inference

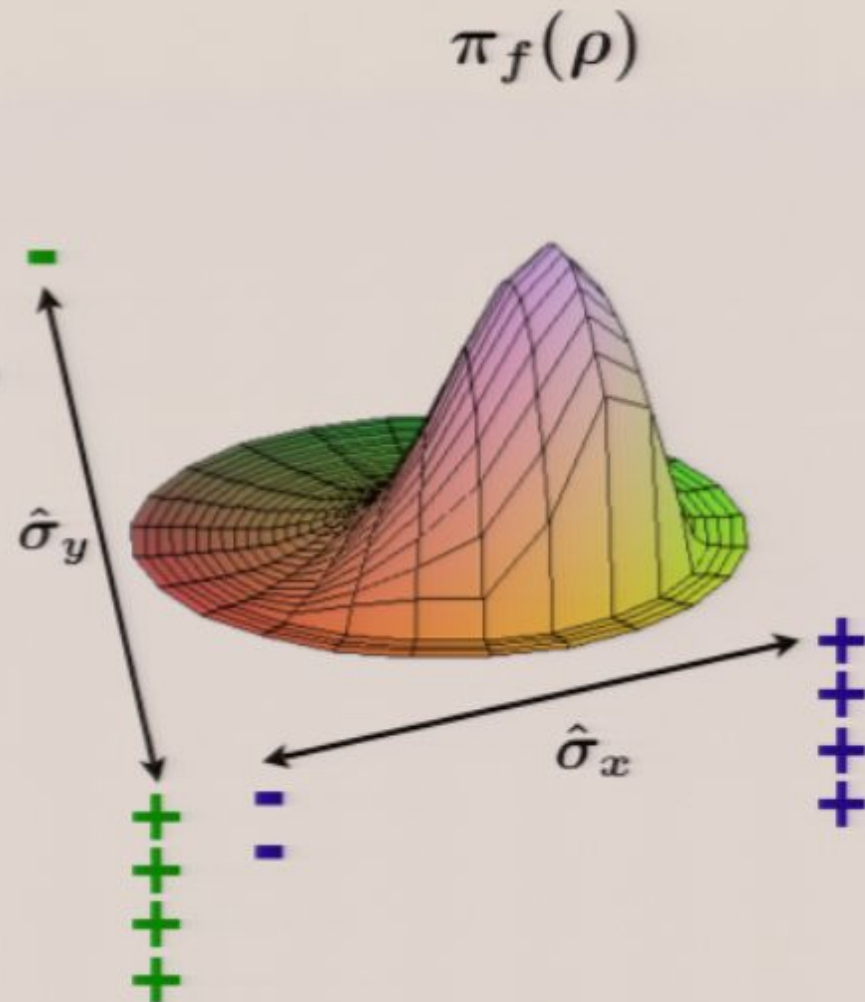
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

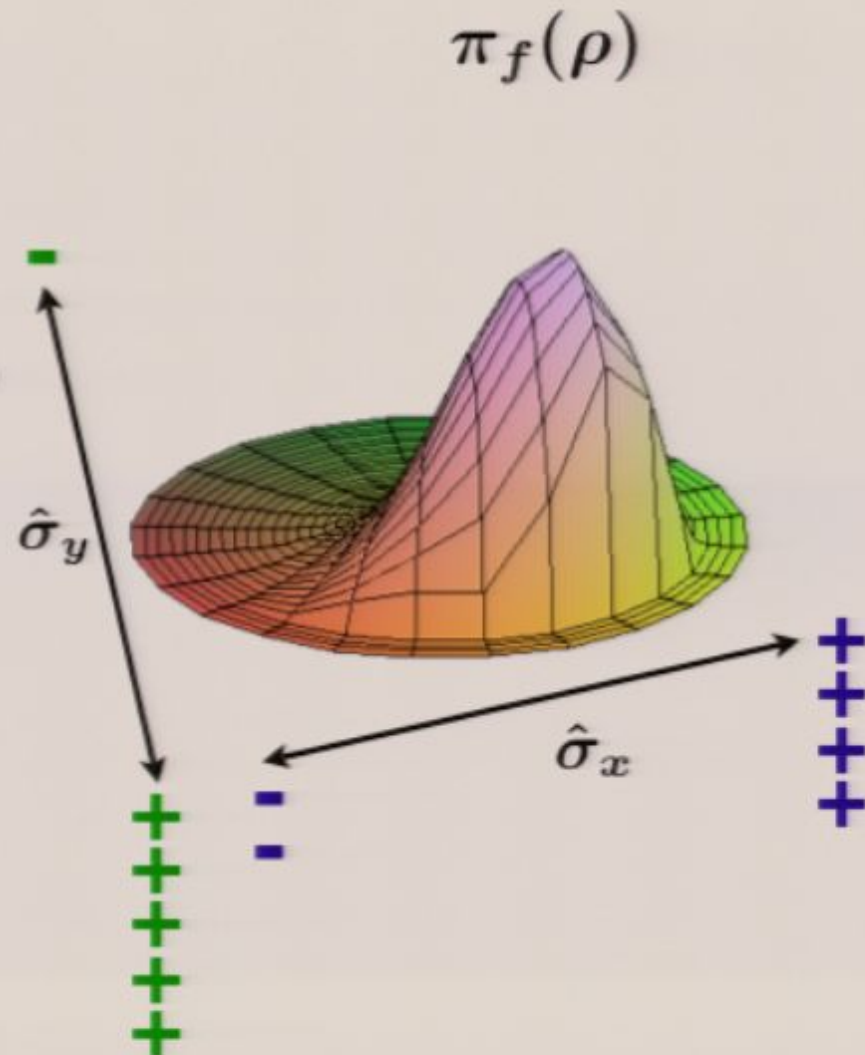
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

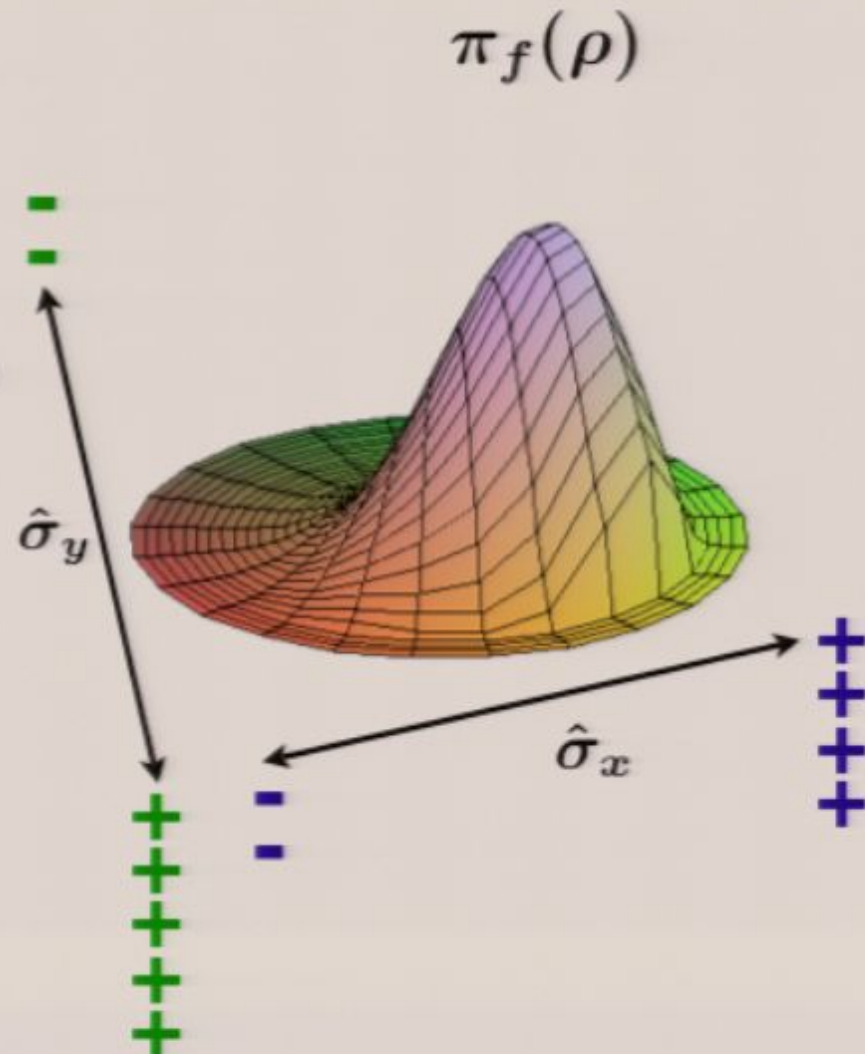
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $\mathbf{M} = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



Bayesian Inference

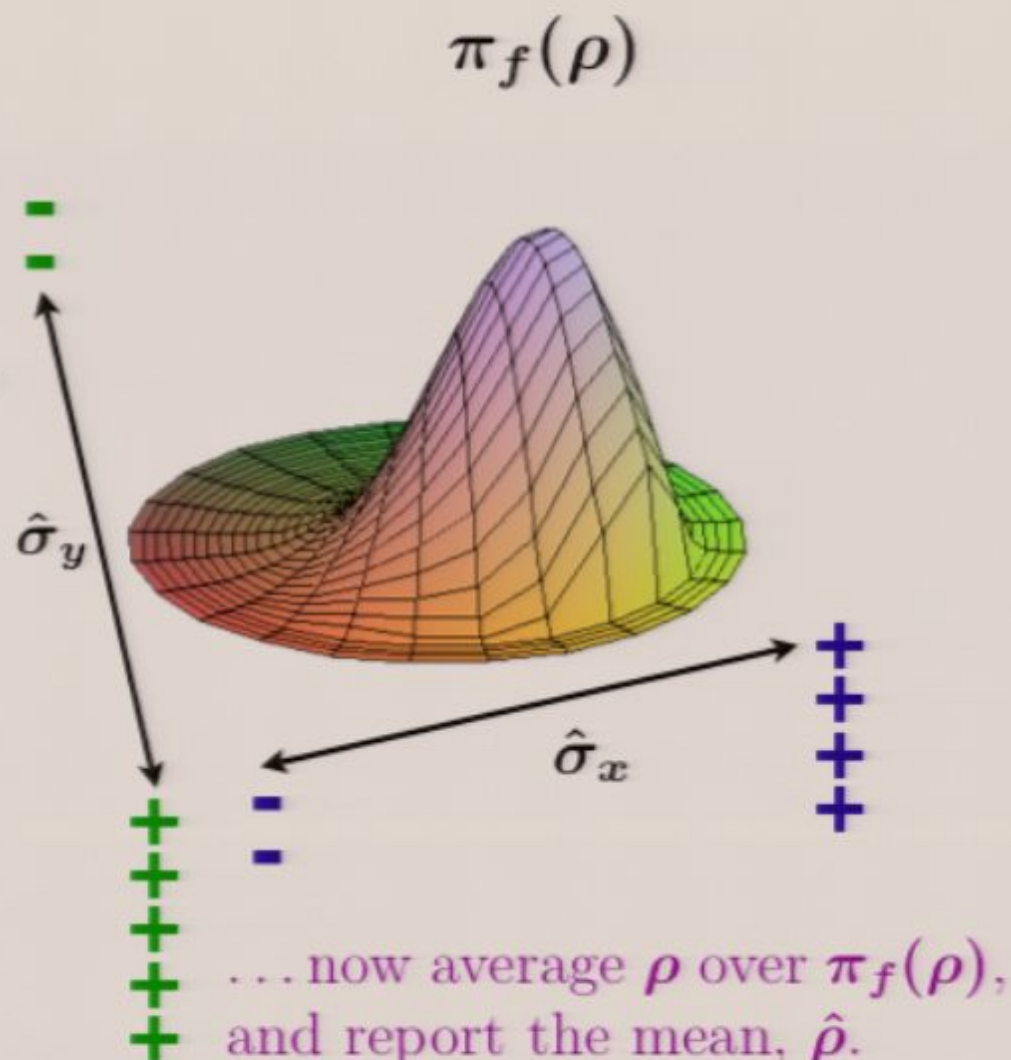
1. Begin with a prior: $\pi_0(\rho)d\rho$.
2. Make N measurements
 $\Rightarrow$ record: $M = \{\hat{M}_1, \hat{M}_2, \dots, \hat{M}_N\}$
3. Bayes' Rule: $\pi_f(\rho)d\rho \propto \mathcal{L}(\rho)\pi_0(\rho)d\rho$
 $= \text{Tr}(\hat{M}_1\rho)\text{Tr}(\hat{M}_2\rho) \dots \text{Tr}(\hat{M}_N\rho)\pi_0(\rho)d\rho$

Example: spin- $1/2$

Measure Y, X (ignore Z)

POVM operators are:

$$\hat{E}_i \in \{|+\rangle\langle+|, |-\rangle\langle-|, |+\rangle\langle+i|, |-\rangle\langle-i|\}$$



ρ_{BME} is strictly positive

- BME never yields a rank-deficient estimate!
- Proof sketch:

For any $|\psi\rangle$, let $\pi_0(\rho)d\rho$ have support on a measurable set of states *not* orthogonal to $|\psi\rangle$.

Each measurement $\hat{M}_i$ removes a measure-zero set of states from the support.

So $\pi_0(\rho)d\rho$ contains states not orthogonal to $|\psi\rangle$.

These contribute to ρ_{BME} , so $\langle\psi|\rho_{\text{BME}}|\psi\rangle > 0$.

$$P(\omega_j) = \text{Tr}(\rho \pi_j)$$

$$\rho = \sum_j p_j \hat{E}_j$$

$$= \text{Tr}(c_j \rho \rho^{-1/2} Q_j \rho^{-1/2})$$

$$\text{Tr}(c_j \hat{Q}_j)$$



$$P(\omega_j) = \text{Tr}(\rho \pi_j)$$

$$\hookrightarrow \rho = \sum_j p_j \hat{E}_j$$

$$= \text{Tr}(\sum_j p_j \rho^{-1/2} Q_j \rho^{-1/2})$$

$$\text{Tr}(\sum_j \hat{Q}_j)$$



ρ_{BME} is strictly positive

- BME never yields a rank-deficient estimate!
- Proof sketch:

For any $|\psi\rangle$, let $\pi_0(\rho)d\rho$ have support on a measurable set of states *not* orthogonal to $|\psi\rangle$.

Each measurement $\hat{M}_i$ removes a measure-zero set of states from the support.

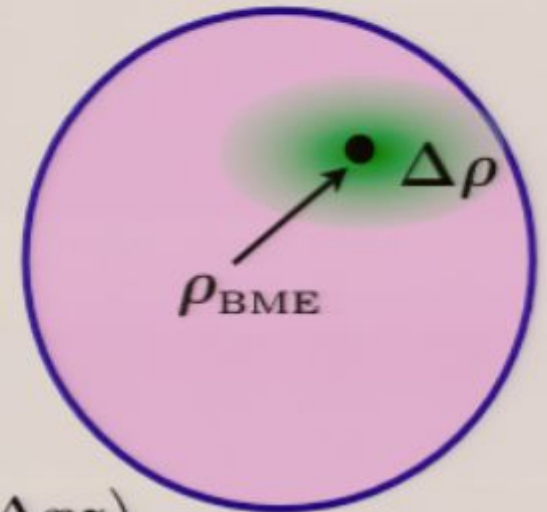
So $\pi_0(\rho)d\rho$ contains states not orthogonal to $|\psi\rangle$.

These contribute to ρ_{BME} , so $\langle\psi|\rho_{\text{BME}}|\psi\rangle > 0$.

Free error bars with every estimate!

- Estimate is a vector in Hilbert-Schmidt space:

$$\rho = \begin{pmatrix} \langle x \rangle \\ \langle y \rangle \\ \langle z \rangle \end{pmatrix}$$



- Uncertainty is a covariance matrix on H-S space:

$$\Delta \rho = \begin{pmatrix} \Delta x^2 & \Delta xy & \Delta xz \\ \Delta xy & \Delta y^2 & \Delta yz \\ \Delta xz & \Delta yz & \Delta z^2 \end{pmatrix}$$

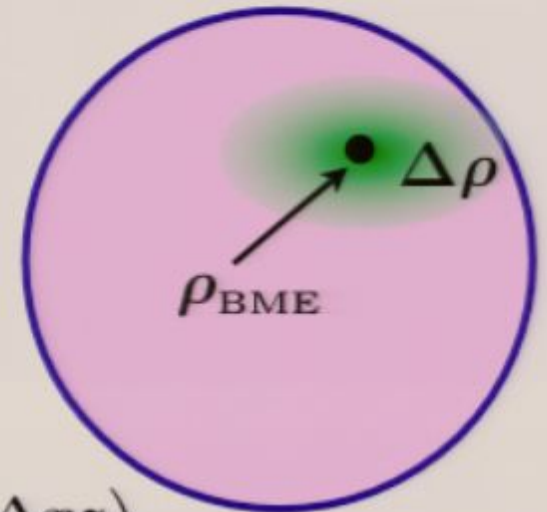
- I.e., a superoperator, with Jamiolkowski form:

$$\begin{aligned} \Delta \rho_{BME} &= \overline{\rho \otimes \rho} - \bar{\rho} \otimes \bar{\rho} \\ &= \int (\rho \otimes \rho) \pi(\rho) d\rho - \rho_{BME} \otimes \rho_{BME} \end{aligned}$$

Free error bars with every estimate!

- Estimate is a vector in Hilbert-Schmidt space:

$$\rho = \begin{pmatrix} \langle x \rangle \\ \langle y \rangle \\ \langle z \rangle \end{pmatrix}$$



- Uncertainty is a covariance matrix on H-S space:

$$\Delta \rho = \begin{pmatrix} \Delta x^2 & \Delta xy & \Delta xz \\ \Delta xy & \Delta y^2 & \Delta yz \\ \Delta xz & \Delta yz & \Delta z^2 \end{pmatrix}$$

- I.e., a superoperator, with Jamiolkowski form:

$$\begin{aligned} \Delta \rho_{BME} &= \overline{\rho \otimes \rho} - \bar{\rho} \otimes \bar{\rho} \\ &= \int (\rho \otimes \rho) \pi(\rho) d\rho - \rho_{BME} \otimes \rho_{BME} \end{aligned}$$

Accuracy:

BME optimizes every operational divergence

If you *know* ρ , the optimal estimate is $\sigma = \rho$.

But what if your knowledge is uncertain (probabilistic)?

$\implies$ the state could be ρ_i with probability π_i .

KEY FACT: $f(\rho : \sigma) = \text{Tr}[\rho \mathcal{R}(\sigma)]$ is *linear* in ρ .

$\implies$ *expected* utility is

$$\bar{f} = \sum \pi_i f(\rho_i : \sigma) = f\left(\sum \pi_i \rho_i : \sigma\right)$$

$\implies$ optimal estimate is $\sigma = \bar{\rho} \equiv \sum \pi_i \rho_i$

This is fairly straightforward, but tedious... see [quant-ph/0603116](https://arxiv.org/abs/quant-ph/0603116)

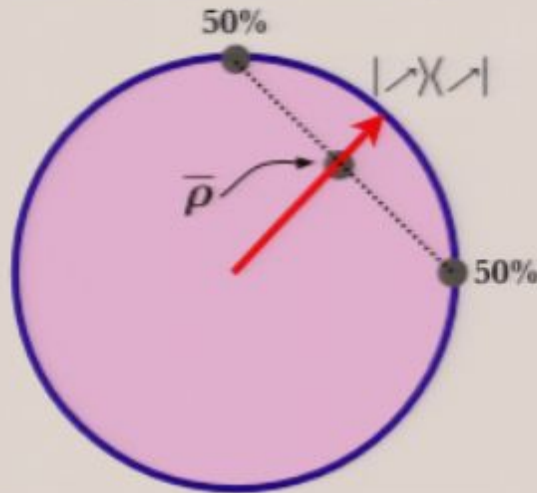
If unknown ρ was selected from distribution $\pi_0(\rho)d\rho$, and measurements $\mathcal{M} = \{E_1, E_2 \dots\}$ were made, then:

1. Your knowledge is $\pi(\rho)d\rho = \frac{p(\mathcal{M}|\rho)\pi_0(\rho)d\rho}{\int p(\mathcal{M}|\rho)\pi_0(\rho)d\rho}$.
2. The optimal estimate is $\bar{\rho} = \int \rho \pi(\rho)d\rho$.

But isn't the mean
an obvious choice?

But isn't the mean an obvious choice? **NO**

(a) Suppose we optimize fidelity.



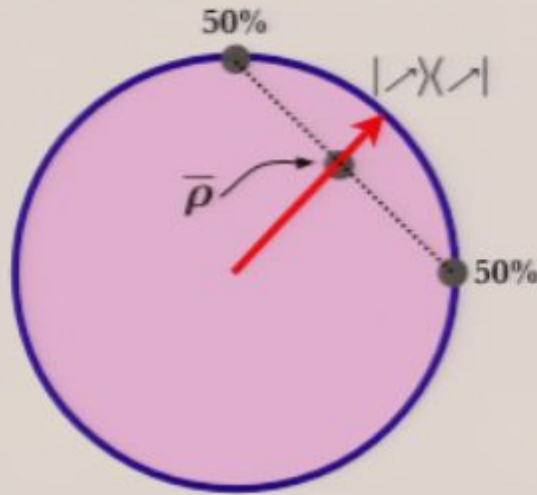
$$\overline{F}(\rho_i, \bar{\rho}) = \frac{3}{4}$$

$$\overline{F}(\rho_i, |\nearrow\rangle\langle\nearrow|) = \frac{3+2\sqrt{2}}{4+2\sqrt{2}} \approx 0.85$$

MORAL: fidelity measures how well σ
simulates ρ , not how well it *estimates* ρ .

But isn't the mean an obvious choice? **NO**

(a) Suppose we optimize fidelity.



$$\overline{F}(\rho_i, \bar{\rho}) = \frac{3}{4}$$

$$\overline{F}(\rho_i, |\nearrow\rangle\langle\nearrow|) = \frac{3+2\sqrt{2}}{4+2\sqrt{2}} \approx 0.85$$

MORAL: fidelity measures how well σ *simulates* ρ , not how well it *estimates* ρ .

(b) Suppose we assume the future will look (statistically) just like the past.

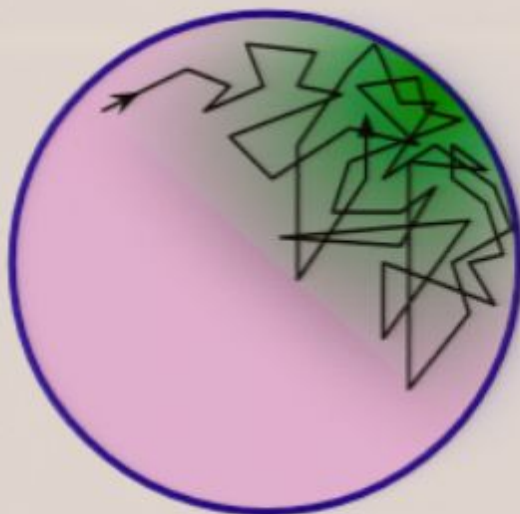
- We know what future datasets will look like
- We can add them to measurement record
- $\mathcal{M} \longrightarrow \mathcal{M} \cup \mathcal{M} \cup \mathcal{M} \cup \dots$
- $p(\mathcal{M}|\rho) \longrightarrow p(\mathcal{M}|\rho)^\infty$
- So $\pi(\rho)d\rho \rightarrow \delta(\rho - \rho_{\text{MLE}})$

MORAL: MLE can be derived by assuming this “frequentist axiom”.

(This explains a lot...)

Implementing BME via Metropolis-Hastings

- Crawl over Hilbert-Schmidt space, integrating as we go.
- Jump from $\rho \rightarrow \rho'$ w/prob $p_{\text{jump}} = \max\left(\frac{\mathcal{L}(\rho')}{\mathcal{L}(\rho)}, 1\right)$



- Seems to scale as N^4 - N^6 (slower than MLE)

Is BME really optimal?

- 1. Motivation:** problems with MLE.
- 2. Foundation:** quantum scoring rules.
- 3. Solution:** the BME algorithm (& its features)
- 4. Testing:** dueling procedures! (numerics)

Is BME really optimal?

1. Motivation: problems with MLE.

2. Foundation: quantum scoring rules.

3. Solution: the BME algorithm (& its features)

4. Testing: dueling procedures! (numerics)

How good is MLE?

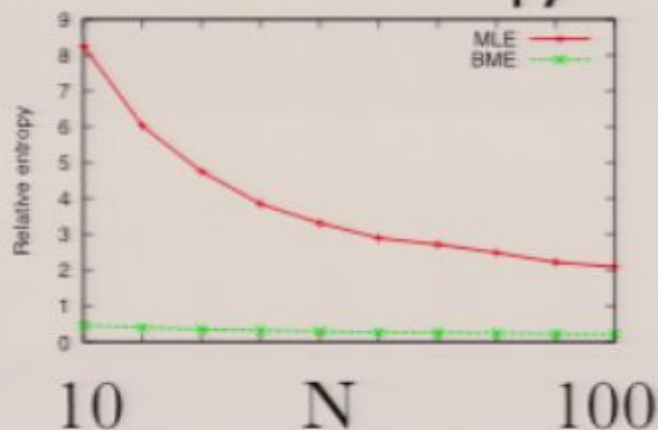
NUMERICS:

1. Generate random (Hilbert-Schmidt measure) mixed states ρ for n qubits,
2. Measure each of 4^{n-1} Pauli observables N times,
3. Analyze measurement record to get σ ,

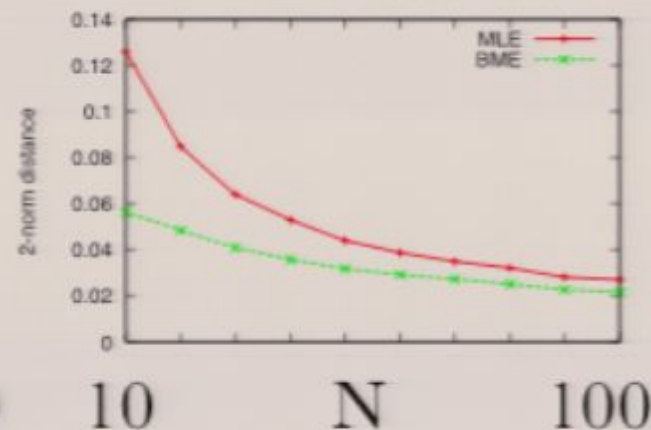
4. Compare σ to ρ using:

- (a) L2-norm,
 $\text{Tr} [(\rho - \sigma)^2],$
- (b) relative entropy,
 $\text{Tr} [\rho(\ln \rho - \ln \sigma)],$
- (c) fidelity,
 $\text{Tr} \left[\sqrt{\sqrt{\sigma} \rho \sqrt{\sigma}} \right]^2$
- (d) L1-norm,
 $\text{Tr} [|\rho - \sigma|].$

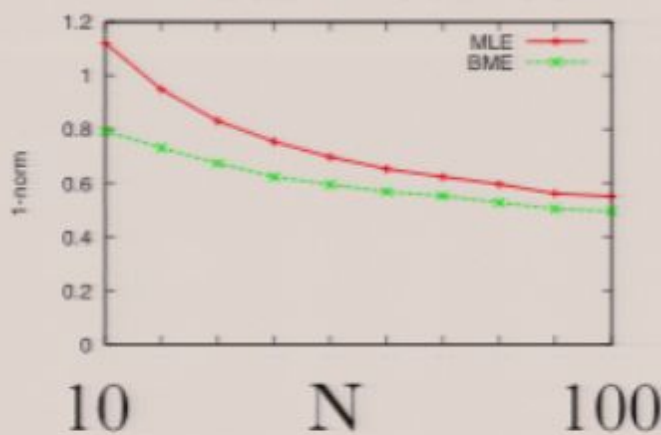
Relative Entropy



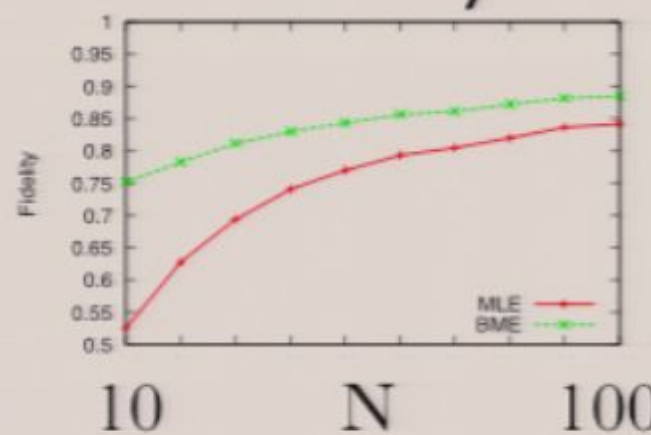
L-2 distance



Trace Distance



Fidelity



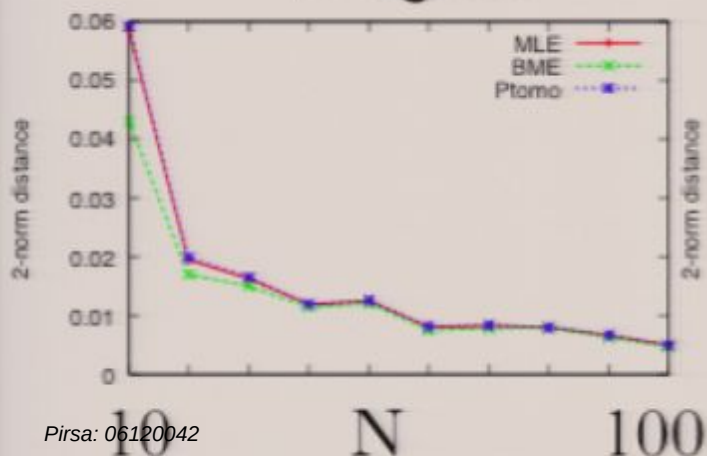
All data is for 3 qubits

Quick & Dirty Tomography

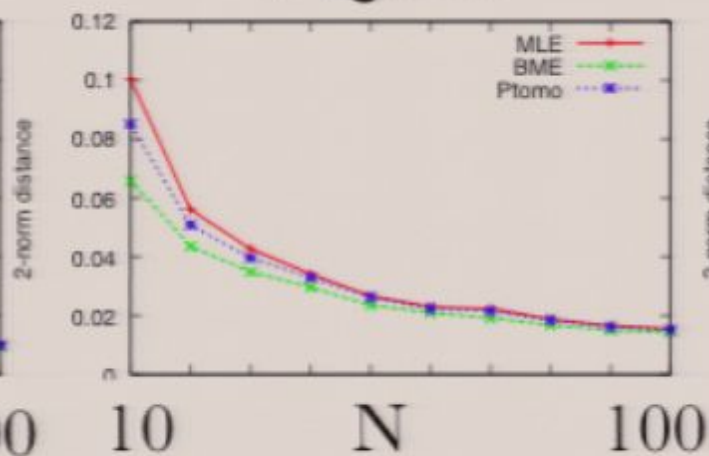
1. MLE gets computationally difficult for large systems.
2. BME is even harder!
3. $\frac{1}{\sqrt{3}}$ ($|\text{D.James}\rangle + |\text{T.Havel}\rangle + |\text{P.Jessen}\rangle$) suggested “quick 'n' dirty tomography”.

$$\sigma_{\text{tomo}} = \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & -\lambda_3 & \\ & & & -\lambda_4 \end{pmatrix} \rightarrow \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & 0 & \\ & & & 0 \end{pmatrix} \rightarrow \frac{1}{\lambda_1 + \lambda_2} \begin{pmatrix} \lambda_1 & & & \\ & \lambda_2 & & \\ & & 0 & \\ & & & 0 \end{pmatrix} = \sigma.$$

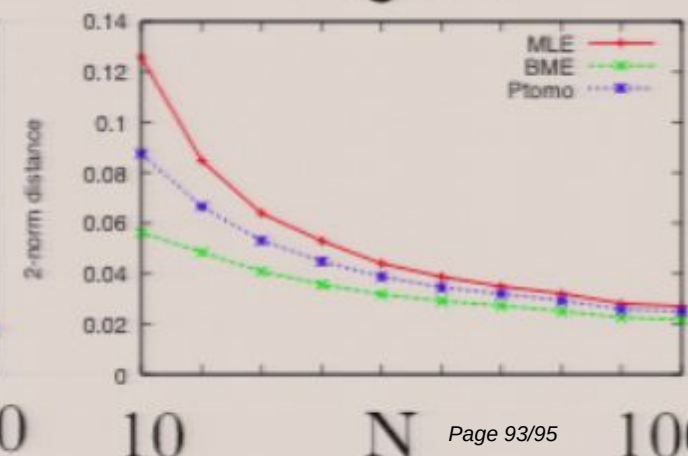
1 Qubit



2 Qubits



3 Qubits



Conclusions

- It's worth thinking carefully and deeply about state estimation.
- Scoring rules provide a *reliable* experimental way to compare multiple estimates .
- BME is optimal, and a baseline for evaluating other (more efficient) approaches.
- Quick & dirty methods can outperform MLE.

How good is MLE?

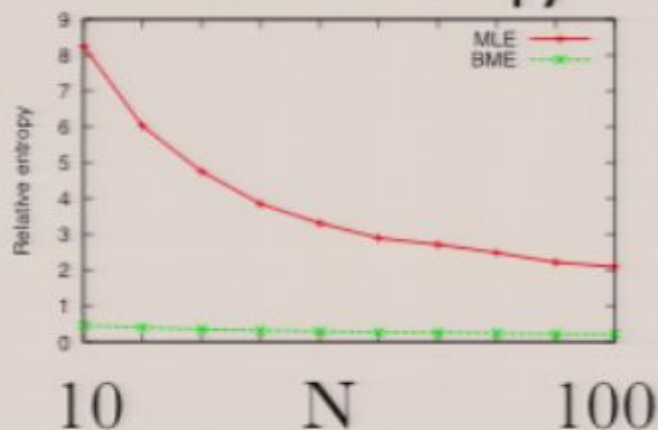
NUMERICS:

1. Generate random (Hilbert-Schmidt measure) mixed states ρ for n qubits,
2. Measure each of 4^{n-1} Pauli observables N times,
3. Analyze measurement record to get σ ,

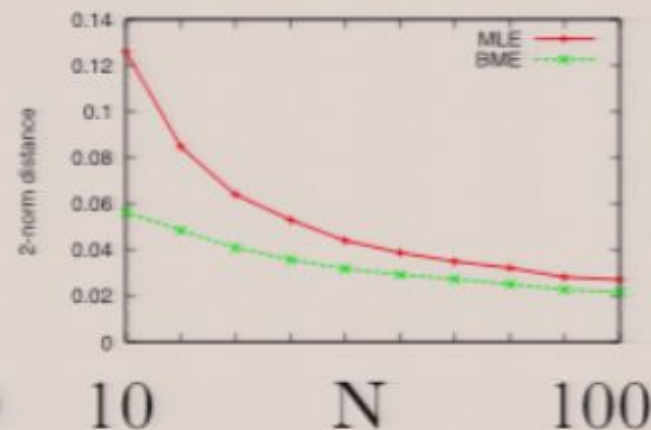
4. Compare σ to ρ using:

- (a) L2-norm,
 $\text{Tr} [(\rho - \sigma)^2],$
- (b) relative entropy,
 $\text{Tr} [\rho(\ln \rho - \ln \sigma)],$
- (c) fidelity,
 $\text{Tr} \left[\sqrt{\sqrt{\sigma} \rho \sqrt{\sigma}} \right]^2$
- (d) L1-norm,
 $\text{Tr} [|\rho - \sigma|].$

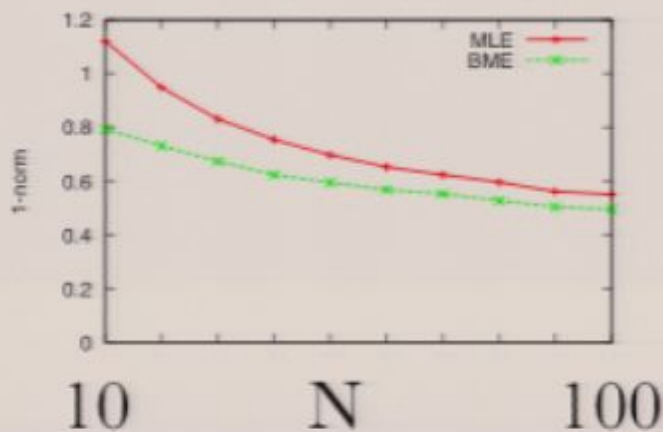
Relative Entropy



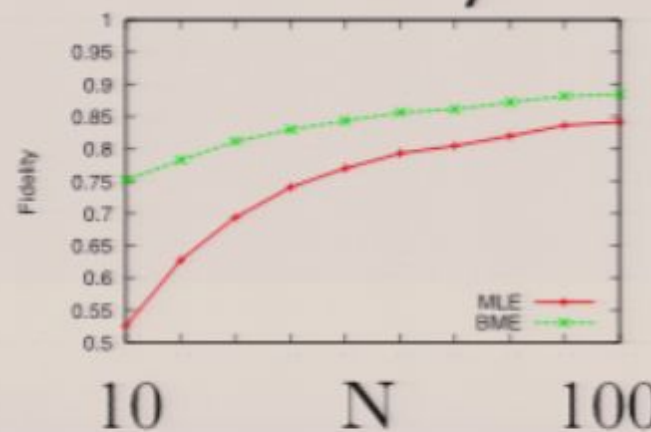
L-2 distance



Trace Distance



Fidelity



All data is for 3 qubits