

Calibrating an Ion-Trap Quantum Computer

by Brennan de Neeve

Supervised by
Vlad Negnevitsky and
Matteo Marinelli

Supervising Professor
Prof. Dr. Jonathan Home
Trapped Ion Quantum Information Group
ETH Zürich

July 2017

Acknowledgements

It's been a lot of fun learning about quantum information and doing experiments with the TIQI group. I'd like to first thank Jonathan for giving me an opportunity to do such interesting research! I'd like to also thank Vlad for all his help throughout my project, for numerous discussions about the control system, C++ and Ionizer, also for recommending me to change my operating system to Arch Linux (a decision I definitely don't regret!), and most of all for his help and patience in realising all the experiments presented in this thesis.

I'd like to thank Matteo for useful discussions about Ionizer, AOM setups, and calibration techniques; Christa for lots of fruitful and interesting discussions about quantum information and related topics, and for helping me debug my magnetic field stabilisation algorithm; Robin for always letting me know about interesting or useful papers; Karan for discussions about calibration techniques and MS gates; Oliver for kicker and ice cream; Thomas Harty for some useful advice during his brief visit to the group; and Renato Renner for giving an excellent lecture on quantum information theory – the concepts I learned there were very useful throughout my project.

Finally, I'd like to thank Evert van Nieuwenburg for some interesting and encouraging discussions about machine learning and neural networks. I only began to explore these techniques a little in my project, but it's likely these methods will become increasingly important to generalise the techniques I investigate here to more complex experiments, and I hope to explore these methods further in the future.

Abstract

Although much of the theory of quantum error correction is now well-established, it remains a challenge to experimentally achieve error rates below the thresholds set by quantum error correction theorems. Maintaining low coherent errors for experiments of increasing complexity is an important part of this challenge, where accurate knowledge about the physical system is key. As a preliminary effort to reduce coherent errors, this thesis investigates the use of phase estimation for single-qubit gate calibration.

In chapter 3 an algorithm used to stabilise the frequency of a $^{40}\text{Ca}^+$ ion qubit in the presence of harmonic noise is described.

In chapter 4, a non-adaptive phase estimation technique is implemented to calibrate laser pulses used to perform single-qubit gates on a $^{40}\text{Ca}^+$ ion qubit. Results show estimation of desired pulse frequencies and pulse times roughly 7 times faster than with standard calibration techniques. Estimates with sample standard deviations of 240 Hz are obtained for the desired pulse frequency, and estimates obtained for the $\pi/2$ -time and π -time show sample standard deviations of 1.56 % and 1.45 %, respectively.

The obtained accuracies are found to be limited by qubit coherence, and possibly also intensity fluctuations of 729 nm laser light used to perform qubit gates. Techniques to further improve accuracies in the presence of such forms of decoherence are investigated and preliminary results are presented.

In chapter 5 a novel adaptive phase estimation protocol is described, and results from simulation and experiment show improved performance compared to the techniques of chapter 4. Results from simulation suggest that improvement in the accuracy of frequency estimation by a factor of 1.69 should be possible, and improvements by a factor of 1.85 are found for the π -time. Experimental results show sample standard deviations of 101 Hz, 0.59 %, and 0.49 % for frequency, $\pi/2$ -time, and π -time, respectively.

Contents

1	Introduction	3
2	Controlling Ion States	5
2.1	The Calcium Ion	5
2.2	The Beryllium Ion	7
2.3	Coherent Control	9
2.3.1	Calcium	9
2.3.2	Beryllium	12
2.4	Single-Qubit Manipulation	14
2.4.1	Quantum Mechanical Description	14
2.4.2	Arbitrary Single-Qubit Gates	16
2.4.3	Calibrating Single-Qubit Gates	16
2.5	Quantum Parameter Estimation	19
3	Magnetic Fields and Complex Regression	21
3.1	The Coherence Problem	21
3.2	Feed-forward Algorithm	21
3.3	Results	25
4	Robust Phase Estimation	28
4.1	The Protocol	28
4.1.1	Fixed Rotation Procedure	28
4.1.2	General Procedure	29
4.1.3	Calibrating Laser Pulses	32
4.2	Preliminary Results	34
5	Adaptive Robust Phase Estimation	41
5.1	The Protocol	41
5.1.1	Fixed Rotation Procedure	41
5.1.2	General Procedure	44
5.1.3	Calibrating Laser Pulses	46
5.2	Simulation	47
5.3	Preliminary Results	50
6	Summary and Outlook	54

A	Improving Accuracy in the Presence of Decoherence	56
A.1	Theory	56
A.2	Experiment	61

Chapter 1

Introduction

In 1982, Richard P. Feynman gave an insightful and entertaining speech about quantum simulation in which he raised some fundamental questions and described some key features on the ability to simulate physical systems [1]. The ideas Feynman presented made an important contribution to the beginnings of research on quantum computing, and since then, the development of both the theoretical underpinnings of quantum computation, as well as the progress in experimental control of quantum systems has been significant. In the beginning, the challenge of achieving the necessary coherent control over quantum systems seemed to be an almost impossible task, and this initially cast doubts as to whether quantum computation was even possible [2]. A key problem was thought to be that errors in a quantum computation, even if rare, would accumulate and overwhelm the computation as its complexity grew. Yet it wasn't long before methods for correcting errors in quantum systems were developed [3–7], and the field of quantum error correction rapidly grew.

New techniques for correcting both coherent and incoherent errors were found, and developments eventually led to error threshold theorems stating targets for the control of individual quantum systems necessary in order to correct errors in a scalable way [8,9]. So it is that, using spin-1/2 systems as *qubits* to form the basic building blocks for a quantum computer, an idea already present in Feynman's speech 35 years ago, we now know it is possible to perform arbitrary quantum operations with only a small set of single and multi-qubit gates [10–14].

At the same time, there has been tremendous progress in experimental control of individual quantum systems with the coming of age of quantum thought experiments as Haroche and Raimond put it in their book [15]. Specific criteria necessary for physical implementation of quantum computation have now been identified [16], and several physical systems meeting these criteria are presently under study.

In ion-trap quantum computing experiments, single-qubit operations can be performed with high accuracy using controlled coherent light interactions, while the work of Mølmer and Sørensen has provided techniques for robustly and accurately coupling multiple ions to achieve high-fidelity multi-qubit gates [17–19].

Although the methods for applying the basic building blocks of quantum computation have been established, it remains an experimental challenge to achieve errors below the bounds set by quantum error correction theorems. One of the key features of these challenges is minimising coherent control errors – errors which can

be seen as the result of mis-calibration of control parameters; they are ultimately due to our lack of knowledge about the system we are trying to control.

Fortunately, the task of learning the parameters of a quantum system is not entirely new; the area of quantum parameter estimation is relevant to many other areas of physics, and has been an active field of research in recent decades [20–23]. Meanwhile, the field of classical control has also seen significant development in recent years with, for example, a revival of the use of neural networks in machine learning to achieve better performance [24].

Successfully realising a universal quantum computer is likely to bring together many of these exciting ideas, and exploring them is the subject of this thesis. In particular, I explore the use of phase estimation techniques, and Bayesian adaptive methods for single-qubit gate calibration, and I present some preliminary experimental results for an ion-trap system using $^{40}\text{Ca}^+$ and $^9\text{Be}^+$ ion qubits.

Chapter 2 briefly reviews the physics of the $^{40}\text{Ca}^+$ and $^9\text{Be}^+$ ions used for the experiments discussed in this thesis, and summarises some relevant details of the experimental system. Then a quantum mechanical description of single-qubit gate operations for a trapped-ion qubit are presented, and finally, a brief review of quantum parameter estimation is also included and relevant bounds on the estimation accuracy are discussed.

Chapter 3 describes a complex linear regression algorithm used to improve the coherence of a $^{40}\text{Ca}^+$ ion qubit in the presence of harmonic noise.

Chapter 4 describes a technique developed by Kimmel *et al.* based on the earlier work of Higgins *et al.* [25, 26] which uses non-adaptive phase estimation for robust calibration of single-qubit gate operations. Experimental results are presented for the calibration of frequency and time of laser pulses used to achieve these gates. Then techniques to improve the accuracy of calibration in the presence of decoherence are also discussed, and some preliminary experimental results are presented.

Chapter 5 describes a novel calibration method combining techniques described in chapter 4, with an adaptive Bayesian Ramsey procedure developed by Andrey Lebedev. The new protocol extends the use of the adaptive procedure for calibration of pulse times as well as frequency, while combining it with methods described by Kimmel *et al.* Results from simulation and experiment demonstrate improved accuracy of estimation for both pulse frequency and pulse time calibration.

Chapter 6 summarises the report and the main results, and discusses possible improvements and generalisations that could be made in the future.

Chapter 2

Controlling Ion States

2.1 The Calcium Ion

Most of the experiments presented in this thesis were performed using a $^{40}\text{Ca}^+$ ion qubit. The ions are trapped using a linear Paul trap with an axial trap frequency of $2\pi \times 2 \text{ MHz}$ ¹. Relevant internal states and transitions of the $^{40}\text{Ca}^+$ ion are shown in figure 2.1. The levels used for the optical qubit are the $|S_{1/2}, m = +1/2\rangle$ ground state and $|D_{5/2}, m = +3/2\rangle$ metastable excited state, where the particular Zeeman sub-levels were chosen to minimise magnetic field sensitivity under some geometric constraints of the system [28].

The qubit levels are coupled by a quadrupole transition at 729 nm with a lifetime on the order of one second [30]. We use a magnetic field of $\simeq 119 \text{ G}$ for which the $D_{5/2}$ levels exhibit Zeeman splittings of 201 MHz. The laser used for coherent manipulation of the qubit is stabilised to a linewidth less than 600 Hz which is well below the axial trap frequency, as required to properly resolve sideband transitions (see section 2.4 below) [28].

Qubit state preparation in the $|S_{1/2}, m = +1/2\rangle$ ground state is achieved by applying a σ -polarised 397 nm laser to couple the $|S_{1/2}, m = -1/2\rangle$ state through a dipole-allowed transition to the $P_{1/2}$ states while simultaneously applying 866 nm and 854 nm lasers to couple the $D_{3/2}$ to the $P_{1/2}$ levels and the $D_{5/2}$ to the $P_{3/2}$ levels, respectively. Since the P levels decay rapidly, this combination of couplings leads to populations of 0.999 in the $|S_{1/2}, m = +1/2\rangle$ ground state after 20 μs [28].

Qubit readout is performed by applying π -polarised 397 nm light while simultaneously applying the 866 nm laser to prevent populating the $D_{3/2}$ levels. For detection times of 300 μs , collection efficiencies in our system lead to average photon counts of ~ 25 when the qubit is in the ground state, and typically fewer than ~ 7 when it is measured in the excited state. This way a measurement outcome in the ground or excited state can be discriminated to better than $1 : 10^5$ [28].

We can also prepare the motional state of the ion by applying Doppler cooling using the 397 nm transition. Following this initial cooling stage by a second one using electromagnetically induced transparency cooling, and finally by applying sideband cooling, the ion can be prepared in the motional ground state [31–33].

¹For a theoretical description of ion trapping see [15] pg. 446, or [27] for a more detailed description; for a description of the trap used in this setup, see [28, 29].

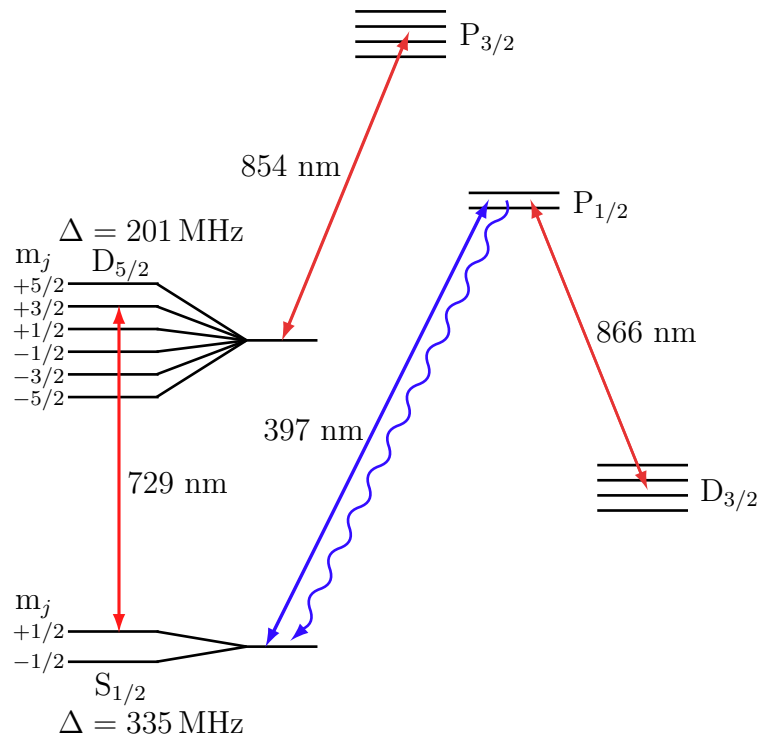


Figure 2.1: $^{40}\text{Ca}^+$ energy levels and transitions. The Zeeman splittings shown are for a magnetic field of ≈ 119 G.

2.2 The Beryllium Ion

Our ion trap has been designed to trap both $^{40}\text{Ca}^+$ and $^9\text{Be}^+$ ions, and some of the experiments in this thesis were also performed using a $^9\text{Be}^+$ ion. The lower mass of $^9\text{Be}^+$ leads to a higher motional frequency which makes it easier to resolve sideband transitions (see section 2.4.1). $^9\text{Be}^+$ has hyperfine structure due to its nuclear spin angular momentum of $\mathbf{I} = 3/2$. The most relevant energy levels are shown in figure 2.2. The hyperfine splittings of the $S_{1/2}$ state are labelled with the total angular momentum quantum number $\mathbf{F} = \mathbf{I} + \mathbf{L} + \mathbf{S}$ and its projection m_F along the z-axis, where $\mathbf{L}$ and $\mathbf{S}$ are the total orbital and total spin angular momentum quantum numbers, respectively. The excited P states are coupled to the ground states by a dipole transition leading to a lifetime of 8.2 ns, and the $P_{1/2}$ and $P_{3/2}$ states are split by 197.2 GHz due to fine structure. All transitions used for state preparation, qubit manipulation, and readout have wavelengths of ~ 313 nm.

To perform qubit operations we choose a pair of levels between the $S_{1/2}$, $\mathbf{F} = 2$ and $\mathbf{F} = 1$ levels. We prepare the ion state by applying σ^+ polarised light at ~ 313 nm near-resonant with transitions from $S_{1/2}$ $|\mathbf{F} = 2, m_F = 1\rangle$ and $|\mathbf{F} = 1, m_F = 1\rangle$ to an excited $P_{1/2}$ state. Since the excited state decays rapidly, the ion is optically pumped into the $|S_{1/2}, \mathbf{F} = 2, m_F = 2\rangle$ state².

Once the ion state is prepared, we coherently couple a given pair of $S_{1/2}$ states via the $P_{1/2}$ states using a stimulated Raman transition. We manipulate the $S_{1/2}$, $|\mathbf{F} = 2, m_F = 2\rangle \leftrightarrow |\mathbf{F} = 1, m_F = 1\rangle$ states, which we refer to as the frequency-dependent qubit (FDQ), and the $S_{1/2}$, $|\mathbf{F} = 2, m_F = 0\rangle \leftrightarrow |\mathbf{F} = 1, m_F = 1\rangle$ states, which we refer to as the frequency-independent qubit (FIQ). We apply a magnetic field of ~ 119.45 G for which the FIQ transition frequency becomes independent of magnetic field to first order. Under this field, the FDQ and FIQ transitions have frequencies of 1018 MHz and 1207 MHz, respectively.

Qubit state readout is performed using a closed cycling transition between the $S_{1/2}$ $|\mathbf{F} = 2, m_F = 2\rangle$ and a $P_{3/2}$ state³. In typical experiments we prepare the $S_{1/2}$ $|\mathbf{F} = 2, m_F = 2\rangle$ state, and then coherently transfer the population to the $|\mathbf{F} = 1, m_F = 1\rangle$ state to perform operations on the FIQ. Since the $|\mathbf{F} = 2, m_F = 2\rangle$ and $|\mathbf{F} = 2, m_F = 0\rangle$ states are close in frequency, population in the latter state can sometimes scatter light during readout. For this reason we typically use a third transition to *shelve* the population from the $|\mathbf{F} = 2, m_F = 0\rangle$ state to the $|\mathbf{F} = 1, m_F = -1\rangle$ state before readout. This transition occurs at 1370 MHz under the magnetic field mentioned above, and we refer to it as the FIQ shelving transition or FIS for short. Thus, after performing qubit manipulation on the FIQ, we coherently transfer the population from the $|\mathbf{F} = 2, m_F = 0\rangle$ state to $|\mathbf{F} = 1, m_F = -1\rangle$, and transfer the population from $|\mathbf{F} = 1, m_F = 1\rangle$ to $|\mathbf{F} = 2, m_F = 2\rangle$, before finally performing readout using the closed cycling transition mentioned above.

To prepare the motional state of the ion we apply doppler cooling using light

²The $P_{1/2}$ states decay to all the $S_{1/2}$ states, but transitions with larger changes in m_F are much less likely. Since the laser light is σ^+ polarised, the two ground states are coupled to the $P_{1/2}$ state with $m_F = 2$.

³We say “closed” here because we excite the $|\mathbf{F} = 3, m_F = 3\rangle$ $P_{3/2}$ state which decays back to the $|\mathbf{F} = 2, m_F = 2\rangle$ $S_{1/2}$ state with high probability.

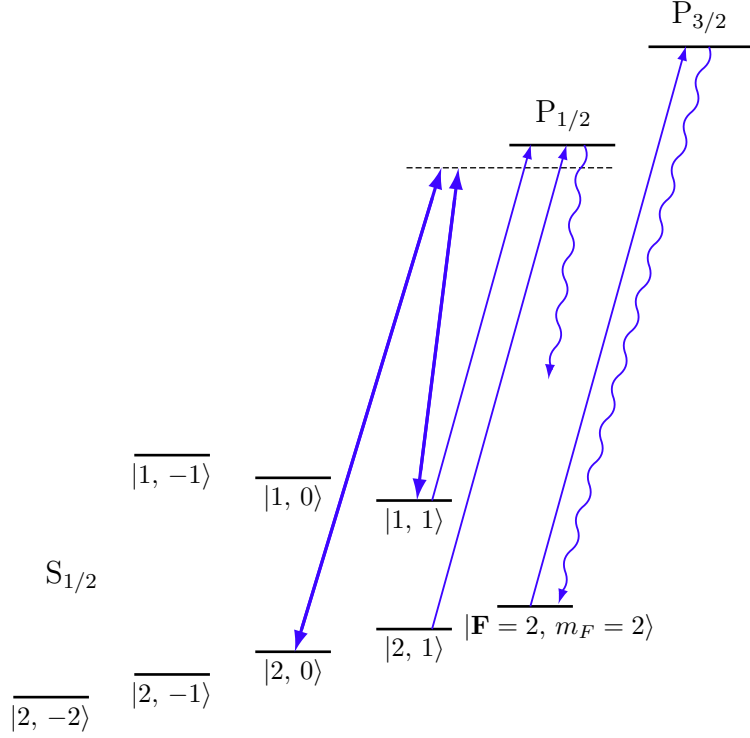


Figure 2.2: $^9\text{Be}^+$ energy levels and transitions. All $S_{1/2}$ to P transitions have a wavelength of ~ 313 nm. Fine structure leads to a splitting of 197.2 GHz between the $P_{1/2}$ and $P_{3/2}$ levels. The leftmost two beams in the figure illustrate the coherent qubit control on the FIQ qubit levels using a stimulated Raman transition, where the two beams are red-detuned from the $P_{1/2}$ levels. The next two beams to the right illustrate σ^+ polarised light near-resonant with transitions from the $S_{1/2}$ $|\mathbf{F} = 2, m_F = 1\rangle$ and $|\mathbf{F} = 1, m_F = 1\rangle$ levels to the $P_{1/2}$ levels, which decay rapidly to prepare the ion in the $S_{1/2}$ $|\mathbf{F} = 2, m_F = 2\rangle$ state. The rightmost beam illustrates the closed cycling transition between $S_{1/2}$ $|\mathbf{F} = 2, m_F = 2\rangle$ and $P_{3/2}$ used for qubit state readout; the same beam is also used for doppler cooling in which case it is red-detuned from the transition.

red-detuned from the closed cycling transition mentioned above. Further cooling can also be applied using sideband cooling. Now you might be wondering how we perform state preparation if there is population in any of the $S_{1/2}$ states with m_F less than one. The answer is that we always perform doppler cooling *before* state preparation. In an initial stage of doppler cooling, the laser power is high enough to broaden the transition such that it becomes resonant with *all* the $S_{1/2}$ states. Since the doppler beam is σ^+ polarised, the combined processes of excitation and stimulated emission tend to *increase* m_F . Thus, after doppler cooling the probability of having population in any of the states with $m_F < 1$ becomes negligible. A detailed description regarding the $^9\text{Be}^+$ ion in our experiments can be found in Hsiang-Yu Lo's PhD thesis [34].

2.3 Coherent Control

2.3.1 Calcium

Accurate control of the laser pulses applied with the 729 nm laser is a necessity for achieving high-fidelity operations on the qubit. Two key devices used to achieve accurate control in our experiments are the acousto-optic modulators (AOMs) used to pulse the 729 nm laser and to control the frequency, phase, and amplitude of the light seen by the ion, and the direct-digital synthesizers (DDSs) used to generate the RF signals that are input to the AOMs.

An AOM is a device that uses the effect of an acoustic wave propagating in a transparent medium to modulate its refractive index. In this way an effective diffraction grating can be generated with parameters that can be tuned by the frequency Ω , phase Φ , and intensity I of the acoustic wave ⁴. The key feature of the AOM is that it leads to Bragg diffraction of the incident beam into different orders of diffracted beams of different frequency and intensity. In the experiments discussed here the positive or negative 1st-order diffracted beam is always used, for which the frequency of the $\pm 1^{\text{st}}$ -order diffraction is given by $\omega_r = \omega \pm \Omega$. Similarly, the phase of the diffracted beam is given by $\phi_r = \phi \pm \Phi$. The ratio R of the output power of the diffracted beam to the total power of the input laser beam has a dependence on the intensity of the acoustic wave that goes as

$$R \propto \left(\sin \sqrt{I} \right)^2. \quad (2.1)$$

Thus, by controlling the frequency, phase, and intensity (proportional to amplitude squared) of the acoustic wave we can tune the corresponding parameters of the light used to manipulate the qubit.

After several stages of amplification and noise cancellation 729 nm light passes through three AOMs on the final stages of its journey to the ion trap as shown in figure 2.3. The first of these, labelled AOM₁ in figure 2.3, is double passed ⁵,

⁴For a short summary of the working principles of an AOM I recommend David Nadlinger's semester thesis report, section 2.2 [35].

⁵Double passing is a common technique in which the diffracted beam is reflected by a mirror back into the AOM where it gets diffracted a second time. After the second pass, the beam of the same order as the first pass will be aligned with the input beam.

and the negative first-order reflected beam is selected. Thus, for incident light of frequency ω and phase ϕ the doubly-diffracted light on its way to the second AOM has frequency and phase

$$\omega_1 = \omega - 2\Omega_1, \quad (2.2)$$

$$\phi_1 = \phi - 2\Phi_1, \quad (2.3)$$

where Ω_1 and Φ_1 are, respectively, the frequency and phase of the acoustic wave in AOM₁. AOM₂ in figure 2.3 is not used in the experiments discussed here. AOM₃ in figure 2.3 is single-passed and the positive first-order diffraction is coupled to an optical fibre and sent to a final focusing stage before making its way to the ion. Therefore, the light that reaches the ion has frequency and phase

$$\omega_3 = \omega_1 + \Omega_3, \quad (2.4)$$

$$\phi_3 = \phi_1 + \Phi_3, \quad (2.5)$$

where Ω_3 and Φ_3 are, respectively, the frequency and phase of the acoustic wave in AOM₃. Both AOM₁ and AOM₃ have a centre acoustic frequency of 200 MHz, and are typically operated within a range of ± 10 MHz around the centre value ⁶.

To set the inputs to the first AOM, we define a DDS pulse instance, `729_master`, in our control system, which has a frequency, phase, amplitude, and time, and we set these parameters in our control computer and send them via an FPGA (Field Programmable Gate Array) to the DDSs. The DDSs then synthesise the radio frequency (RF) pulses used to generate acoustic waves in the AOM. Similarly, we define DDS pulses `729_a` and `729_b` to control AOM₃; in this case we can combine the `729_a` and `729_b` signals to generate two superimposed pulses with different frequencies on the ion ⁷.

To apply single-tone pulses the `729_b` parameters are not needed. Before a pulse is applied we set the `729_a` to a fixed amplitude output at a default frequency which ensures that we are far detuned from qubit resonance. At the same time the `729_master` is set to zero amplitude ⁸. Now to apply a laser pulse to the qubit, we first switch the frequency of the `729_a` to a non-default value which brings the qubit close to resonance when the `729_master` is calibrated to the correct frequency. We then switch on the `729_master` for the desired duration of the pulse, and finally switch the `729_a` back to an off-resonant default value after the `729_master` pulse is over. To tune the frequency, phase and amplitude of the pulse we only change the values for the `729_master`, and keep the phase and amplitude of the `729_a` constant. The `729_a` amplitude is kept at a constant non-zero value in order to stabilise the temperature of AOM₃. This is because the deflection angle of the AOM has a temperature dependence causing the coupling to the output optical fibre after

⁶Using wider ranges will lead to notable losses in power of the output beam.

⁷This is used for example to perform some multi-qubit operations requiring application of two pulses with slightly different frequencies simultaneously.

⁸Although the amplitude is zero in principle, RF-noise can cause small fluctuations in the amplitude which is one reason to keep the `729_a` off-resonant when no pulses are being applied.

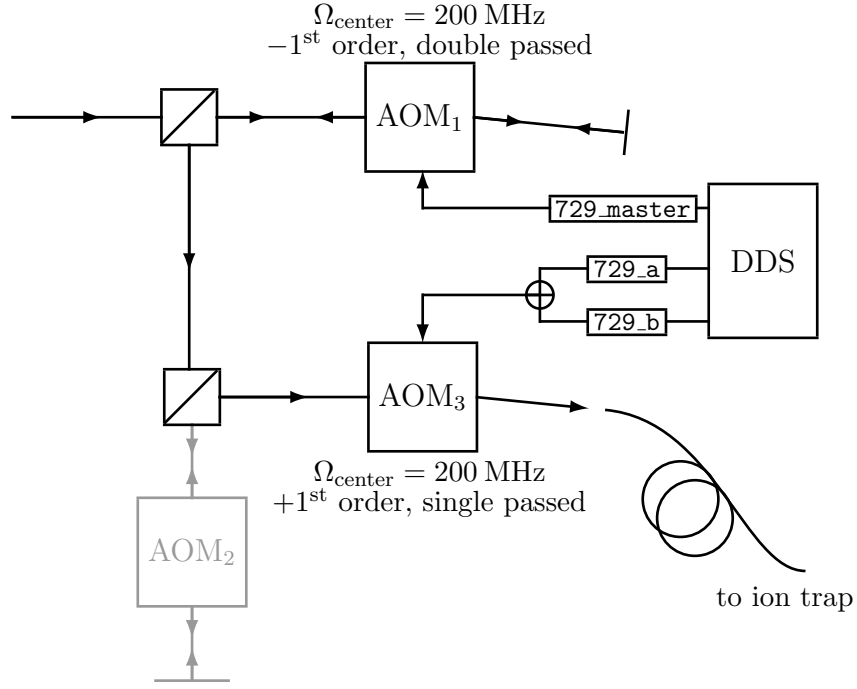


Figure 2.3: Schematic of the final stages of the 729 nm light on its journey to the ion trap. Three AOMs are used to control the frequency, phase, amplitude, and timing of the pulses applied to the ion. In the experiments discussed in this thesis, only AOM₁ and AOM₃ are used.

frequency	$\frac{1}{2^{32}}$ GHz = 0.2328 Hz
phase	$\frac{360}{2^{16}}$ deg = 0.0055 deg
amplitude	$\frac{100}{2^{14}}$ % = 0.0061 %
time	8 ns

Table 2.1: Resolutions of DDS outputs.

AOM₃ to vary (see figure 2.3). For AOM₁ changes in beam alignment depend much more weakly on AOM temperature since, to first order, double-passing the beam cancels any fluctuations in reflection angle.

The accuracies with which we can set the pulse parameters in our experiments are ultimately limited by the digital resolution of the DDS outputs listed in table 2.1. Note that the dependence of the laser amplitude at the ion on the amplitude of the DDS output is not linear, in part due to the nonlinear relation of the reflection coefficient (2.1), and because we use AOM₁ in a double pass configuration⁹. Since AOM₁ is in double pass configuration, the frequency and phase relations given in equations (2.2) and (2.3) effectively halves the resolution for these parameters. This gives a frequency resolution of 0.4657 Hz, and a phase resolution of 0.011 degrees for the light at the ion. In addition to the time resolution listed in table 2.1, pulses in our experiments are limited to a minimum duration of 1.4 μ s due to the time it takes for communication (via serial peripheral interface) between the FPGA and the DDSs.

2.3.2 Beryllium

A conceptual schematic of the AOM setup used to tune the parameters of the 313 nm laser light for coherent control of the $^9\text{Be}^+$ ion qubit states is shown in figure 2.4. Since coherent control of $^9\text{Be}^+$ states is achieved using a stimulated Raman transition, we need to address the ion with two beams whose frequency difference is equal to the qubit transition frequency. Let the wave vectors of the two beams used to address the ion be

$$\begin{aligned}\vec{k}_1 &= \frac{2\pi}{\lambda_1} \hat{n}_1, \\ \vec{k}_2 &= \frac{2\pi}{\lambda_2} \hat{n}_2,\end{aligned}$$

where $\hat{n}_1$ and $\hat{n}_2$ are normal vectors pointing in the direction of the phase velocity for each beam. In order to address sideband transitions (discussed in section 2.4.1) we require the difference of the wave vectors of the two beams, $\vec{k}_1 - \vec{k}_2$, to have a large component along the direction of ion motion. For this reason we apply the two beams at 45° and 135° from the axis of motion of the ion when we want to

⁹In particular, for low intensities the dependence of output laser power to input power is roughly quadratic, which leads to a quartic dependence when the AOM is in double pass configuration.

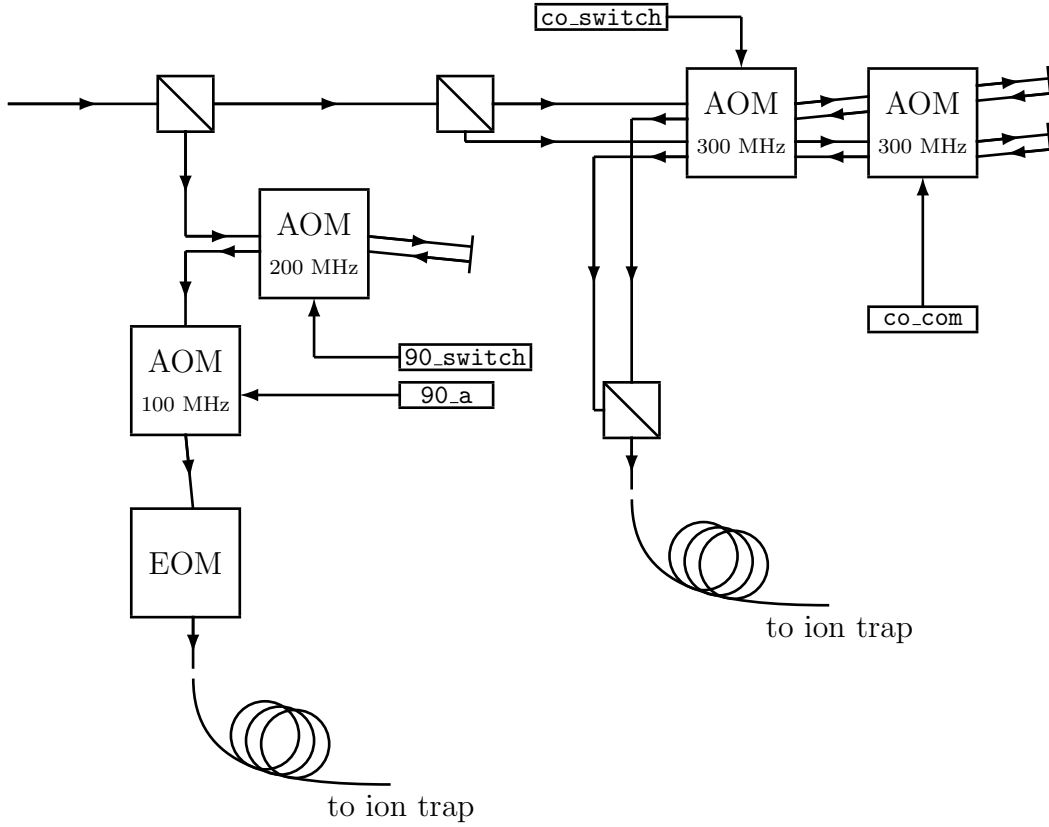


Figure 2.4: Schematic of the final stages of the 313 nm Raman beams used to manipulate ${}^9\text{Be}^+$ qubit states. The first polarising beam splitter (PBS) is used to split the beam into two parts which will address the ion at 45° and 135° from the trap axis. The AOMs are labelled with their respective centre frequencies. The DDS pulses indicate the AOM input side; the beams are deflected in a direction dependent on the order of the reflected beam used.

drive sideband transitions¹⁰. Conversely, when we want to become insensitive to the motion of the ion, we overlap the beams in order to minimise the difference of the wave vectors. As shown in figure 2.4, a first polarising beam splitter (PBS) divides the light into two paths. We use the light from each path to address the ion from the two directions mentioned above.

Lets first consider the trajectory of the transmitted beam; the light passes through a second PBS and both the transmitted and reflected beams are passed through two AOMs¹¹. For our purposes we can consider that only one beam passes through each of these AOMs. They are both in double-pass configuration, and are addressed by the DDS pulses `co_com` and `co_switch`, where “co” refers to the fact that these two beams will be co-propagating when they reach the ion. For one

¹⁰In our experiments we use the motion along the trap axis to perform sideband transitions.

¹¹The fact that both beams pass through both AOMs is not relevant here. This configuration allows phase insensitive transitions for two qubit gates; details regarding this feature can be found in Hsiang Yu Lo’s PhD thesis [34].

of the beams we pick the negative first order reflected beam from the first AOM addressed by the DDS pulse `co_switch`, while the other beam is unaffected. For the second AOM addressed by the DDS pulse `co_com`, the beam reflected by the first AOM is transmitted unaffected, while for the beam that was transmitted by the first AOM we choose the positive first order reflected beam. Both AOMs are operated in double-pass configuration, and the centre frequencies are also shown in figure 2.4. Finally, these two beams are recombined using another PBS and coupled to a fibre leading to the ion trap.

The beam reflected by the first PBS is double-passed through an AOM addressed by the DDS pulse `90_switch`, and then finally is single-passed through an AOM addressed by the DDS pulse `90_a`. Here “90” refers to the fact that these beams will be propagating perpendicular to the “co” beams when they reach the ion.

Similarly to the setup for $^{40}\text{Ca}^+$, we keep the `90_a` on at constant amplitude and only switch the frequency when we send pulses. To perform experiments where we would like to drive sideband transitions, we send pulses to the `co_com` and the `90_switch`. When we would like to be insensitive to ion motion we send pulses to the `co_com` and `co_switch`. To calibrate the frequency and phase of the pulses, we only tune the parameters of the `co_com`.

2.4 Single-Qubit Manipulation

2.4.1 Quantum Mechanical Description

The bare Hamiltonian describing the internal qubit state of a single ion and its axial motional state in the trap is given by

$$H_0 = \frac{\hbar\omega_0}{2}\sigma_z + \hbar\omega_m a^\dagger a, \quad (2.6)$$

where ω_0 is the frequency of the qubit transition, ω_m is the frequency of the axial motion of the ion in the trap, σ_z is the Pauli z operator acting on the ion qubit states, and a and $a^\dagger$ are respectively the annihilation and creation operators acting on the Hilbert space of the ion motion.

The effect of laser light near-resonant with the qubit transition can be described by the interaction Hamiltonian

$$H_{int} = \frac{\Omega}{2} \left(\sigma_+ e^{-i\phi} e^{i(\eta(a+a^\dagger)-\omega t)} + \sigma_- e^{i\phi} e^{-i(\eta(a+a^\dagger)-\omega t)} \right), \quad (2.7)$$

where ω and ϕ are the frequency and phase of the laser, respectively, and σ_+ and σ_- are the raising and lowering operators, respectively, acting on the qubit Hilbert space of the qubit. η is the *Lamb-Dicke parameter* given by

$$\eta = \sqrt{\frac{\hbar k_l^2}{2m\omega_m}} \cos \theta,$$

where k_l is the wavevector of the laser, m is the mass of the ion, and θ is the angle formed between the wavevector of the laser and the axis of motion of the ion for the motional mode we are considering. For most of the experiments of interest for ion trap quantum computation η is small enough that we can approximate the exponential terms in the interaction Hamiltonian (2.7) by their first-order Taylor expansion

$$e^{\pm i\eta(a+a^\dagger)} \approx 1 \pm i\eta a \pm i\eta a^\dagger.$$

With this we can rewrite the interaction Hamiltonian as

$$\begin{aligned} H_{int} = & \frac{\Omega}{2} (\sigma_+ e^{-i\phi} e^{-i\omega t} + \sigma_- e^{i\phi} e^{i\omega t}) + \\ & \frac{i\eta\Omega}{2} (\sigma_+ a e^{-i\phi} e^{-i\omega t} + \sigma_- a^\dagger e^{i\phi} e^{i\omega t}) + \\ & \frac{i\eta\Omega}{2} (\sigma_+ a^\dagger e^{-i\phi} e^{-i\omega t} + \sigma_- a e^{i\phi} e^{i\omega t}). \end{aligned}$$

Going into an interaction picture with respect to the bare Hamiltonian H_0 , we obtain [36]

$$\begin{aligned} H_I = & \frac{\Omega}{2} (\sigma_+ e^{-i\phi} e^{-i(\omega-\omega_0)t} + \sigma_- e^{i\phi} e^{i(\omega-\omega_0)t}) + \\ & \frac{i\eta\Omega}{2} (\sigma_+ a e^{-i\phi} e^{-i(\omega-\omega_0+\omega_m)t} + \sigma_- a^\dagger e^{i\phi} e^{i(\omega-\omega_0+\omega_m)t}) + \\ & \frac{i\eta\Omega}{2} (\sigma_+ a^\dagger e^{-i\phi} e^{-i(\omega-\omega_0-\omega_m)t} + \sigma_- a e^{i\phi} e^{i(\omega-\omega_0-\omega_m)t}). \end{aligned} \quad (2.8)$$

When the linewidth of the laser is much narrower than the motional frequency ω_m , a single mode of laser light will couple only one of the three terms in the interaction Hamiltonian (2.8), as long as $\Omega \ll \omega_m$. The first term will lead to evolution only of the ion's internal qubit state, while the second and third terms will lead to evolution of both the qubit state as well as the motional state of the ion. These three resulting evolutions are referred to as *carrier transitions*, and *red* and *blue sideband transitions*, respectively.

We define the detuning to the qubit transition $\delta = \omega - \omega_0$. Now assuming $\Omega \ll \omega_m$ and choosing ω such that $\delta \ll \omega_m$, we neglect the sideband terms in the Hamiltonian (2.8) and consider only the carrier transition. We can take a similar approach as described above, except that we go into an interaction picture w.r.t. the Hamiltonian

$$H_{0,c} = \frac{\hbar\omega}{2} \sigma_z + \hbar\omega_m a^\dagger a = \frac{\hbar(\omega_0 + \delta)}{2} \sigma_z + \hbar\omega_m a^\dagger a \quad (2.9)$$

to obtain a time-independent Hamiltonian for the carrier transition

$$H_c = \frac{\hbar\delta}{2} \sigma_z + \frac{\Omega}{2} (\sigma_+ e^{-i\phi} + \sigma_- e^{i\phi}).$$

This leads to a time evolution of the qubit's internal state described by the unitary operator

$$U_c(t) = \cos\left(\frac{\Omega_{\text{eff}} t}{2}\right) \mathbb{1} - \frac{i}{\Omega_{\text{eff}}} \sin\left(\frac{\Omega_{\text{eff}} t}{2}\right) [\delta\sigma_z + \Omega(\cos\phi\sigma_x + \sin\phi\sigma_y)] , \quad (2.10)$$

where $\Omega_{\text{eff}} = \sqrt{\Omega^2 + \delta^2}$, $\mathbb{1}$ is the identity operator, and σ_x , σ_y , and σ_z are the Pauli operators acting on the Hilbert space of the qubit.

2.4.2 Arbitrary Single-Qubit Gates

To apply single-qubit gates, we tune the frequency of the laser so that $\delta \approx 0$, in which case the evolution (2.10) simplifies to

$$U_{\text{pulse}}(\Omega t, \phi) = \cos\left(\frac{\Omega t}{2}\right) \mathbb{1} - i \sin\left(\frac{\Omega t}{2}\right) (\cos\phi\sigma_x + \sin\phi\sigma_y) . \quad (2.11)$$

Considering the qubit as a spin-1/2, this evolution describes rotation of the spin around an axis $\cos\phi\hat{x} + \sin\phi\hat{y}$ that lies in the x-y-plane, with frequency Ω , referred to as the *Rabi frequency*. Since two rotations with arbitrary phase along two different axes form a basis for any mapping between single-qubit pure states, we can perform arbitrary unitary operations by applying laser pulses with different phases and times.

2.4.3 Calibrating Single-Qubit Gates

To achieve high-fidelity single-qubit operations experimentally, we need to estimate the ideal parameters to apply to the AOMs in our setup, which determine the detuning δ , and the phase of the qubit rotation Ωt . Unlike the detuning and Rabi frequency, the laser pulse phase ϕ is only a property of the laser, so we can choose the phase of the first pulse in a sequence to correspond to $\phi = 0$, and all later pulse phases are referenced to the first pulse. Therefore the accuracy in controlling the laser phase is only limited by the resolution of the DDS output discussed in section 2.3¹². In contrast, the Rabi frequency Ω depends on the power of the laser at the ion which can drift continuously. Similarly, the detuning will change both due to the Zeeman effect coupling the qubit transition frequency to magnetic field drifts at the ion, and due to drifts in laser frequency.

To calibrate the detuning and qubit rotation phase accurately we can perform operations and measurements on the qubit to estimate the values of the parameters δ and Ωt directly. Then we can adjust the DDS output parameters appropriately to minimise the *error* between the operation performed on the qubit and the desired operation.

To calibrate the frequency of the laser pulse, a standard method is to scan the frequency of the pulse, and collect measurement results for different input

¹²Although the control of the phase is limited by the resolution of the DDS output, since we know the resolution, the error in the applied pulse could in principle be reduced to the inherent DDS error which is much lower.

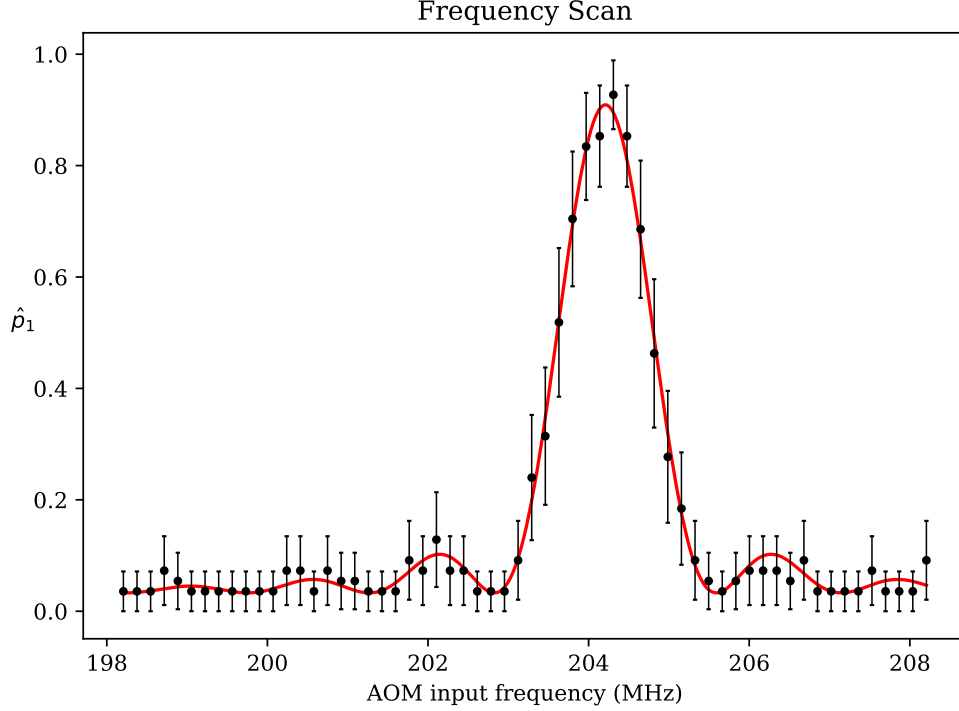


Figure 2.5: Example data obtained by scanning the frequency of a laser pulse over 60 equally-spaced values applied close to resonance with the qubit transition of an ion. This data was obtained from simulation using the evolution operator (2.10) with 50 shots for each data point.

frequencies. We select an input frequency range to scan, and divide the range into N equally-spaced values. For each value of the input frequency we then prepare the qubit in the ground state, and apply the pulse described by the evolution (2.10). Then we measure as described in section 2.1 which collapses the state into the $\{|0\rangle, |1\rangle\}$ basis, where $|0\rangle$ denotes the qubit ground state, and $|1\rangle$ the excited state. We refer to a single realisation of the experiment (preparation, pulse, measurement) as a single *shot* of the experiment. By repeating many shots for each of the N frequency values and plotting the average number of shots per value for which we measure the state $|1\rangle$, we will get something that looks like the data plotted in figure 2.5. It's easy to show from the evolution (2.10) that we expect the probability to measure $|1\rangle$ to take a functional form similar to a Sinc function, which we can use to fit the data as in figure 2.5. This gives us an estimate of the detuning as a function of AOM input frequency. To better approximate the evolution (2.11) we would then set the AOM input frequency to the value given by the peak of the fit.

To calibrate the phase of qubit rotation induced by a laser pulse we take a similar approach as for the frequency calibration only now we scan the time of the pulse instead. In this case we observe Rabi oscillations as plotted in figure 2.6. From the data we can directly estimate the Rabi frequency or the time required to apply a pulse with a given phase. For example, the π -time¹³ can be estimated as one half the period of oscillation of the fit.

¹³i.e. the time t_π such that $\Omega t_\pi = \pi$.

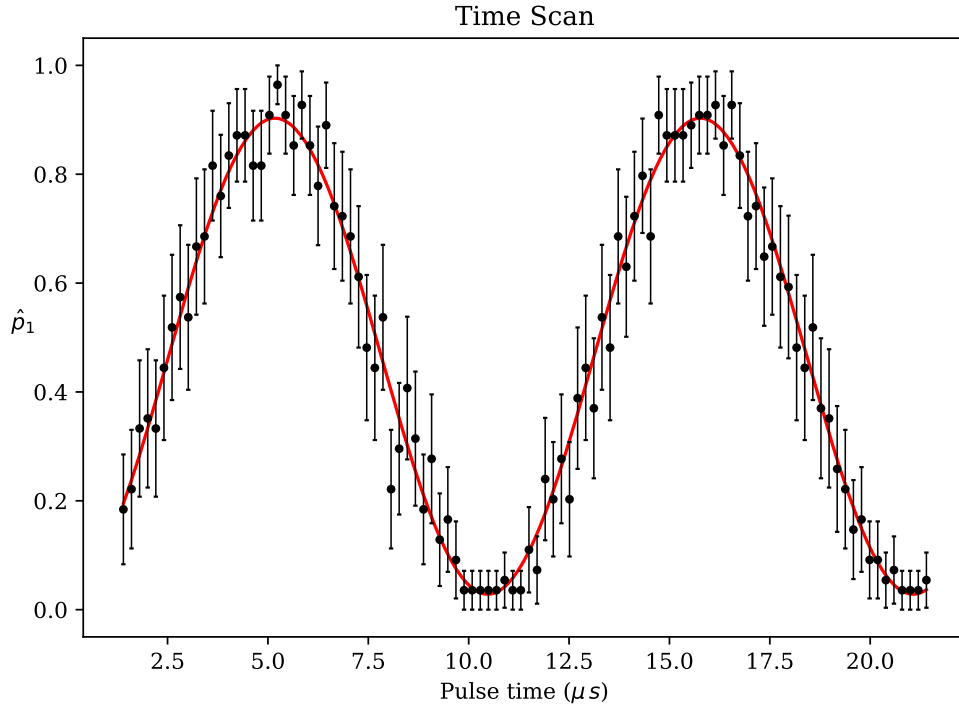


Figure 2.6: Example data obtained by scanning the time of a laser pulse with 100 equally-spaced input pulse time values. This data was obtained from simulation using the evolution operator (2.10) (with a small non-zero detuning) with 50 shots for each data point.

Another standard technique to calibrate the frequency, which can in many cases be more accurate than the method of scanning the frequency described above, is to apply a *Ramsey experiment*. In this case we first prepare the qubit in the ground state and then apply a pulse with $\Omega t = \pi/2$, and which is ideally on resonance with the qubit. We refer to this as a $\pi/2$ -pulse, and it ideally puts the qubit in an initial state

$$|\psi_i\rangle = \frac{1}{\sqrt{2}} (|0\rangle + |1\rangle) , \quad (2.12)$$

which evolves under the bare Hamiltonian (2.6)¹⁴. In an interaction picture w.r.t. (2.9) the free evolution is described by the unitary

$$U_{\text{wait}}(\delta, t) = \cos\left(\frac{\delta t}{2}\right) \mathbb{1} + i \sin\left(\frac{\delta t}{2}\right) \sigma_z . \quad (2.13)$$

Applied to the state (2.12), this leads to the time dependent state

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}} (|0\rangle + e^{-i\delta t}|1\rangle) .$$

By applying another $\pi/2$ -pulse with the same phase as the first pulse, we obtain the final state

$$|\psi_f\rangle = i \sin\left(\frac{\delta t}{2}\right) |0\rangle + \cos\left(\frac{\delta t}{2}\right) |1\rangle .$$

That is, we have converted the phase into a population and observe oscillations in the probability of measuring the state $|0\rangle$ or $|1\rangle$. We can then get an estimate for the frequency in an analogous way as for the time calibration method described above. In practice this technique gives a more accurate estimate of the qubit frequency since a laser pulse leads to AC Stark shifts of the ionic energy levels. This makes the frequency estimation method by scanning pulse frequency dependent on laser power, and generally more sensitive to any errors in laser parameters¹⁵.

2.5 Quantum Parameter Estimation

In section 2.4.3 we saw that calibrating quantum gates in our experiments involves estimating parameters in the expected quantum evolution so that we can tune the system to perform a desired operation. This raises the questions of how well we can estimate parameters for a physical system, and what differences there may be in estimating parameters for a quantum system versus a classical system. Estimating physical parameters is one of the main problems in physics, and these questions

¹⁴The motional part of the Hamiltonian (2.6) is irrelevant here.

¹⁵However, sometimes we would like to calibrate the frequency in the presence of the Stark shifts caused by the laser, or we do not know the $\pi/2$ -time accurately enough to do a Ramsey experiment. In such cases, scanning the pulse frequency is preferred to the Ramsey technique.

have already been considered by numerous physicists in many areas. I will mention here a few of the main findings stemming from the general body of work related to this topic.

The estimation of all parameters associated with Hamiltonian dynamics can be cast into a phase estimation [37]. For example, for single-qubit calibrations we estimate the frequency and time by estimating the phases $\omega_0 t$ and Ωt , respectively.

If we describe both the physical system we are measuring, and the measurement process, classically, then there is no fundamental limit in the precision of phase estimation since in such a description parameters can be measured in principle with arbitrary precision. In contrast, if we describe the measurement process in terms of quantum theory, then the detection process becomes stochastic. In particular, if we still describe the physical system classically, and only the measurement process in terms of quantum mechanics, then we will be limited by Poissonian statistics of the measured variable. For example, if our measurement involves counting N photons then the relative uncertainty of our phase estimate will be limited by

$$\frac{\Delta N}{\langle N \rangle} = \frac{1}{\sqrt{\langle N \rangle}}.$$

This $1/N^{1/2}$ type scaling is known as the *standard quantum limit* (SQL) or *shot noise limit*.

On the other hand, if we also describe the physical system quantum mechanically it is possible to beat the SQL, and this was initially shown by proposals using squeezed states of light [20, 21]. Scaling as good as $1/N$ was eventually shown [22, 23], and has come to be known as *Heisenberg scaling*. Although Heisenberg scaling was thought to perhaps be a fundamental quantum limit, work in nonlinear quantum metrology has shown that scaling better than $1/N$ is possible in some cases [37].

Chapter 4 describes an estimation procedure used to calibrate single-qubit gates that exhibits Heisenberg scaling. But first, chapter 3 describes a method for removing harmonic noise from the transition frequency of a $^{40}\text{Ca}^+$ qubit. This problem could be considered as a classical parameter estimation problem, however, in this case we also perform quantum parameter estimation as part of the procedure (mainly for experimental reasons rather than any fundamental physical reason).

Chapter 3

Magnetic Fields and Complex Regression

3.1 The Coherence Problem

As described in section 2.1, the $^{40}\text{Ca}^+$ ion qubit transition frequency is dependent on the magnetic field at the position of the ion through the Zeeman splittings of the $|S_{1/2}, m = +1/2\rangle$ ground state and the $|D_{5/2}, m = +3/2\rangle$ excited state. Thus, magnetic field fluctuations lead to reduced coherence times when performing qubit gates. Fluctuations in qubit transition frequency can also be seen as a limitation to the accuracy with which we can calibrate the laser pulse frequency to perform qubit gates. A significant source of fluctuations in our experiments is due to mains AC at 50 Hz, and this leads to magnetic field noise at the ion with frequencies of 50 Hz and higher harmonics (i.e. 100 Hz, 150 Hz, ...).

In order to reduce the effects of mains noise on qubit frequency we adjust the current of the coils used to generate the magnetic field at the ion by shifting the set-point of the current controller. Additional magnetic field coils were placed close to the ion to measure the magnetic field [38]; however, for periodic noise oscillating slower than the time resolution with which we can perform experiments on the $^{40}\text{Ca}^+$ qubit, it is possible to get more accurate information about the magnetic field at the ion by measuring fluctuations in the transition frequency directly, at various phases of the noise oscillation. This can be done by repeatedly calibrating the frequency of the 729 nm laser to be resonant with the qubit transition. An example of the qubit frequencies measured over time using Ramsey experiments is shown in figure 3.1 ¹.

3.2 Feed-forward Algorithm

In order to best eliminate the effects of mains noise at the ion, we use the measured oscillations of the qubit resonance frequency to determine the amplitudes of the currents to apply to the coils used to stabilise the field. Since we expect these

¹The detuning estimates for these experiments are determined with an optimised Ramsey procedure using the method described in section 5.1.1, though they could be similarly determined by repeating a standard Ramsey measurement.

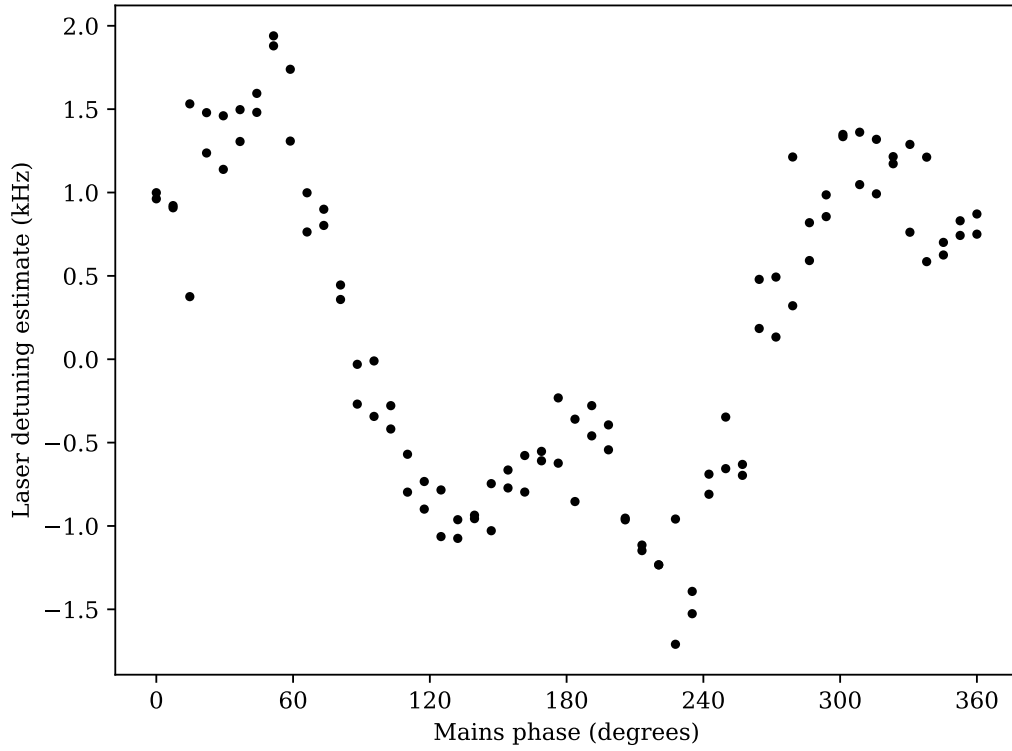


Figure 3.1: Estimated detuning of the 729 nm laser at resonance with the $^{40}\text{Ca}^+$ ion qubit transition. To obtain the x-axis of the plot the experiments are line-triggered according to direct measurement of the mains phase, so the equivalent time for a 1-degree phase shift is $1 \text{ degree} = 20 \text{ ms}/360$. Each point is estimated from 100 shots of a Ramsey experiment with a wait-time of $80 \mu\text{s}$.

oscillations to have frequencies of 50 Hz and higher harmonics, we will need to apply currents with the same frequencies to the coils in order to eliminate the noise, and we fit the measured oscillations using a sum of sinusoidal functions at these frequencies,

$$\sum_{n=1}^N Y_{c,n} \cos(\omega_n t) + Y_{s,n} \sin(\omega_n t), \quad (3.1)$$

where $\omega_n = 2\pi \times 50n$ Hz, the time t is in seconds, and the amplitudes $Y_{c,n}$ and $Y_{s,n}$ are in units of frequency (i.e. since we're fitting the estimated qubit frequencies).

Given the measured amplitudes our goal is to determine the amplitudes of the current that must be applied to the feed-forward coils. One way to think of this problem is to break it into three parts: 1. to quantify the effects of the noise on the qubit transition frequency, 2. to quantify the effect of the feed-forward field on the qubit transition frequency, and 3. to use the information from 1 and 2 to calculate the current to apply to the feed-forward in order to cancel the noise.

To first determine the effects of the noise on the qubit, we can calibrate the laser frequency with the current for the feed-forward coils set to zero. This will yield data like that shown in figure 3.1. Then we can apply current inputs of different amplitudes to the feed-forward coils and repeat the same experiments to measure the qubit frequency, only now we are measuring the effects of both the noise and the feed-forward. In particular, we can apply currents at the same harmonics of 50 Hz as we use to fit the measured data, as mentioned above:

$$\sum_{n=1}^N X_{c,n} \cos(\omega_n t) + X_{s,n} \sin(\omega_n t), \quad (3.2)$$

where in this case the amplitudes $X_{c,n}$ and $X_{s,n}$ are those of the currents applied to the feed-forward coils. If we repeat this several times indexed by $m = 0, 1, 2, \dots, M$, where I will index the case when no currents are applied to the feed-forward with $m = 0$, then we can estimate the effects of the feed-forward on the qubit transition frequency by the measured amplitudes with the noise subtracted, that is $\hat{Y}_{c,n}^m = Y_{c,n}^m - Y_{c,n}^0$ and $\hat{Y}_{s,n}^m = Y_{s,n}^m - Y_{s,n}^0$. Because there is generally a phase shift between an input current oscillation and the measured output, an input amplitude $X_{c,n}$ is coupled to *both* output amplitudes $\hat{Y}_{c,n}^m$ and $\hat{Y}_{s,n}^m$ (likewise for $X_{s,n}$). We would now like to apply linear fits to the input-output data, however this coupling complicates the dependence.

Fortunately, we can make two simplifications to better solve the problem. During the development of our algorithm, we originally tried to isolate the effect of the feed-forward by subtracting the noise we measured when the current was set to zero, as described above. We simplified our task by only considering the input to the feed-forward and the resulting output. Our goal is simply to obtain an estimate for the input amplitudes that minimise the output amplitudes. This simplification also reduced errors in the final current estimates since subtracting the zero feed-forward output from our other outputs was also adding the errors of the two data-sets. The second simplification is that we can describe the problem more conveniently in terms of complex numbers. In our control system we input

the values of the currents for the feed-forward as an amplitude and a phase instead of two amplitudes as described here, but we can determine the amplitudes for the cos and sin terms in equation 3.2 from the trigonometry:

$$X_n \cos(\omega_n t - \phi_n) = X_n \cos(\phi_n) \cos(\omega_n t) + X_n \sin(\phi_n) \sin(\omega_n t).$$

This relation also makes it more apparent that we can represent the amplitudes of the inputs by the complex numbers

$$x_n = X_{c,n} + iX_{s,n} = X_n e^{i\phi_n}.$$

Similarly we can write the amplitudes of the outputs as

$$y_n = Y_{c,n} + iY_{s,n}.$$

Our goal is then to find the complex linear relations between the inputs and the outputs

$$\hat{y}_n = \hat{a}_n x_n + \hat{b}_n,$$

where $\hat{a}_n$ and $\hat{b}_n$ are complex numbers. Writing $\hat{a}_n = A_n e^{i\alpha_n}$ we have $\hat{a}_n x_n = A_n X_n e^{i(\phi_n + \alpha_n)}$. So we see that $\hat{a}_n$ has the effect of scaling the magnitude of the input and changing the phase. This is precisely the effect we expect the input to the feed-forward to have on the ion frequency. $\hat{b}_n$ represents a constant offset in the complex plane; in our case this corresponds to the noise we want to cancel. Our original description to the problem, namely 1. to determine the noise, 2. to determine the effect of the feed-forward, and 3. to cancel the noise, can now be stated simply as 1. determine the values of $\hat{b}_n$, 2. determine the values of $\hat{a}_n$, and 3. solve the equations $\hat{a}_n \hat{x}_n^{\text{opt}} + \hat{b}_n = 0$ for the optimal input estimates. We can determine $\hat{a}_n$ and $\hat{b}_n$ by solving a complex least-squares regression in a similar way to a real least-squares regression. We seek the coefficients

$$\hat{\gamma}_n = \begin{pmatrix} \hat{b}_n \\ \hat{a}_n \end{pmatrix} = \arg \min_{\gamma_n} \sum_{m=1}^M |y_n^m - \hat{y}_n^m|^2,$$

which we can obtain by setting the gradient of this objective function to zero since it is convex. This gives the well known *normal equations*, which we solve by constructing matrices with the inputs and outputs of each experiment as

$$U_n = \begin{pmatrix} 1 & x_n^1 \\ 1 & x_n^2 \\ \vdots & \vdots \\ 1 & x_n^M \end{pmatrix}, \quad V_n = \begin{pmatrix} y_n^1 \\ y_n^2 \\ \vdots \\ y_n^M \end{pmatrix},$$

and calculating the desired coefficients as

$$\hat{\gamma}_n = (U_n^* U_n)^{-1} U_n^* V_n,$$

where a $*$ denotes the conjugate transpose. This is the essentially the same as a real least-squares regression with the transpose replaced by the conjugate transpose. Finally, we set the optimal inputs to the feed-forward as

$$\hat{x}_n^{\text{opt}} = \frac{-\hat{b}_n}{\hat{a}_n}.$$

3.3 Results

To implement this method experimentally, we use Ramsey experiments to estimate the frequency as mentioned above. In a Ramsey experiment, using a longer wait-time leads to a larger phase shift and hence to a more accurate estimate of the frequency. However, we typically apply the complex regression above in several iterations to prevent a Ramsey phase of magnitude greater than π ; if δt in equation (2.13) is outside the range $(-\pi, \pi]$, then we can obtain the wrong frequency estimate corresponding to $\delta' t = \delta t \pm 2\pi l$, $l \in \{1, 2, \dots\}$. This correction could be accounted for, but having much larger ranges in the frequency is also undesirable since the $\pi/2$ -pulses applied in the Ramsey sequence will become inaccurate as they will be further off-resonant.

For this reason it is best to apply the complex regression in several iterations. First we can apply a shorter wait-time to prevent large phase shifts, and apply the regression to the data to obtain estimates for the optimal feed-forward parameters. Then we can apply the regression again using input values in a restricted range around the optimal values found in the previous iteration. Since we stay close to the optimal feed-forward inputs, the detunings oscillate over a much smaller range and we can use much longer wait-times to increase the accuracy of our qubit frequency estimates.

Data obtained for estimated qubit frequency after performing two iterations of complex regression is plotted in figure 3.2. We first performed a quick ($\sim 1 : 20$ mins) and rough correction to cancel the majority of the noise. For this rough correction we estimated the qubit frequency by performing Ramsey experiments with a wait-time of $80 \mu\text{s}$, using 20 points (different values of the mains phase), and 50 shots per point. For the regression we used $n \in \{50 \text{ Hz}, 100 \text{ Hz}, 150 \text{ Hz}\}$ and $M = 4$. For a second iteration we used a much longer wait-time of $400 \mu\text{s}$. We also increased the number of points to 60, and the number of shots per point to

Frequency	Amplitude	Phase (deg)
50 Hz	581.8	222.8
100 Hz	31.5	237.5
150 Hz	333.7	298.7
200 Hz	31.4	201.3
250 Hz	17.5	330.9

Table 3.1: Optimal values obtained for the feed-forward current inputs using complex regression on February 15th, 2017.

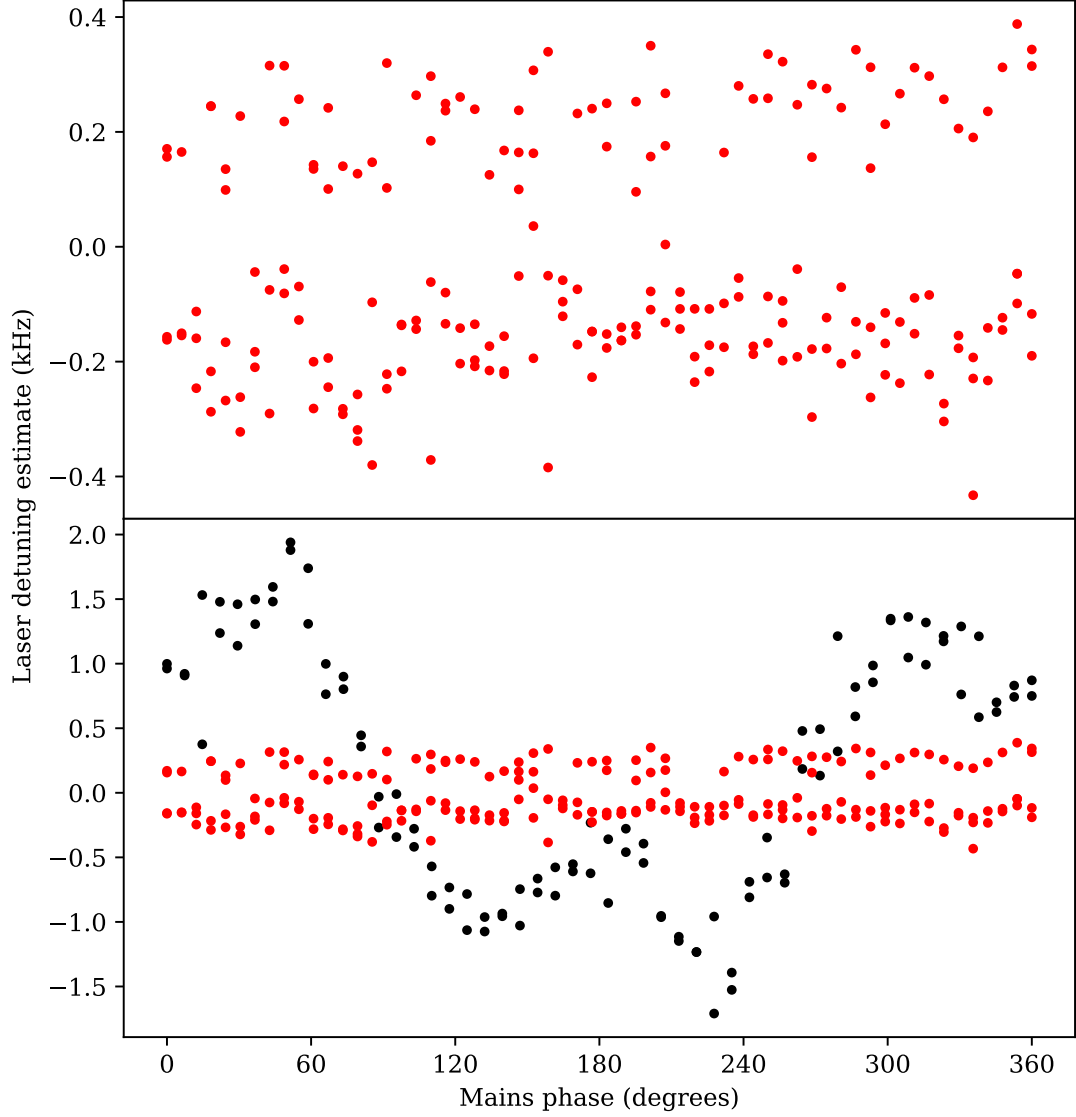


Figure 3.2: Plotted detuning estimates as a function of mains phase after the complex regression method has been applied to estimate the optimal input amplitudes and phases for the feed-forward (red). In the bottom subplot the same data (red) is plotted with the uncorrected frequency estimates (also shown in figure 3.1) (black). The remaining fluctuations are at much higher frequency and most of the fluctuations due to mains noise have been eliminated. The data (red) was obtained using a Ramsey wait-time of $600\text{ }\mu\text{s}$ with 50 shots per point.

100². For the second iteration we also increased the number of frequencies to use for the regression, by taking $n \in \{50 \text{ Hz}, 100 \text{ Hz}, 150 \text{ Hz}, 200 \text{ Hz}, 250 \text{ Hz}\}$ ³, and we again used $M = 4$. For the different values of $m \in \{1, 2, \dots, M\}$, we applied random values for the input amplitudes and phases to the feed-forward. For the first iteration, the amplitudes were chosen uniformly between 0 and 800⁴, and the phases were chosen uniformly between 0 and 360 degrees. For the second iteration, lower ranges were used to prevent large amplitude fluctuations; the centre values for the ranges were taken as the optimal values obtained from the first iteration, or zero (for 200 and 250 Hz). The second iteration lasted $\sim 10 : 20$ mins for a total of $\sim 11 : 40$ mins. The optimal values for the feed-forward obtained from the regression are listed in table 3.1. We see in figure 3.2 that the harmonic oscillations at 50 Hz harmonics have mostly been eliminated (compare the amplitudes with those in figure 3.1).

²For those who run experiments: this was actually 2 scans with 50 shots.

³Special thanks to Vlad for adding in the 200 Hz and 250 Hz components to our control system.

⁴The units for the current are arbitrary.

Chapter 4

Robust Phase Estimation

As a starting point for investigating single-qubit gate calibrations, I experimented with a phase estimation protocol described by Kimmel *et al.* [25] which I will refer to as *robust phase estimation* (RPE). The protocol describes a general method for estimating single-qubit gate parameters and provides an attractive combination of computational simplicity, robustness to errors, and Heisenberg scaling without requiring resources such as entanglement ¹.

4.1 The Protocol

The RPE protocol applies a subroutine to estimate the phase by applying a fixed number of the rotations every shot of the experiment; this subroutine is explained in section 4.1.1. To improve the accuracy of estimation achievable in a given time, which we refer to as the *scaling*, RPE uses this subroutine at each step of the procedure to estimate higher and higher multiples of the phase; i.e. the phase is applied more and more times per shot at later and later steps of the procedure. This is described in detail in section 4.1.2. Section 4.1.3 explains how the protocol can be applied to calibrate the frequency and time of a laser pulse to perform single-qubit gates.

4.1.1 Fixed Rotation Procedure

First I describe the method used in the RPE protocol to estimate the phase of a qubit rotation after a given number of shots $2S$, where each shot involves applying the rotation a *fixed* number of times. Given a single-qubit rotation of phase θ_0 around an arbitrary axis, applied to a pure state whose Bloch vector lies in the plane normal to the axis of rotation, we define two families of measurements. A $|0\rangle$ *measurement* is defined as one along the axis of the initial qubit state vector; a $|+\rangle$ *measurement* is defined as one along the axis perpendicular both to the initial state vector of the qubit, and to the axis of rotation. This is illustrated on the Bloch sphere in figure 4.1 ².

¹Many phase estimation protocols use entanglement as a resource in order to achieve Heisenberg scaling [37].

²Kimmel *et al.* instead define two families of experiments: $|0\rangle$ *experiments* and $|+\rangle$ *experiments* which are specified only by the probabilities of successful outcomes. But we

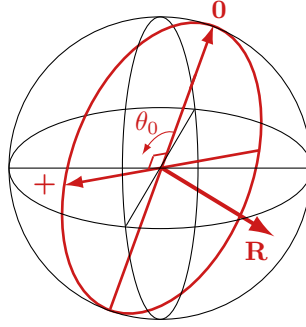


Figure 4.1: The phase θ_0 of an arbitrary qubit rotation, about an axis $\mathbf{R}$, is represented on the Bloch sphere. The measurement axis for the $|0\rangle$ measurements is defined along the initial pure state used by the procedure, and the axis for the $|+\rangle$ measurements lies in the plane of rotation at an angle of $\pi/2$ from the initial state vector.

Now if we prepare the qubit in the initial state, and apply the rotation whose phase we wish to estimate, then the probabilities of the *successful* outcomes for the $|0\rangle$ and $|+\rangle$ measurements will be

$$p_0(\theta_0) = \frac{1 + \cos \theta_0}{2}$$

$$p_+(\theta_0) = \frac{1 + \sin \theta_0}{2}$$

Repeating S shots for each family of measurements, we can get an estimate for the phase

$$\hat{\theta}_0 = \text{atan2}(n_+ - S/2, n_0 - S/2) \in (-\pi, \pi]$$

where n_0 and n_+ are the total number of measured successes for each of the measurement families.

4.1.2 General Procedure

The following theorem is due to Higgins *et al.* [25, 26].

Theorem 1. *Say that we can perform two families of experiments, $|0\rangle$ experiments and $|+\rangle$ experiments, indexed by $k \in \mathbb{Z}$, whose probabilities of success are, respec-*

will see that applying the measurements as I define them here will result in precisely two such families of experiments.

tively,

$$p_0(A, k) = \frac{1 + \cos(kA)}{2},$$

$$p_+(A, k) = \frac{1 + \sin(kA)}{2}.$$

Also assume that performing either of the k^{th} experiments takes time proportional to k . Then, an estimate $\hat{A}$ of A with standard deviation $\sigma(\hat{A})$ can be obtained in time $T = O(1/\sigma(\hat{A}))$ using non-adaptive measurements.

Now I will describe a protocol from Kimmel *et al.*³ which achieves the claim of the theorem⁴. In general, the procedure is performed in K steps. At each step $j \in 1, 2, 3, \dots, K$, a single shot of the experiment consists of preparing the qubit in the initial state, applying the rotation a number of times k_j , and performing a measurement. And at every step j , we perform S_j shots of the experiment for each family of measurements ($|0\rangle$ and $|+\rangle$). Using the fixed rotation procedure described in section 4.1.1 we can get an estimate for the phase $\widehat{k_j A}$ modulo 2π . To determine an estimate for A with Heisenberg scaling, we choose $k_j = 2^{j-1}$, and at each step j of the procedure we calculate an estimate $\hat{A}_j = \widehat{k_j A}/k_j$ restricted to the range $(\hat{A}_{j-1} - \pi/2^{j-1}, \hat{A}_{j-1} + \pi/2^{j-1}]$. This restricts the possible values of $\hat{A}_j$ based on the estimate from the previous step $j-1$ in order to determine the principal range of $\widehat{k_j A}$.

To determine the number of shots S_j at each step j , Kimmel *et al.* show that for an estimate with standard deviation $\sigma(\hat{A})$, and assuming a total time required to obtain the estimate

$$T = 2 \sum_{j=1}^K 2^{j-1} S_j, \quad (4.1)$$

then we can attain the bound

$$\sigma(\hat{A})T < 10.7\pi \quad (4.2)$$

by choosing

$$S_j = \lceil \alpha(K-j) + \beta \rceil, \quad (4.3)$$

with $\alpha = 5/2$ and $\beta = 1/2$. The general procedure described in terms of $|0\rangle$ and $|+\rangle$ measurements, as I have done here, is summarised below.

³This is a slightly modified version of the original protocol by Higgins *et al.*.

⁴I will however not prove the theorem; a full proof can be found in [25].

Procedure Robust Phase Estimation

Input : Unitary evolution $U_r(A)$ which rotates a single-qubit state by an angle A about an axis r .
Pure state $|\psi_0\rangle$ which lies in the plane perpendicular to r .
Total number of steps K .

Output: Estimate $\hat{A}$ of the angle A .

```
for  $j \leftarrow 1$  to  $K$  do
     $k_j \leftarrow 2^{j-1}$ ;
     $S_j \leftarrow \lceil \alpha(K - j) + \beta \rceil$ ;
     $n_0 \leftarrow 0$ ;
     $n_+ \leftarrow 0$ ;
    for  $m \leftarrow 1$  to  $S_j$  do
        for  $e \in \{0, +\}$  do
            prepare the state  $|\psi_0\rangle$ ;
            apply  $U_r(A)$   $k_j$  times;
            perform an  $|e\rangle$  measurement;
            if  $|e\rangle$  measurement is successful then
                 $n_e \leftarrow n_e + 1$ ;
            end
        end
    end
    end
    if  $j == 1$  then
         $\hat{A}_j \leftarrow \text{atan2}(n_+ - S_j/2, n_0 - S_j/2)$ ;
    else
        /* determine the principal range of  $\hat{A}_j$  using  $\hat{A}_{j-1}$  */
         $\hat{A}'_j \leftarrow \text{atan2}(n_+ - S_j/2, n_0 - S_j/2) / k_j$ ;
         $q_1 \leftarrow \text{quotient}(\frac{\hat{A}_{j-1}}{2\pi/k_j})$ ;
         $q_2 \leftarrow q_1 + 1$ ;
        if  $q_1 2\pi/k_j + \hat{A}'_j \in (\hat{A}_{j-1} - \pi/k_j, \hat{A}_{j-1} + \pi/k_j]$  then
             $\hat{A}_j \leftarrow q_1 2\pi/k_j + \hat{A}'_j$ ;
        else if  $q_2 2\pi/k_j + \hat{A}'_j \in (\hat{A}_{j-1} - \pi/k_j, \hat{A}_{j-1} + \pi/k_j]$  then
             $\hat{A}_j \leftarrow q_2 2\pi/k_j + \hat{A}'_j$ ;
        else
             $\hat{A}_j \leftarrow \hat{A}_{j-1} + \pi/k_j$ ;
        end
    end
end
return  $\hat{A}_K$ 
```

Although the RPE protocol provides good scaling in the standard deviation of the estimate as ensured by the bound (4.2), there nevertheless exist more accurate phase estimation techniques [25]. The strength of the RPE protocol is that it is

also robust to errors as quantified by the following theorem due to Kimmel *et al.*.

Theorem 2. *Suppose that we can perform two families of experiments, $|0\rangle$ experiments and $|+\rangle$ experiments, indexed by $k \in \mathbb{Z}^+$, whose probabilities of success are, respectively,*

$$\begin{aligned} p_0(A, k) &= \frac{1 + \cos(kA)}{2} + \delta_0(k), \\ p_+(A, k) &= \frac{1 + \sin(kA)}{2} + \delta_+(k). \end{aligned} \quad (4.4)$$

Also assume that performing either of the k^{th} experiments takes time proportional to k and that

$$\sup_k \{|\delta_0(k)|, |\delta_+(k)|\} < 1/\sqrt{8}. \quad (4.5)$$

Then an estimate $\hat{A}$ of $A \in (-\pi, \pi]$ with standard deviation $\sigma(\hat{A})$ can be obtained in time $T = O(1/\sigma(\hat{A}))$ using non-adaptive experiments. On the other hand, if $|\delta_0(k)|$ and $|\delta_+(k)|$ are less than $1/\sqrt{8}$ for all $k < k^$, then it is possible to obtain an estimate $\hat{A}$ of A with $\sigma(\hat{A}) = O(1/k^*)$ (with no promise on the scaling of the procedure).*

It is possible to show that many of the errors that can occur in typical experiments can be written in the form of additive errors as in (4.4). This includes, for example, errors in preparing the initial state $|\psi_0\rangle$, measurement errors (e.g. if the measurement projectors do not project exactly to the desired states, but to some pure states in the vicinity of the desired ones), and depolarising errors of the form (for a qubit state ρ)

$$\Lambda_\gamma(\rho) = \gamma\rho + (1 - \gamma)\mathbb{1}/2.$$

In particular if we have an experiment with a sequence of k gates, and we write the probability of an outcome without errors as $1/2 + r$, $r \leq 1/2$, then the probability of obtaining the same outcome in the presence of depolarising errors will be $1/2 + r\gamma^k$. This leads to an additive gate error of

$$|r|(1 - \gamma^k) \leq (1 - \gamma^k)/2. \quad (4.6)$$

An analysis quantifying state-preparation and measurement errors as additive errors can be found in Kimmel *et al.* [25]. It is also noted in the same work that, for the case of depolarising errors, a more precise bound than that in (4.6) can be obtained, but they relegate the analysis to a later work (not yet published at the time of writing this report).

4.1.3 Calibrating Laser Pulses

Now I will explain how the RPE protocol can be used to calibrate laser pulses to perform single-qubit gates. I will use some common notation to denote relevant qubit states:

$$\begin{aligned}
|+\rangle &= \frac{1}{\sqrt{2}} (|0\rangle + |1\rangle), \\
|-\rangle &= \frac{1}{\sqrt{2}} (|0\rangle - |1\rangle), \\
|\rightarrow\rangle &= \frac{1}{\sqrt{2}} (|0\rangle + i|1\rangle), \\
|\leftarrow\rangle &= \frac{1}{\sqrt{2}} (|0\rangle - i|1\rangle).
\end{aligned}$$

By choosing the initial state

$$|\psi_0\rangle = |+\rangle,$$

and letting the qubit evolve freely under the Hamiltonian (2.6) as for the Ramsey experiment described in section 2.4.3, then we can choose the $|0\rangle$ measurements described by the projectors $\{|+\rangle\langle+|, |-\rangle\langle-|\}$, where we take a measurement in the $|+\rangle$ state as a successful outcome. Similarly, we can choose $|+\rangle$ measurements described by the projectors $\{|\rightarrow\rangle\langle\rightarrow|, |\leftarrow\rangle\langle\leftarrow|\}$ and take the successful outcome to be a measurement in the state $|\rightarrow\rangle$. Now applying the robust phase estimation procedure, we can get an estimate $\hat{\delta t}_{\text{wait}}$ for δt_{wait} ⁵. To bring the laser pulse closer to resonance with the qubit, we can then update the AOM frequency by assigning a new estimate for the desired input frequency given by

$$\hat{f}_{\text{desired}} = f_{\text{input}} - \frac{\hat{\delta t}_{\text{wait}}}{2\pi t_{\text{wait}}}.$$

To calibrate the phase of rotation induced by a laser pulse, we can take the initial state

$$|\psi_0\rangle = |0\rangle,$$

and choose the $|0\rangle$ measurements $\{|0\rangle\langle 0|, |1\rangle\langle 1|\}$, and $|+\rangle$ measurements $\{|+\rangle\langle+|, |-\rangle\langle-|\}$. In this case applying a laser pulse described by the evolution (2.10) with the phase of the laser set to $\phi = \pi/2$ ⁶, we can use RPE to get an estimate for Ωt .

To update the control parameters for the DDS using an estimate for Ωt , we must consider that there is an unknown time offset in the pulses due to delays caused by AOM rise and fall time, such that

$$t_{\text{pulse}} = t_{\text{input}} - t_{\text{offset}}. \quad (4.7)$$

In this case we can compute an estimate for the Rabi frequency given an estimate $\hat{t}_{\text{offset}}$ for the time offset

⁵Strictly speaking the protocol will return an estimate for $-\delta t_{\text{wait}}$ in this case.

⁶The phase of the laser is set so that the evolution performs a rotation around the y-axis, at least when the detuning is zero. This is to be consistent with the measurements I've chosen to use for the RPE protocol.

$$\hat{\Omega} = \frac{\widehat{\Omega}t_{\text{pulse}}}{t_{\text{input}} - \hat{t}_{\text{offset}}},$$

and estimate the desired input time to the DDS as

$$\hat{t}_{\text{desired}} = t_{\text{input}} + \frac{(\Omega t_{\text{pulse}})_{\text{desired}} - \widehat{\Omega}t_{\text{pulse}}}{\hat{\Omega}}.$$

It is of course also possible to adjust the phase of qubit rotation by changing the Rabi frequency Ω which we can control accurately by changing the amplitude of the pulse as described in section 2.3. In the experiments I will be discussing in later sections, the time was always used to control the phase of the gate.

In the context of calibrating the laser pulses we use to perform single-qubit gates I have described how we can calibrate the time and frequency of the pulse using the RPE protocol. It is also possible to use RPE to estimate the angle between the axes of two rotations; however, since we have very accurate control of the phase of laser pulses in our experiments, we do not *expect* the error to be detectable with the qubit.

4.2 Preliminary Results

The calibration of single-qubit gates by robust phase estimation was implemented with both a $^{40}\text{Ca}^+$ and a $^9\text{Be}^+$ qubit. Most of the results shown here are for $^{40}\text{Ca}^+$, but some preliminary results for $^9\text{Be}^+$ are also mentioned. Results from experiments to collect statistics on the precision of the estimators $\hat{f}_{\text{desired}}$ and $\hat{t}_{\text{desired}}$ are presented, where the time calibrations were performed for the times $t_{\pi/2}$ and t_{π} required to apply qubit rotations of phase $\pi/2$ and π . First, standard methods were used to obtain initial estimates for the pulse parameters, and then the RPE calibration protocol was applied on initial values for f_{input} and t_{input} uniformly distributed over ranges centred on the initial estimates. For the frequency calibration, a somewhat conservative range of 10 kHz was used, while for the time calibrations the ranges were set to 30% of the centre value. For the frequency calibrations a base wait-time for the free evolution of the qubit of $3.0\text{ }\mu\text{s}$ was used.

Results for 50 trials of frequency calibration with K from 6 to 12 are plotted in figure 4.2⁷. Histograms of the estimates are plotted as well as the resulting sample standard deviations which are compared in the figure to the predicted values from the theoretical upper bound (4.2). The upper bounds expected for the frequency estimates were calculated as

$$\sigma_f = \frac{10.7\pi}{T 2\pi t_{\text{wait}}}, \quad (4.8)$$

where T depends on K and is calculated from (4.1).

For frequency calibration we approximately obtain the scaling theoretically predicted when $K < 10$; at $K = 10$ the theoretical bound is reached, and little or no

⁷The mean value of the frequency estimates $\hat{f}_{\text{desired}}$ is not particularly relevant to the current discussion, but for the sake of completeness the input frequency is defined in our control system such that the input to AOM₁ in figure 2.3 is set as $f_{\text{AOM}} = 203.8315 - f_{\text{input}}/2$ MHz.

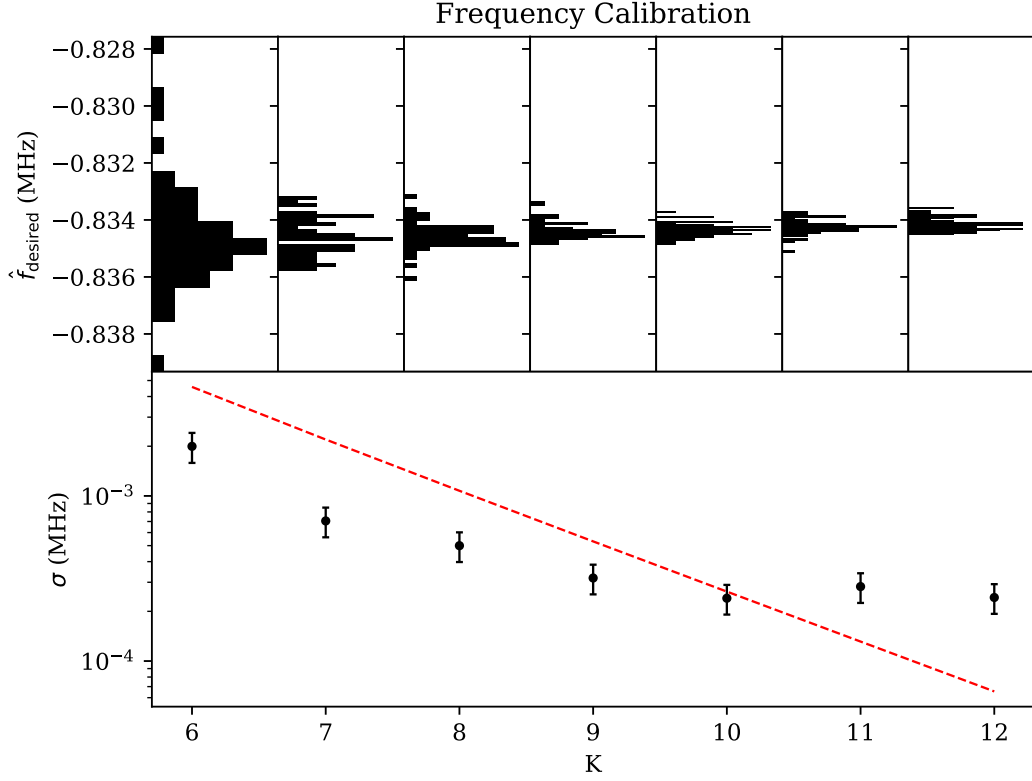


Figure 4.2: RPE frequency estimates obtained for 50 calibration trials with $^{40}\text{Ca}^+$. In the upper plot, histograms of the estimates are shown for different values of the input K to the RPE protocol. In the lower plot, the sample standard deviations (computed from the uniformly minimum variance unbiased (UMVU) estimate for the variance) are plotted (black points) with error bars showing twice the approximate standard error for 50 trials [39]. The red dashed line in the lower plot shows the theoretical upper bound on the standard deviation computed from equation (4.8).

improvement is obtained by increasing K further. This is an expected result for our system due to the finite coherence time of the qubit. As described in section 4.1.2 the RPE protocol progressively applies higher and higher multiples of the phase we wish to estimate. In particular, the largest multiple of the phase applied will be 2^{K-1} . With a wait-time of $3\mu\text{s}$, this leads to a total wait-time for the final steps of the protocol of $768\mu\text{s}$ when $K = 9$, and $1536\mu\text{s}$ when $K = 10$. We expect our coherence to be roughly on the order of 1 ms for the $^{40}\text{Ca}^+$ qubit so that for wait-times above this value, little or no information is obtained since the probabilities in (4.4) will converge to $1/2$. Possible ways of improving the scaling in the presence of decoherence are discussed in appendix A.

The estimates at $K = 9$ have a sample standard deviation of 318 Hz . Using an estimated π -time of $3.9\mu\text{s}$ this corresponds to a sample standard deviation of the normalised detuning, δ/Ω , of 0.04% . If an estimate which is more accurate than $\sim 300\text{ Hz}$ is desirable then without changing any parameters in the RPE protocol, we can already get further improvements in accuracy with a scaling of roughly $\sigma \propto 1/\sqrt{T}$ by averaging over the results of several estimates obtained using $K = 9$ ⁸. The lowest sample standard deviation obtained for the data in figure 4.2 was 240 Hz (0.03% normalised) for $K = 10$.

Results for 250 trials of $\pi/2$ -time calibration for K from 3 to 9 are plotted in figure 4.3. The theoretical upper bound for the standard deviation of the time estimate is approximated from

$$\sigma_t = \frac{10.7\pi}{T\hat{\Omega}} = \frac{10.7\pi\hat{t}_\pi}{T\pi} = \frac{10.7\hat{t}_\pi}{T}. \quad (4.9)$$

The π -time for this calculation was estimated to be roughly $3.5\mu\text{s}$ by taking the average value of the estimates for $K = 9$ and compensating for a time offset of approximately 80 ns (see equation (4.7)).

For the calibration of the $\pi/2$ -time the theoretical bound is reached at about $K = 7$, and little or no improvement is obtained by increasing the value of K further. This is a similar result to that obtained for the frequency calibration, only in this case the absence of improvement above $K = 7$ may also be due to intensity fluctuations of the 729 nm laser. The $\pi/2$ -time results indicate that we do not get additional information when applying pulse sequences with total times greater than $\sim 100\mu\text{s}$. We obtain estimates for the $\pi/2$ -time with a sample standard deviation as low as 28 ns for $K = 8$ from the data plotted in figure 4.3. Previous data taken by David Nadlinger when performing randomised benchmarking experiments on calcium show that longer pulse sequences were performed in the past with lower errors [41]. This suggests that it should be possible to increase the effectiveness of RPE at higher K values. It's also worth noting that David found it likely that the main source of errors in his experiments were due to fast fluctuations of the qubit frequency rather than intensity fluctuations of the laser light. Therefore it's likely that improved performance of RPE for pulse-time calibrations could be achieved by carefully tracking down the sources of error in these experiments.

Results for 50 trials of π -time calibration for K from 3 to 9 are plotted in figure 4.4. For the prediction of the theoretical upper bound, the π -time was

⁸For a sequence of random variables X_i with variance σ_i^2 , the distribution of the sum $\sum_i X_i$ will have variance $\sum_i \sigma_i^2$ [40].

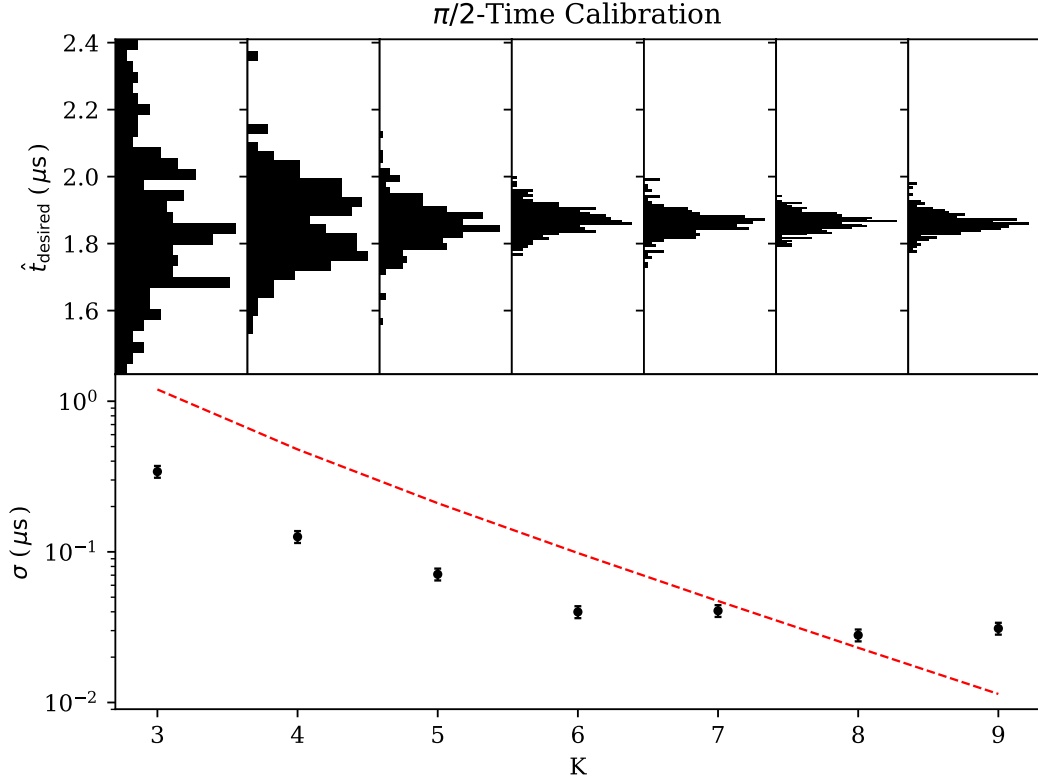


Figure 4.3: RPE $\pi/2$ -time estimates obtained for 250 calibration trials with $^{40}\text{Ca}^+$. In the upper plot, histograms of the estimates are shown for different values of the input K to the RPE protocol. In the lower plot, the sample standard deviations are plotted with error bars showing twice the approximate standard error for 250 trials. The red dashed line in the lower plot shows the theoretical upper bound on the standard deviation computed from equation (4.9) with $\hat{t}_\pi = 3.5 \mu\text{s}$.

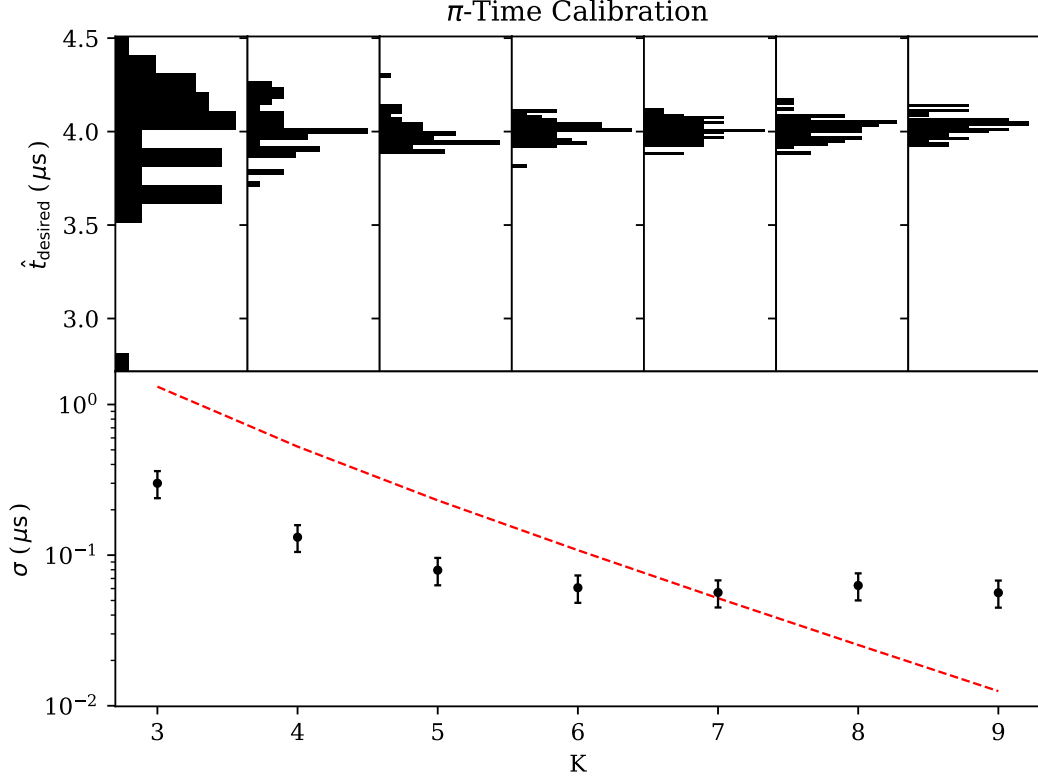


Figure 4.4: RPE π -time estimates obtained for 50 calibration trials with $^{40}\text{Ca}^+$. In the upper plot histograms of the estimates are shown for different values of the input K to the RPE protocol. In the lower plot, the sample standard deviations are plotted with error bars showing twice the approximate standard error for 50 trials. The red dashed line in the lower plot shows the theoretical upper bound on the standard deviation computed from equation (4.9) with $\hat{t}_\pi = 3.9 \mu\text{s}$.

estimated at about $3.9 \mu\text{s}$. In this case the theoretical bound is reached for a slightly lower K value than for the $\pi/2$ -time calibration, although not quite as low as we might expect; since the π -time is roughly double the $\pi/2$ -time (in theory it is exactly double of course) and if we expect the incoherent errors to have roughly the same effect for the same total pulse times, then we also expect to stop seeing improvements in scaling at a K value *lower by one* for the π -time compared to the $\pi/2$ -time (since the pulse is applied 2^{K-1} times at the end of the protocol). However the fact that the bound is apparently reached for a slightly higher K value than expected may only be due to other sources of fluctuation in the data. Indeed, for the π -time, a sample standard deviation as low as 57 ns was obtained for $K = 7$, which agrees surprisingly well with the minimum value obtained from the $\pi/2$ -time data.

To compare with standard calibration methods, we can take the π -time estimates at $K = 5$ which have a sample standard deviation of roughly 1.9% of the centre value for $\hat{t}_\pi$, and $T = 182$ from (4.1). I find similar accuracies by simulating a time scan with a time from 0 to $15.5 \mu\text{s}$ using 25 points with 25 shots per point. This leads to

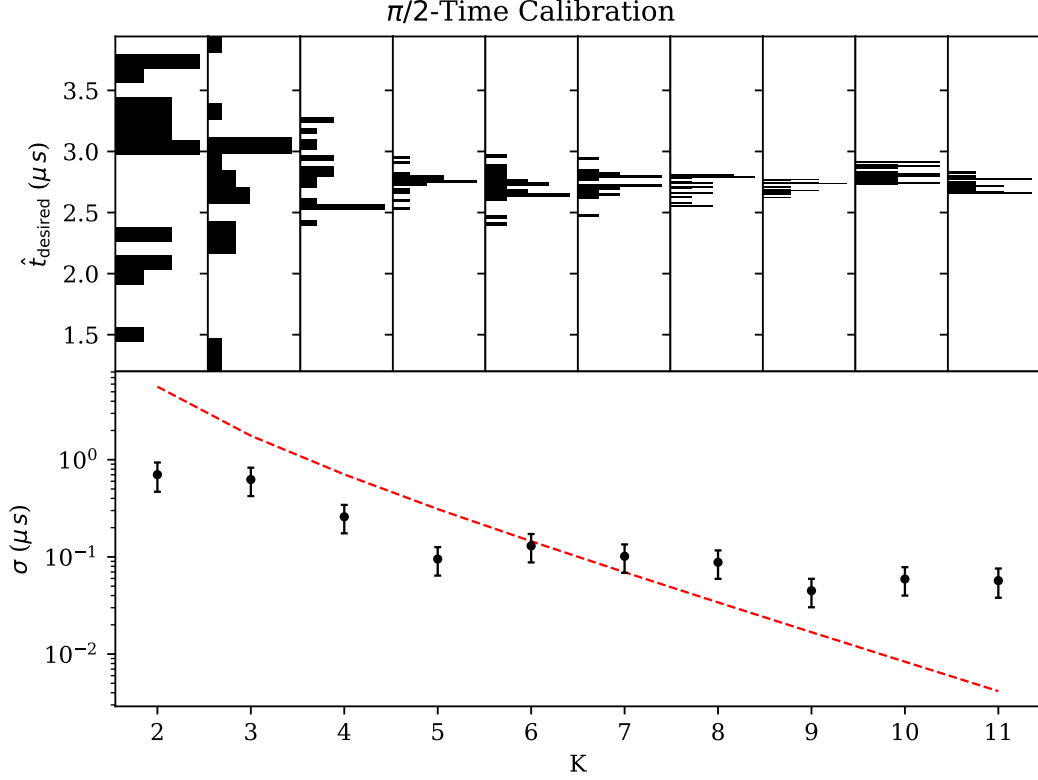


Figure 4.5: RPE $\pi/2$ -time estimates obtained for 20 calibration trials with the $^9\text{Be}^+$ FDQ transition. In the upper plot histograms of the estimates are shown for different values of the input K to the RPE protocol. In the lower plot, the sample standard deviations are plotted with error bars showing twice the approximate standard error for 20 trials. The red dashed line in the lower plot shows the theoretical upper bound on the standard deviation computed from equation (4.9) with $\hat{t}_\pi = 5.28 \mu\text{s}$.

$$T = \sum_{t \in \{2.5, 5, \dots, 15.5\}} \frac{25t}{3.91 \mu\text{s}} \approx 1300,$$

which shows that the RPE estimate is obtained roughly 7 times faster in this case. Also note that here we are comparing experimental data for RPE with an ideal simulation where the detuning is set to zero.

Finally, some preliminary results for $\pi/2$ -time calibration with the $^9\text{Be}^+$ FDQ transition are shown in figure 4.5. The data shows similar behaviour to the $\pi/2$ -time calibration for $^{40}\text{Ca}^+$. Here the theoretical bound is reached at $K = 6$, and a minimum sample standard deviation of 45 ns is reached at $K = 9$. This indicates slightly larger errors than for $^{40}\text{Ca}^+$. It would be interesting to investigate the performance further with $^9\text{Be}^+$, and in particular to see how the frequency calibration performs for the FIQ transition where we expect coherence to be orders of magnitude better than with $^{40}\text{Ca}^+$.

In these experiments the number of shots used throughout the protocol were calculated from equation (4.3) with the optimal values for α and β mentioned in section 4.1.2. This optimisation was performed by calculating the “time” of the pro-

protocol according to equation (4.1). However, this approximation for the experiment time may not be optimal for many of our experiments where there are additional delays due to FPGA performance, communication between hardware, additional cooling stages performed in experimental sequences and so on. Therefore, it may be possible to improve the scaling for specific experiments by re-optimising the number of shots of the protocol based on the exact run times for those experiments.

Chapter 5

Adaptive Robust Phase Estimation

In this chapter I present an adaptive protocol based on the RPE protocol described in chapter 4 and using work by Andrey Lebedev [42] on phase estimation. I will also present some results from simulation that show improved accuracy compared with the non-adaptive RPE protocol. Finally, I will present some preliminary results from implementing the protocol with a $^{40}\text{Ca}^+$ qubit.

5.1 The Protocol

5.1.1 Fixed Rotation Procedure

The adaptive RPE (ARPE) protocol I describe here is similar to the RPE protocol described in chapter 4 only the fixed rotation procedure described in section 4.1.1 is replaced with an adaptive procedure based on the work of Andrey Lebedev.

In the fixed rotation procedure for ARPE we assume that if we wish to measure a phase θ_0 , then we can prepare the qubit in a state for which the probabilities of the outcomes are equal to

$$p_\xi(\alpha, \theta_0) = \frac{1 + \xi \cos(\alpha - \theta_0)}{2}, \quad (5.1)$$

where ξ is the spin of the measured outcome (take $\xi = +1$ for $|0\rangle$, $\xi = -1$ for $|1\rangle$), and α is a phase which we can choose freely. We represent our knowledge of the phase θ_0 at a step $s = 1, 2, \dots$ of the procedure by a Fourier series

$$P_s(\theta) = \sum_{n=-\infty}^{\infty} c_n^{(s)} e^{in\theta}, \quad (5.2)$$

where $c_n = c_{-n}^*$, and by normalisation $c_0 = 1$. We define the phase estimator

$$\hat{\theta}_0 = \arg \int_0^{2\pi} \frac{d\theta}{2\pi} P(\theta) e^{i\theta} = \arg(c_{-1}), \quad (5.3)$$

which gives an estimate for the phase θ_0 , and is justified by the fact that $\hat{\theta}_0$ is the position of the maximum of $P(\theta)$ when $P(\theta)$ is a Gaussian distribution.

At step $s - 1$ of the procedure we update our knowledge of θ_0 using Bayes theorem, which tells us

$$P_s(\theta) = P_s(\theta|\xi_s, \alpha) \propto p(\xi_s|\theta, \alpha)P_{s-1}(\theta|\alpha) = p_{\xi_s}(\alpha, \theta)P_{s-1}(\theta),$$

where in the last step we use the fact that the distribution at the previous step is independent of α . Rewriting this in terms of the Fourier representation of $P(\theta)$, we have

$$\begin{aligned} \sum_{n=-\infty}^{\infty} c_n^{(s)} e^{in\theta} &\propto (1 + \xi_s \cos(\alpha - \theta)) \sum_{n=-\infty}^{\infty} c_n^{(s-1)} e^{in\theta} \\ &= \sum_{n=-\infty}^{\infty} \left[c_n^{(s-1)} + \frac{\xi_s}{2} \left(e^{i\alpha} c_{n+1}^{(s-1)} + e^{-i\alpha} c_{n-1}^{(s-1)} \right) \right] e^{in\theta}. \end{aligned}$$

Therefore we can update the distribution $P(\theta)$ conditioned on the measurement result ξ_s according to the rule

$$\begin{aligned} \tilde{c}_n^s &\leftarrow c_n^{(s-1)} + \frac{\xi_s}{2} \left(e^{i\alpha} c_{n+1}^{(s-1)} + e^{-i\alpha} c_{n-1}^{(s-1)} \right), \\ c_n^s &\leftarrow \frac{\tilde{c}_n^s}{\tilde{c}_0^s}, \end{aligned} \tag{5.4}$$

where the second step is to ensure normalisation holds.

Now we define our goal as choosing the value of α which maximises the expected entropy gain

$$\Delta_s H(\alpha) = \sum_{\xi_s} \pi(\xi_s|\alpha) (H[P_{s-1}(\theta)] - H[P_s(\theta|\alpha, \xi_s)])$$

at step s of the procedure, where H is the Shannon entropy

$$H[P(\theta)] = - \int_0^{2\pi} \frac{d\theta}{2\pi} P(\theta) \ln[P(\theta)/2\pi],$$

and

$$\pi(\xi_s|\alpha) = \int_0^{2\pi} \frac{d\theta}{2\pi} p_{\xi_s}(\alpha, \theta) P_{s-1}(\theta)$$

is the expected probability of measuring outcome ξ_s given a phase α . In other words we wish to set the value α_s at the step s of the procedure as

$$\alpha_s = \alpha : \max_{\alpha} \Delta_s H(\alpha).$$

It can be shown that

$$\Delta_s H(\alpha) = \sum_{m=1}^{\infty} \frac{\Re\{c_{2m}^{(s-1)} e^{2im\alpha}\}}{m(4m^2 - 1)} - \sum_{\xi} p_{\xi}(\alpha) \ln p_{\xi}(\alpha), \tag{5.5}$$

Procedure Adaptive Fixed Rotation Phase Estimation

Input : Experiment with probabilities of measured outcomes given by equation (5.1), where we can freely choose the value of α .

Prior distribution $P_0(\theta)$.

Total number of shots S .

Output: Estimate $\hat{\theta}_0$ of the angle θ_0 .

for $s \leftarrow 1$ **to** S **do**

if $s == 1$ **and** *prior distribution is uniform* **then**

$\alpha_s \leftarrow 0$;

else

 compute the optimal angle α_s by maximising the entropy gain (5.5);

end

 perform the measurement described by (5.1) and update the distribution according to (5.4);

end

return $\hat{\theta}_0$ *given by* (5.3)

up to an additive constant. In the implementation of the procedure described here we maximise this function numerically at each step of the procedure to obtain the optimal value of α . The adaptive fixed rotation procedure is summarised below.

5.1.2 General Procedure

The general procedure of the ARPE protocol is similar to the RPE protocol, only with the fixed rotation procedure described in section 4.1.1 replaced by the one described in section 5.1.1. In addition, the description of our current knowledge of the phase in terms of a probability distribution has the added advantage of allowing us to define prior distributions which reflect our current knowledge, and to directly use information from previous steps of the protocol to initialise the distribution for later steps.

For the protocol described here, I will sometimes set the initial distribution to a Gaussian which is thought to reflect our current knowledge of the phase. We can initialise the distribution (5.2) as a Gaussian with mean μ and variance σ^2 by setting the coefficients to

$$c_n = e^{-2n^2\sigma^2 - in\mu}. \quad (5.6)$$

I will also sometimes want to shift the centre of the distribution; in general we can shift the distribution by an angle β by updating the coefficients in (5.2) according to

$$c_n \leftarrow c_n e^{-in\beta}.$$

Finally, we can increase the standard deviation of the distribution (5.2) by a factor a by setting the values of the coefficients according to

$$c_n \leftarrow |c_n|^{a^2-1} c_n, \quad (5.7)$$

which will be used to modify the distribution between steps of the protocol.

With this I can describe the ARPE protocol. First, we set the initial distribution (5.2) to a uniform distribution $P_0(\theta) = 1$, if we have no prior knowledge of the phase. Otherwise, we set P_0 to a Gaussian distribution according to (5.6) with appropriate values for μ and σ . As for RPE, the procedure is then performed in K steps. At each step $j \in 1, 2, 3, \dots, K$, a single shot of the experiment consists of preparing the qubit in the initial state, applying the rotation a number of times k_j , and performing a measurement such that the probability of outcomes is given by

$$p_\xi^j(\alpha, k_j A) = \frac{1 + \xi \cos(\alpha - k_j A)}{2}. \quad (5.8)$$

We perform S_j shots of the experiment at every step j , and we get an estimate $\widehat{k_j A}$ of the phase $k_j A$ modulo 2π using the adaptive fixed rotation procedure of section 5.1.1 and applying equation (5.3). Now we use the same procedure as for RPE to determine an estimate for A by choosing $k_j = 2^{j-1}$ at each step j of the procedure, and computing an estimate $\hat{A}_j = \widehat{k_j A}/k_j$ restricted to the range $(\hat{A}_{j-1} - \pi/2^{j-1}, \hat{A}_{j-1} + \pi/2^{j-1}]$. In addition, after every step j , we estimate the phase for the next step $j+1$ as $k_{j+1}\hat{A}_j$, and shift the distribution by the difference between the current estimate of the distribution and the estimate for the next step:

$$\begin{aligned}\delta\theta &\leftarrow k_{j+1}\hat{A}_j - \arg(c_{-1}), \\ c_n &\leftarrow c_n e^{-in\delta\theta}.\end{aligned}$$

Finally, we spread out the distribution by a factor a according to equation (5.7) after every step j . The general procedure for ARPE is summarised below.

Procedure Adaptive Robust Phase Estimation

Input : Experiment with probabilities of measured outcomes given by equation (5.8).
Prior distribution $P_0(\theta)$.
Factor a by which to spread the distribution after every step j , according to equation (5.7).
Total number of steps K .

Output: Estimate $\hat{A}$ of the angle A .

```

for  $j \leftarrow 1$  to  $K$  do
     $k_j \leftarrow 2^{j-1}$ ;
    assign number of shots  $S_j$ ;
    for  $s \leftarrow 1$  to  $S_j$  do
        | apply adaptive fixed rotation procedure;
    end
    if  $j == 1$  then
        |  $\hat{A}_j \leftarrow \arg(c_{-1})$ ;
    else
        | /* determine the principal range of  $\hat{A}_j$  using  $\hat{A}_{j-1}$  */
         $\hat{A}'_j \leftarrow \arg(c_{-1})/k_j$ ;
         $q_1 \leftarrow \text{quotient}(\frac{\hat{A}_{j-1}}{2\pi/k_j})$ ;
         $q_2 \leftarrow q_1 + 1$ ;
        if  $q_1 2\pi/k_j + \hat{A}'_j \in (\hat{A}_{j-1} - \pi/k_j, \hat{A}_{j-1} + \pi/k_j]$  then
            |  $\hat{A}_j \leftarrow q_1 2\pi/k_j + \hat{A}'_j$ ;
        else if  $q_2 2\pi/k_j + \hat{A}'_j \in (\hat{A}_{j-1} - \pi/k_j, \hat{A}_{j-1} + \pi/k_j]$  then
            |  $\hat{A}_j \leftarrow q_2 2\pi/k_j + \hat{A}'_j$ ;
        else
            |  $\hat{A}_j \leftarrow \hat{A}_{j-1} + \pi/k_j$ ;
        end
    end
     $\delta\theta \leftarrow k_{j+1}\hat{A}_j - \arg(c_{-1})$ ;
     $c_n \leftarrow c_n e^{-in\delta\theta}$ ;
     $c_n \leftarrow |c_n|^{a^2-1} c_n$ ;
end
return  $\hat{A}_K$ 

```

5.1.3 Calibrating Laser Pulses

In order to use the ARPE protocol described in section 5.1.2 to calibrate single-qubit operations, we need to apply the phase we wish to estimate in such a way as to obtain the outcome probabilities given by equation (5.8). In addition, we need to be able to accurately control the parameter α in (5.8); in all the calibration methods described here this is done by changing the phase of the laser, since we have good control over this parameter in our experiments (see sections 2.3 and 2.4.3). In the following, I will use the notation for a laser pulse as in equation (2.11), and for the evolution during a Ramsey wait as in equation (2.13).

For the frequency calibration, we start with a qubit prepared in the ground state $|0\rangle$, and apply a Ramsey sequence as ¹

$$\begin{aligned}
& U_{\text{pulse}}(\pi/2, \pi/2) U_{\text{wait}}(\delta, k_j t_{\text{wait}}) U_{\text{pulse}}(\pi/2, -\pi/2 - \alpha) |0\rangle \\
&= U_{\text{pulse}}(\pi/2, \pi/2) U_{\text{wait}}(\delta, k_j t_{\text{wait}}) \frac{1}{\sqrt{2}} (|0\rangle - e^{-i\alpha} |1\rangle) \\
&= U_{\text{pulse}}(\pi/2, \pi/2) \frac{1}{\sqrt{2}} (|0\rangle - e^{-i(\alpha + k_j \delta t_{\text{wait}})} |1\rangle) \\
&= \cos\left(\frac{\alpha + k_j \delta t_{\text{wait}}}{2}\right) |0\rangle + i \sin\left(\frac{\alpha + k_j \delta t_{\text{wait}}}{2}\right) |1\rangle,
\end{aligned}$$

for which the outcome probabilities for the resulting state are

$$p_\xi = \frac{1 + \xi \cos(\alpha + k_j \delta t_{\text{wait}})}{2}.$$

This is the same as in equation (5.8) with the replacement $A \rightarrow -\delta t_{\text{wait}}$.

To calibrate the pulse time we apply three pulses to the qubit prepared in the ground state:

$$\begin{aligned}
& U_{\text{pulse}}(\pi/2, \pi/2) U_{\text{pulse}}(\pi/2, -\pi/2 - \alpha) U_{\text{pulse}}(\Omega k_j t_{\text{pulse}}, -\alpha) |0\rangle \\
&= U_{\text{pulse}}(\pi/2, \pi/2) U_{\text{pulse}}(\pi/2, -\pi/2 - \alpha) \cdot \\
&\quad \cdot \left(\cos\left(\frac{k_j \Omega t_{\text{pulse}}}{2}\right) |0\rangle - i e^{-i\alpha} \sin\left(\frac{k_j \Omega t_{\text{pulse}}}{2}\right) |1\rangle \right) \\
&= U_{\text{pulse}}(\pi/2, \pi/2) \frac{1}{\sqrt{2}} (|0\rangle - e^{-i(\alpha - k_j \Omega t_{\text{pulse}})} |1\rangle) \\
&= \cos\left(\frac{\alpha - k_j \Omega t_{\text{pulse}}}{2}\right) |0\rangle + i \sin\left(\frac{\alpha - k_j \Omega t_{\text{pulse}}}{2}\right) |1\rangle.
\end{aligned}$$

In this case the outcome probabilities for the resulting state are

$$p_\xi = \frac{1 + \xi \cos(\alpha - k_j \Omega t_{\text{pulse}})}{2},$$

which is equivalent to equation (5.8) with the replacement $A \rightarrow \Omega t_{\text{pulse}}$. Given estimates $\widehat{\delta t_{\text{wait}}}$ and $\widehat{\Omega t_{\text{pulse}}}$, we can compute estimates for the desired pulse frequency, $\widehat{f_{\text{desired}}}$, and time $\widehat{t_{\text{desired}}}$, as described in section 4.1.3.

¹The resulting states are written up to a global phase.

5.2 Simulation

In this section I present results from simulating the ARPE protocol and I compare with simulation results for the RPE protocol as an initial characterisation of performance. Note that I have not specified the number of shots S_j to be performed at each step j of the protocol in section 5.1.2; this is because the analysis by Kimmel *et al.* to obtain S_j according to equation (4.3) is no longer valid for ARPE. Nevertheless, as an initial test, simulations were performed using S_j given by (4.3) ².

Results for $\hat{\sigma}$ and the ME (as defined in appendix A.1) obtained by simulating frequency calibration are compared with RPE in figure 5.1 for 1000 calibration trials with $K = 3, \dots, 10$. The initial distribution $P_0(\theta)$ is set to a uniform distribution, and after each step j the distribution is spread out by a factor of $a = 2.3$ in (5.7). The distribution $P(\theta)$ was represented using equation 5.2 with 100 coefficients $\{c_n; n = 0, \dots, 100\}$. For each trial of the calibration, the input pulse frequency was chosen from a uniform distribution with a range of 100 kHz, centred on the true value used for the simulation.

A similar simulation was performed for π -time calibration, and the results are shown in figure 5.2. In this case, the initial value for the π -time was chosen from a uniform distribution centred on the true value used by the simulation with a range of 30 % of the true value. In addition, the initial distribution $P_0(\theta)$ was set to a Gaussian according to equation (5.6) with $\mu = \pi$, and $\sigma = \pi/8$, and a value of $a = 2.0$ in (5.7) was used to spread the distribution after each step of the protocol.

For both frequency and time calibrations $\hat{\sigma}$ and the ME of the estimates from ARPE are lower than those for RPE. Fitting a functional form $\propto 1/T$ to $\hat{\sigma}$ and the ME of the estimates for both protocols, I obtain an approximation for the relative scaling of the protocols. For the frequency I find

$$\begin{aligned} (\sigma(\hat{A})T)_{\text{ARPE}} &\approx 0.53 (\sigma(\hat{A})T)_{\text{RPE}} , \\ (\text{ME} \cdot T)_{\text{ARPE}} &\approx 0.59 (\text{ME} \cdot T)_{\text{RPE}} , \end{aligned} \quad (5.9)$$

from the data shown in figure 5.1. And for the π -time I find

$$\begin{aligned} (\sigma(\hat{A})T)_{\text{ARPE}} &\approx 0.52 (\sigma(\hat{A})T)_{\text{RPE}} . \\ (\text{ME} T)_{\text{ARPE}} &\approx 0.54 (\text{ME} T)_{\text{RPE}} . \end{aligned} \quad (5.10)$$

from the data shown in figure 5.2.

However, we have not optimised the number of shots S_j at step j to obtain optimal scaling with this protocol. Therefore, it may be possible to obtain further improvements in the scaling.

²In fact, I set S_j to *twice* the value given by (4.3) to obtain the same number of shots for each j and the same total “time” as the RPE protocol for a given value of K .

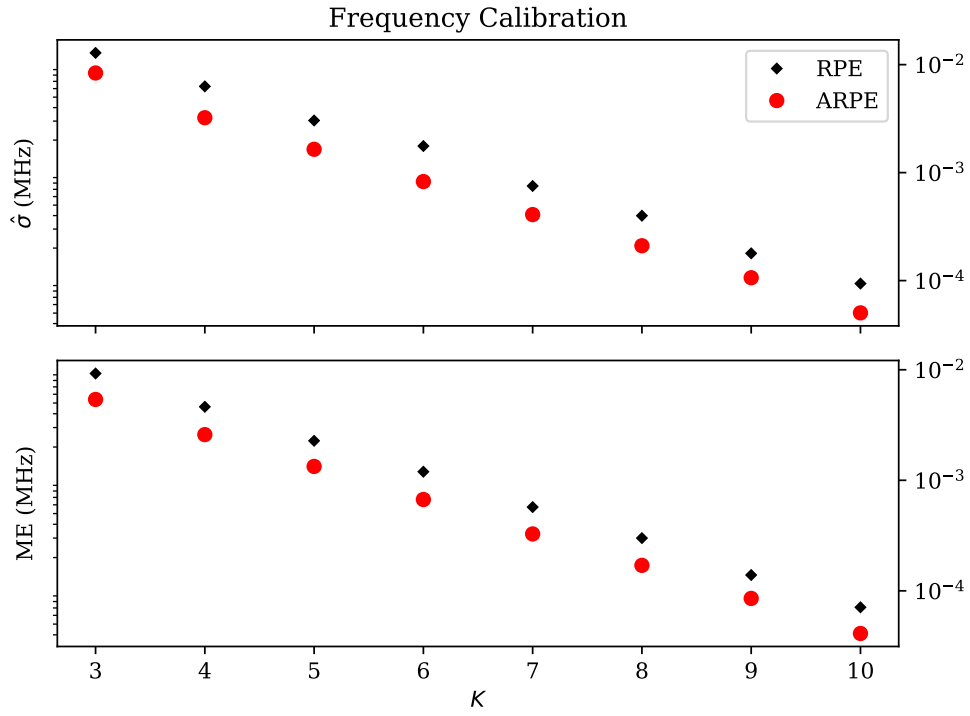


Figure 5.1: Comparison of the sample standard deviations and the mean errors from a simulated sample of 1000 frequency calibration trials for the ARPE and RPE protocols. The distribution for the phase is initially uniform, and a value of $a = 2.3$ is used in equation (5.7) to spread the distribution after each step of the ARPE protocol.

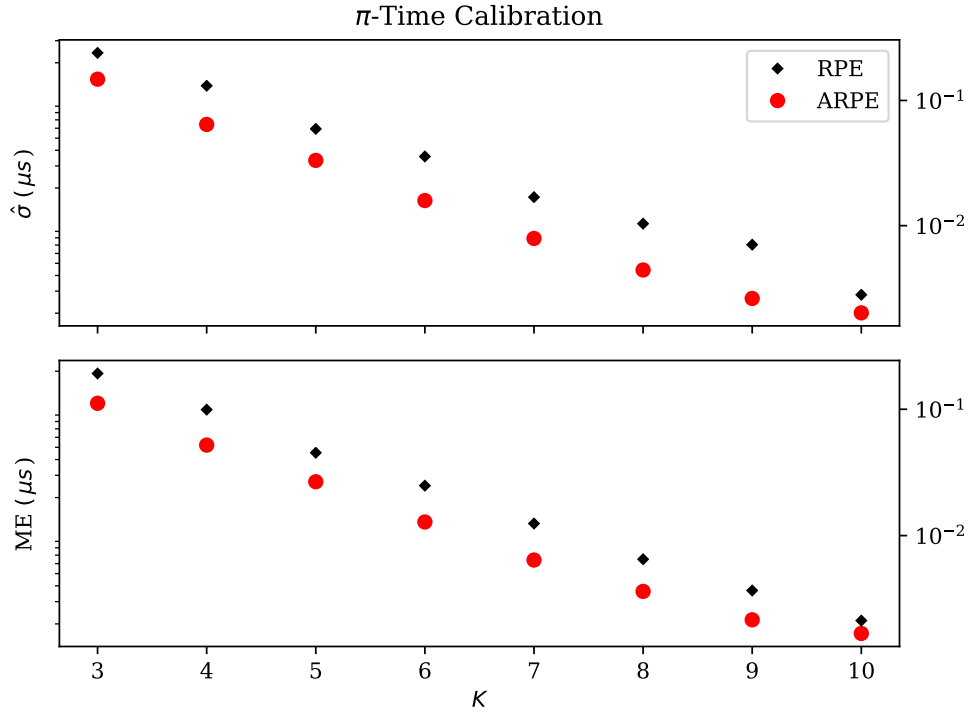


Figure 5.2: Comparison of the sample standard deviations and the mean errors from a simulated sample of 1000 π -time calibration trials for the ARPE and RPE protocols. The distribution for the phase is initially set to a Gaussian as in (5.6) with $\mu = \pi$, $\sigma = \pi/8$, and a value of $a = 2.0$ is used in equation (5.7) to spread the distribution after each step of the ARPE protocol.

5.3 Preliminary Results

The ARPE protocol was also tested experimentally for the case that the number of shots is the same as for the RPE protocol. Figure 5.3 shows the histogrammed estimates from 50 trials of frequency calibration for K values from 3 to 12. Also included in the figure is a plot showing the sample standard deviation of the estimates compared with some data for frequency calibration using RPE. The improvement seen for the ARPE estimates is probably not as good as it looks in this case since the data obtained for the RPE protocol used a shorter wait-time of $3.0 \mu\text{s}$ compared to a wait-time of $7.2 \mu\text{s}$ used for the ARPE data. However the minimum sample standard deviation attained should be approximately independent of the wait-time used by the protocol. Therefore, the lower standard deviations seen at high K values give a less ambiguous indication of improved performance. A sample standard deviation of 120 Hz ($\Delta\delta/\Omega = 0.012\%$) is observed at $K = 10$, and a minimum of 102 Hz ($\Delta\delta/\Omega = 0.010\%$) is observed at $K = 12$.

Figures 5.4, and 5.5 also show histogrammed estimates for 50 trials of $\pi/2$ -time and π -time calibration, respectively, as well as plotted sample standard deviations. We see that the ARPE protocol consistently delivers estimates with lower standard deviation than the RPE protocol. For the $\pi/2$ -time, a minimum sample standard deviation of 9 ns is observed at $K = 8$, compared to 17 ns obtained with RPE at the same K -value. For the π -time, a minimum sample standard deviation of 16 ns is observed at $K = 8$, compared to 25 ns obtained with RPE at the same K -value.

We can also calculate the *process fidelity* of quantum operations from the values for the standard deviations of our estimates³. The process fidelity of an applied operation G can be expressed as

$$F(G, U) = \frac{1}{n(n+1)} \left(\text{tr} \left(\sum_k M_k^\dagger M_k \right) + \sum_k |\text{tr}(M_k)|^2 \right)$$

where U is the ideal unitary target operation, n is the dimension of the Hilbert space, and $M_k = U^\dagger G_k$ is calculated from the target unitary and the Kraus operators $\{G_k\}_k$ of G . If we assume our applied gate is given by (2.10) with finite detuning δ and angular rotation error ϵ due to imperfect pulse-time calibration, the process fidelity for a π -pulse can be expressed as

$$F(G, U(\pi)) = \frac{1}{3} \left(\frac{\Omega}{\Omega_{\text{eff}}} \right)^2 \left(2 + \left(\frac{\delta}{\Omega} \right)^2 - \cos \left(\frac{\Omega_{\text{eff}}}{\Omega} (\epsilon + \pi) \right) \right).$$

$\delta = 102\text{Hz}$ corresponds to $\delta/\Omega \approx 1 \times 10^{-4}$, and a 16 ns π -time sample standard deviation to $\epsilon/\pi = 4.9 \times 10^{-3}$. Substituting these values leads to a process fidelity for a π -pulse of 0.999996 or equivalently an infidelity of 4×10^{-6} . Previous analysis by David Nadlinger using randomised benchmarking found *measured* process infidelities were larger (1.69×10^{-4}) mainly due to fast frequency fluctuations of the qubit. The above calculation for the process fidelity only includes static calibration

³The process fidelity can be seen as the average over all possible initial pure states of the state fidelity between a target and applied quantum operation. See for example [41, 43].

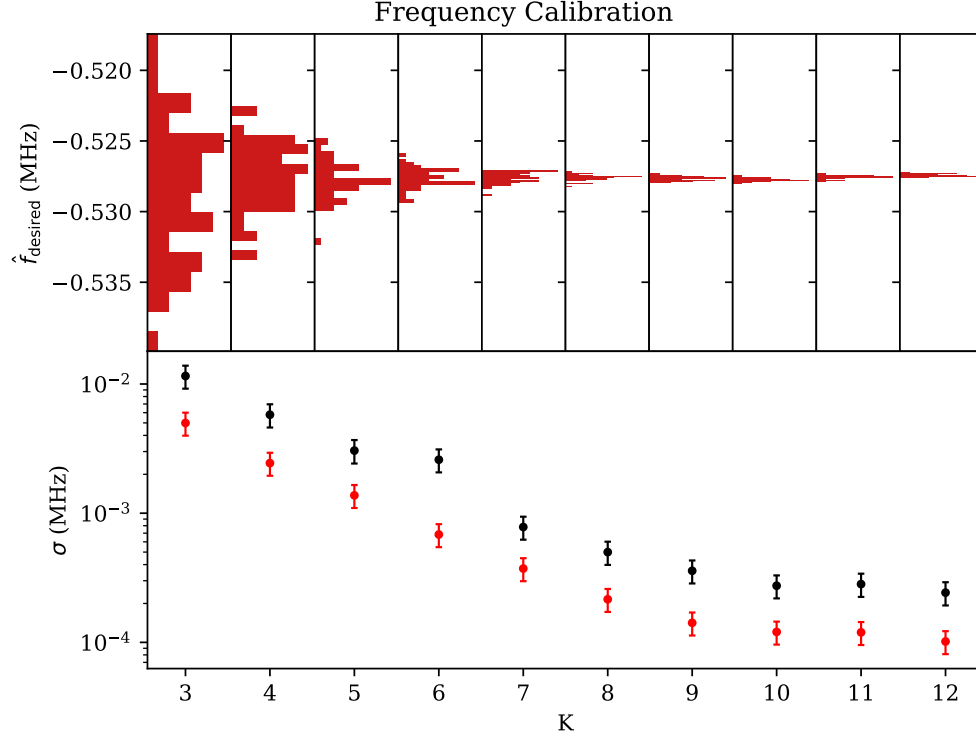


Figure 5.3: Experimental data comparing frequency calibration using ARPE versus RPE. The ARPE data was taken with an input uniform distribution of frequency estimates with a range of 100 kHz about a centre value obtained using standard calibration techniques before the trials were performed. A uniform initial distribution was used for $P(\theta)$, and a factor of $a = 2.5$ was used in equation (5.7). The RPE data is not from the same day, and was taken with a smaller range of 10 kHz. The wait-times used for the protocols were $7.2 \mu s$ and $3.0 \mu s$ for ARPE and RPE, respectively.

errors and not other errors due to qubit frequency fluctuations or laser intensity noise. Based on the number of pulses performed before overwhelming the error bounds for RPE, the actual gate infidelities were likely larger here than previously measured using randomised benchmarking. This suggests better accuracies with RPE and ARPE may be achievable in our system.

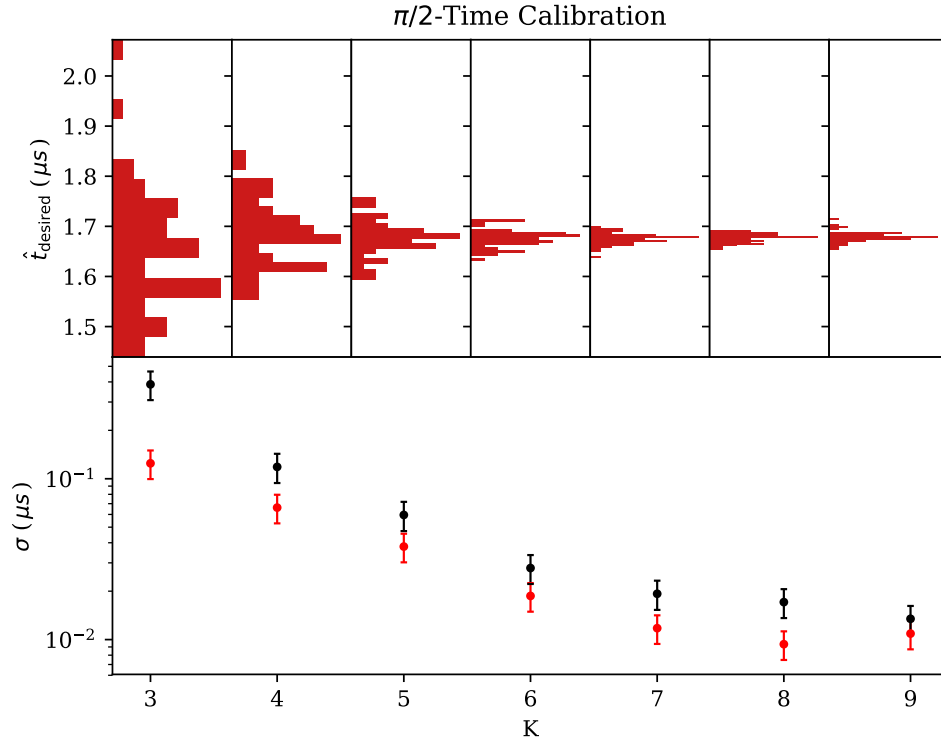


Figure 5.4: Experimental data comparing $\pi/2$ -time calibration for ARPE versus RPE. The data for both ARPE and RPE were taken using a uniform input distribution of estimates for the $\pi/2$ -time. This distribution was taken over a relative range of 30% of a centre value for the $\pi/2$ -time obtained using standard calibration techniques before the trials were performed. For ARPE, the input distribution $P(\theta)$ was set to a Gaussian according to equation (5.6) with $\mu = \pi/2$, and $\sigma = \pi/10$.

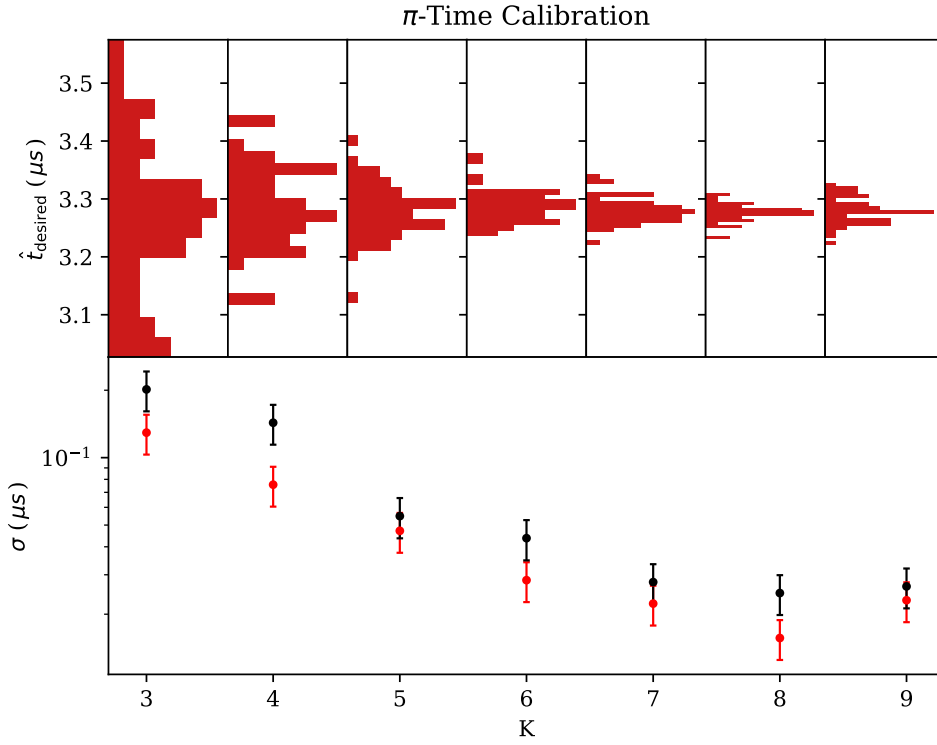


Figure 5.5: Experimental data comparing π -time calibration for ARPE versus RPE. The data for both ARPE and RPE were taken using a uniform input distribution of estimates for the π -time. This distribution was taken over a relative range of 30% of a centre value for the π -time obtained using standard calibration techniques before the trials were performed. For ARPE, the input distribution $P(\theta)$ was set to a Gaussian according to equation (5.6) with $\mu = \pi$, and $\sigma = \pi/8$.

Chapter 6

Summary and Outlook

In chapter 4 I showed how the robust phase estimation (RPE) technique described by Kimmel *et al.* can be used to calibrate laser pulses to perform arbitrary single-qubit gates with a $^{40}\text{Ca}^+$ ion qubit. From the results of applying this technique to our experimental system we find that we can achieve calibration of pulse frequency with sample standard deviations as low as 240 Hz. For $\pi/2$ -time and π -time calibrations we find sample standard deviations as low as 28 ns (1.56 % of measured $\pi/2$ -time), and 57 ns (1.45 % of measured π -time), respectively. The standard deviations of the obtained estimates for the frequency are limited by the finite coherence time of the qubit. For the time calibrations, we may also be limited by intensity fluctuations of the 729 nm laser light resonant with the qubit transition.

On the one hand, we can make improvements to our experiments in order to reduce fluctuations in laser frequency or qubit frequency to improve qubit coherence. In chapter 3, I described a technique using complex linear regression which we used to remove harmonic noise from our qubit transition frequency without which the accuracies for both frequency and time calibrations using the techniques of chapters 4 and 5 would have been much lower.

However, since decoherence is nearly always present in experimental systems, optimising the calibration protocols in these circumstances is also an important problem to solve in order to achieve better accuracy calibrations. An analysis of the effects of increasing the number of shots S_K performed at the last step of the RPE protocol was presented showing that we may improve performance in the presence of decoherence. In particular, although it is possible to improve the accuracy of estimation at a rate $1/\sqrt{T}$ by averaging the estimates obtained by applying RPE with a number of shots specified by equation (4.3), we can get better scaling by increasing S_K . Preliminary results from applying the RPE protocol experimentally with different numbers of shots S_K demonstrates improved accuracy in the estimation that roughly agrees with results from simulation. We also find that, in the presence of the sources of decoherence mentioned above, the optimal values of the parameter K used to apply the protocol are 9, 6, and 5 for the frequency, $\pi/2$ -time, and π -time, respectively.

A novel phase estimation protocol which I refer to as adaptive robust phase estimation (ARPE) was described in chapter 5. This technique combines adaptive Bayesian phase estimation with the general procedure used in RPE. Results from simulation show improved accuracy of estimation. An increase in accuracy by a

factor of 1.69 is found from simulation results for frequency calibration, and improvement by a factor of 1.85 is found for π -time calibration compared to RPE. By performing an optimisation on the number of shots used, as well as the parameter used to spread the distribution after each step of the procedure, further improvements in performance may be achievable.

Experimental results from applying ARPE to calibrate pulse frequency and time show improved accuracies compared to RPE. A minimum sample standard deviation of 102 Hz is observed for frequency estimation. For $\pi/2$ -time estimation a minimum of 9 ns (0.59 % of measured $\pi/2$ -time) is observed, and for the π -time, a minimum of 16 ns (0.49 % of measured π -time). From these estimation accuracies the contribution of static calibration errors to the overall process infidelity of a π -pulse is estimated to be 4×10^{-6} .

These results demonstrate accurate calibration of single-qubit gates in much less time than standard techniques, and ultimately should enable higher accuracy single-qubit operations to be achieved. As more complex quantum information processing experiments and algorithms are performed on larger systems, these or similar techniques will become necessary in order to maintain accurate control. As such this thesis has investigated methods which will become increasingly important to minimise coherent control errors and which could eventually allow quantum error correction thresholds to be achieved in more complex experiments. In particular, it should also be possible to devise similar procedures for multi-qubit operations where their benefit may be even greater.

Appendix A

Improving Accuracy in the Presence of Decoherence

A.1 Theory

Results presented in section 4.2 for frequency and time calibration using RPE showed that the accuracy of estimation is either limited by the finite coherence of the qubit in the case of the frequency, or possibly also by fluctuations in laser intensity for time calibrations. Although averaging over estimates can improve the accuracy further at a rate of $\sigma \propto 1/\sqrt{T}$, it may be possible to get better scaling by changing some of the parameters in the RPE protocol. In particular, I will show results obtained by simulating the RPE protocol that suggest that better than $1/\sqrt{T}$ scaling can be achieved mainly by increasing the number of shots S_K performed at the last step of the protocol. The number of shots given by (4.3), with $\alpha = 5/2$ and $\beta = 1/2$, has been optimised assuming no decoherence. Roughly speaking, the best scaling is achieved by increasing the number of applied qubit rotations from k_j to k_{j+1} as soon as the current estimate for the phase is accurate enough that the principal range of the estimate $\widehat{k_{j+1}A}$ can be determined correctly with high enough probability. By applying a number of shots according to equation (4.3), we have $S_K = \lceil 1/2 \rceil = 1$. When decoherence is present this is no longer optimal since increasing the number of applied rotations per shot further gives no advantage. Therefore, equation (4.3) should only be used to determine the number of shots when K is low enough that decoherence has little or no effect. When the maximum K value is reached before decoherence takes effect, then we may continue to obtain improvements in the accuracy of the estimates by increasing S_K .

Simulations of the RPE protocol were performed to obtain quantitative estimates for the effectiveness of this strategy. The maximum likelihood estimate for the variance was used to estimate the standard deviation:

$$\hat{\sigma} = \sqrt{\frac{1}{N} \sum_i^N (\hat{e}_i - \bar{e})^2}$$

where $\hat{e}_i$ is the estimate for the desired parameter (frequency or time), and $\bar{e}$ is the sample mean. The mean error over N calibration trials was calculated as

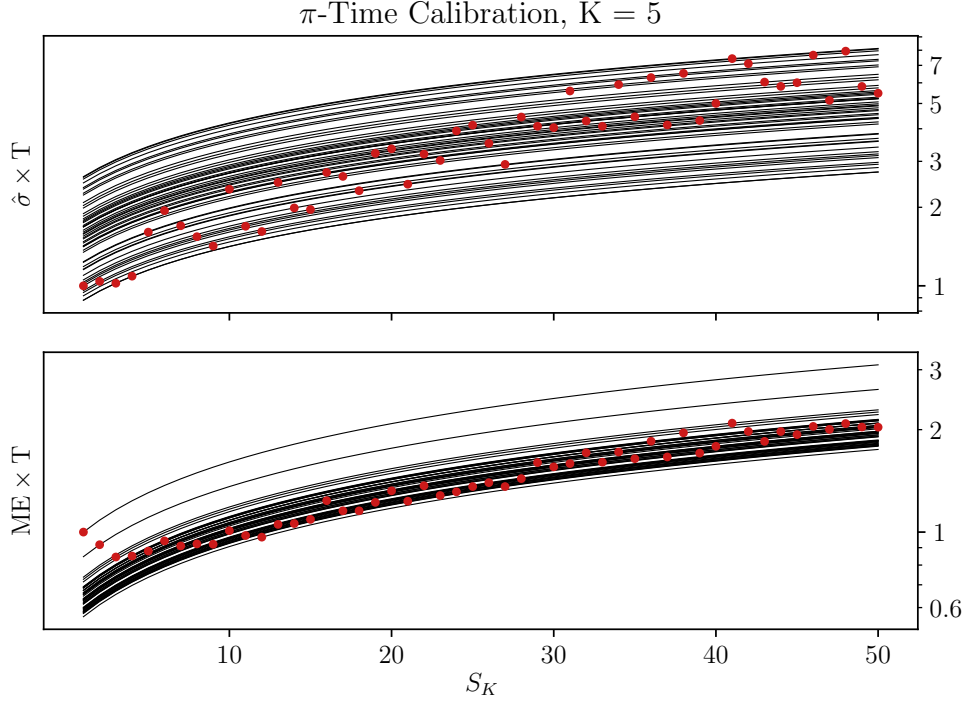


Figure A.1: Sample standard deviations and mean errors obtained by simulating 10000 trials of π -time calibration with RPE as we increase the number of shots S_K . The red points correspond to the scaling (σT or $ME \times T$) computed from the simulated samples and normalised by the value obtained for $S_K = 1$; a black line $\propto 1/\sqrt{T}$ passes through each point indicating the scaling that would result from averaging the estimates for that value of S_K .

$$ME = \frac{1}{N} \sum_i^N |\hat{e}_i - e_i|$$

where e_i is the true value used by the simulation.

Results from simulating the RPE protocol for different numbers of shots S_K are shown in figure A.1 for π -time calibrations. For each value of S_K the values of $\hat{\sigma}$ and the ME were computed for samples of 10000 calibration trials. To obtain the values plotted, the total “time” computed from (4.1) was multiplied by $\hat{\sigma}$ (ME) for each point, and then all values were normalised by the value of σT ($ME \times T$) for the $S_K = 1$ case which corresponds to the unchanged RPE protocol. Therefore, a horizontal line $y = 1$ on the plots corresponds to the scaling of the RPE protocol in the absence of decoherence. For each of the points on the plots, the $1/\sqrt{T}$ scaling curve passing through that point is also plotted; this represents the expected scaling we would obtain if we were to average the estimates obtained at that point ¹.

The results for π -time calibration plotted in figure A.1 show that increasing the value of S_K may initially improve the standard deviation of the estimates. But

¹Although the x-axes on the plots corresponds to the number of shots S_K , I plot the curves $\propto 1/\sqrt{T}$ by calculating the “time” $T(S_K)$ as a function of the number of shots.

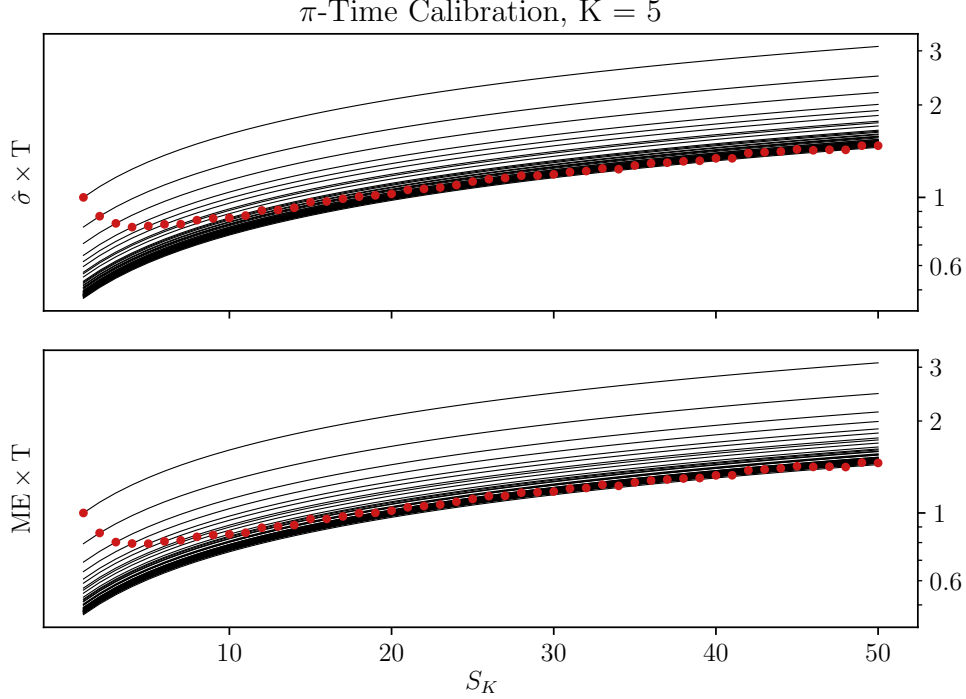


Figure A.2: Sample standard deviations and mean errors obtained by simulating 10000 trials of π -time calibration with RPE as we increase the number of shots S_K , and after removing outlying estimates. The red points correspond to the scaling (σT or $ME \times T$) computed from the simulated samples and normalised by the value obtained for $S_K = 1$; a black line $\propto 1/\sqrt{T}$ passes through each point indicating the scaling that would result from averaging the estimates for that value of S_K .

it's not clear that this will give much advantage over averaging, and when S_K is increased above ~ 10 we see that averaging estimates is likely to give better scaling than increasing S_K . These results are consistent with the analysis by Kimmel *et al.* since we are unable to obtain better scaling than the RPE procedure when the number of shots performed are always given by (4.3) (none of the points in the plot for $\hat{\sigma}$ have a y-value less than one).

Upon closer inspection, the results for increased S_K were found to have a small number (less than $\sim 2\%$) of significant outlying estimates which are often easily identified in practice. This is no surprise since the number of shots at previous steps has not been re-optimised with the new values of S_K . Removing the outlying estimates results in the data plotted for the π -time in figure A.2. The underlying distribution of estimates improves in both $\hat{\sigma}$ and the ME. So a simple way to improve performance would be to remove outlying estimates whenever enough data is taken in practice to be able to easily detect any outliers. Results for frequency calibration are analogous.

By following the work by Kimmel *et al.* to optimise the number of shots S_j , we can try reduce the number of outliers by performing a more detailed analysis. They define the probability of obtaining an error

$$p_{\text{error}}(k_j A) \equiv P \left[k_j(\hat{A}_j - A) \geq \frac{\pi}{2} \vee k_j(\hat{A}_j - A) < -\frac{\pi}{2} \right],$$

and derive the upper bound

$$p_{\text{max}}(S_j) \equiv \frac{1}{\sqrt{2\pi S_j} 2^{S_j}} > p_{\text{error}}(k_j A). \quad (\text{A.1})$$

Conditioned on having no errors up to a step $j - 1$ of the procedure, then the error at step j is bounded by $\pi/2^{j-1}$. Combining this with the upper bound (A.1), we can write an upper bound for the variance of the phase estimate

$$\sigma^2(\hat{A}) \leq [1 - p_{\text{max}}(S_K)] \left(\frac{\pi}{2^K} \right)^2 + \sum_{j=1}^K \left(\frac{\pi}{2^{j-1}} \right)^2 p_{\text{max}}(S_j). \quad (\text{A.2})$$

Combining this with (4.1) and setting $\delta_{S_j}(\sigma^2(\hat{A})T^2) = 0$ leads to the form of equation (4.3). Since we have changed the number of shots S_K , we need to re-optimize the number of shots performed at each step j of the procedure. However equation (A.2) will not capture the effect of increasing S_K since the first term rapidly converges to a *constant* value. We would need to generalise the upper bound in equation (A.1) to arbitrary bounding ranges (above the range is $(-\frac{\pi}{2}, \frac{\pi}{2}]$). Here I make some assumptions to get a rough estimate for such a bound. I compute the exact expected standard deviations of the estimates for the fixed rotation procedure with $S \in \{1, 2, \dots, 30\}$ and fit the result to obtain $\sigma(\hat{A}) \propto S^{-0.48}$. Assuming a normal distribution of estimates we can approximate the bound (A.1) by integrating the distribution outside the range $(-\frac{\pi}{2}, \frac{\pi}{2}]$ to get

$$\hat{p}_{\text{max}} = 1 + \text{erf} \left(\frac{-\pi/2}{\sigma(\hat{A})\sqrt{2}} \right).$$

Using the functional form I found for $\sigma(\hat{A})$ I determine the proportionality constant which gives the best agreement between $\hat{p}_{\text{max}}$ and p_{max} to get

$$\hat{p}_{\text{max}} = 1 + \text{erf} \left(\frac{-(\pi/2)S^{0.48}}{1.225\sqrt{2}} \right).$$

Taking a range $(-\varepsilon, \varepsilon]$, we can get an estimate for the error probabilities of this range as

$$\hat{p}_{\text{max}} = 1 + \text{erf} \left(\frac{-\varepsilon S^{0.48}}{1.225\sqrt{2}} \right). \quad (\text{A.3})$$

Now I use the estimate for the probability of an error (A.3) to derive a more appropriate bound for the scaling as we change the number of shots S_K . Letting $\gamma = S_K$ and choosing the error range given by

$$\varepsilon(\gamma) = \frac{\pi}{2\gamma^{0.3}},$$

leads to

γ	α	β	$\sigma(\hat{A})T$
1	3	1	11.31π
2	3	1	10.42π
3	3	2	10.15π
4	3	2	10.14π
5	3	3	10.25π

Table A.1: Optimal values obtained for α and β by minimising equation (A.4) for values of γ from 1 to 5. The values for $\sigma(\hat{A})T$ are the approximate upper bounds computed from equation (A.4).

$$\hat{p}_{\max}(\gamma) = 1 + \operatorname{erf}\left(\frac{-\pi\gamma^{0.18}}{2.45\sqrt{2}}\right).$$

Combining this with equation (A.2) I write a new approximate bound for the variance

$$\begin{aligned} \sigma^2(\hat{A}) \sim & \leq [1 - p_{\max}(\gamma)] \left(\frac{\pi}{2^K}\right)^2 (\hat{p}_{\max}(\gamma) + [1 - \hat{p}_{\max}(\gamma)]\gamma^{-0.6}) + \\ & \left(\frac{\pi}{2^{K-1}}\right)^2 p_{\max}(\gamma) + \sum_{j=1}^{K-1} \left(\frac{\pi}{2^{j-1}}\right)^2 p_{\max}(S_j). \end{aligned}$$

Taking $\delta_{S_j}(\sigma^2(\hat{A}T^2)) = 0$ leads to the same functional form (4.3) for S_j when $j < K$, only now we are explicitly assuming a discontinuous function. Making use of equations (4.3) and (4.1) I find the approximate bound for the scaling

$$\begin{aligned} \sigma(\hat{A})T \sim & \leq \pi(2\alpha + \beta + \gamma) \left\{ [1 - p_{\max}(\gamma)] [\gamma^{-0.6} + (1 - \gamma^{-0.6})\hat{p}_{\max}(\gamma)] + \right. \\ & \left. 4p_{\max}(\gamma) + 16p_{\max}(\alpha + \beta) \left(\frac{2^\alpha}{2^\alpha - 4}\right) \right\}^{\frac{1}{2}}. \end{aligned} \quad (\text{A.4})$$

Now we can try to optimise the scaling as we increase γ by minimising equation (A.4) over the possible values of α and β . I found this to give rough agreement with simulation over a range of γ . Optimal values obtained for γ from 1 to 5 and the corresponding approximate bounds are listed in table A.1.

A notable result is that we find optimal scaling for $\gamma = 4$, $\alpha = 3$, and $\beta = 2$ which is in fact better than that found by Kimmel *et al.*. This result was confirmed by simulation, and can perhaps be explained by the fact that, since we have assumed a discontinuous function, we are exploring a solution that could not have been obtained by setting $\delta_{S_j}(\sigma^2(\hat{A}T^2)) = 0$, and assuming a continuous function. It's also not so surprising that we can find better solutions than the optimal case found by Kimmel *et al.* since they optimised a theoretical upper bound, whereas what we really care about and what I measured in simulation is the *average* performance.

Thus we are able to maintain good scaling for larger values of γ , and the standard deviation of the estimates can be reduced with better scaling than we would get by averaging.

To make these results most useful, it may be possible to determine a function for the optimal values of α and β in terms of γ . This could be done by minimising equation A.4 or a more rigorous bound on the scaling. But better performance can probably be found by optimising the values in simulation and fitting the result, or deriving a theoretical expression for the *average* scaling since what we really care about is the best average performance rather than the best scaling in the worst case (i.e. an upper bound).

A.2 Experiment

To experimentally verify the performance of the RPE protocol with an increased number of shots S_K , estimates for the frequency, $\pi/2$ -time, and π -time were obtained for 50 calibration trials at several different K values. For each K value, calibrations were performed with a number of shots S_K set to 1, 7, 20, and 45 (i.e. 50 trials for each). These values for the numbers of shots roughly correspond to repeated doubling of the “time” T in (4.1).

The experimental results in figures A.5, A.6, and A.7 can be compared visually with expected results from simulating the same S_K values for 50000 calibration trials shown in figures A.3 and A.4. The simulated calibrations are for the case with no decoherence, so we expect the experimental results to roughly match when the effects of decoherence are small.

For the frequency, similar results to simulation are obtained for $K = 9$; increasing K further we see that changing S_K does not clearly have any effect. This result agrees with the results presented in section 4.2; as we apply longer wait-times the qubit decoheres and we no longer obtain any information from the measurement results ².

For the $\pi/2$ -time and π -time calibrations, similar effects are observed, and the results again agree roughly with the results in section 4.2. Here, we find that the best performance is obtained when increasing S_K for $K = 6$ and $K = 5$ for the $\pi/2$ -time and π -time calibrations, respectively.

We can conclude from the analysis and results that the best accuracies are probably obtained for our system by choosing the K values of 9, 6, and 5 for frequency, $\pi/2$ -time, and π -time calibrations, respectively, when using $^{40}\text{Ca}^+$.

²I have not shown the results for $K = 10$ here; they are somewhat better than for $K = 11$, but the decrease in variance with S_K is not as significant as for $K = 9$ indicating that decoherence probably begins to take effect.

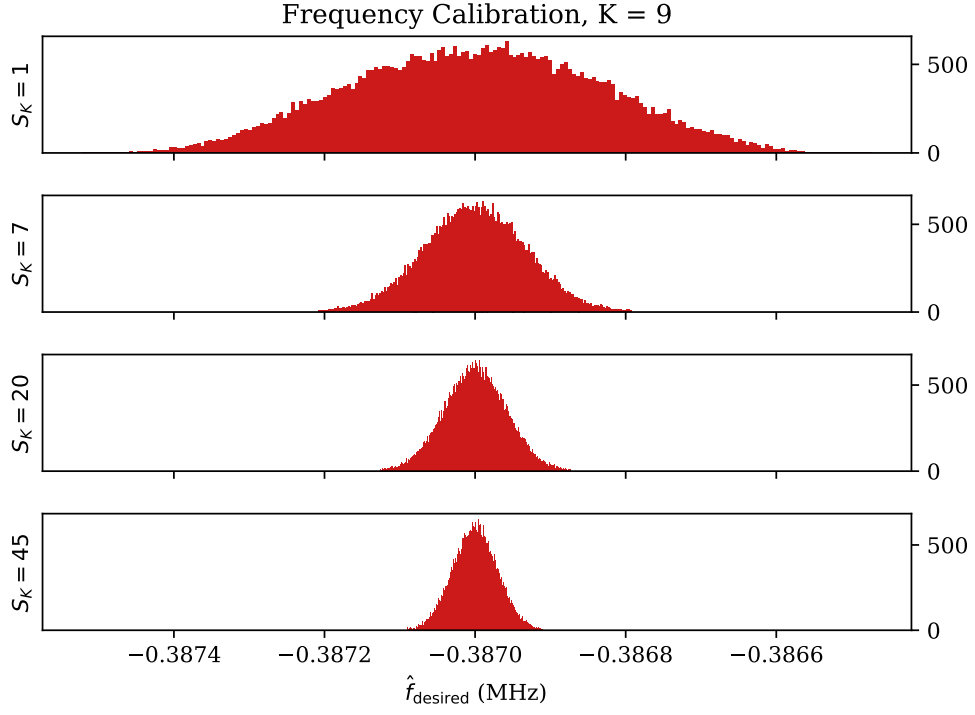


Figure A.3: Histogramed estimates from simulating 50000 trials of frequency calibration with RPE as we increase the number of shots S_K .

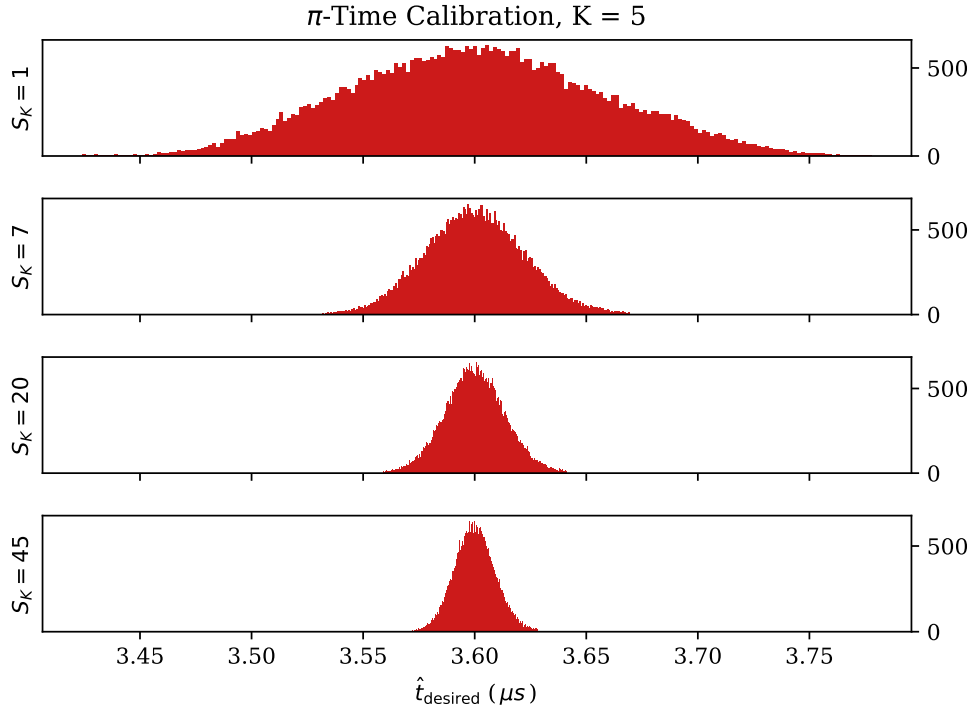


Figure A.4: Histogramed estimates from simulating 50000 trials of π -time calibration with RPE as we increase the number of shots S_K .

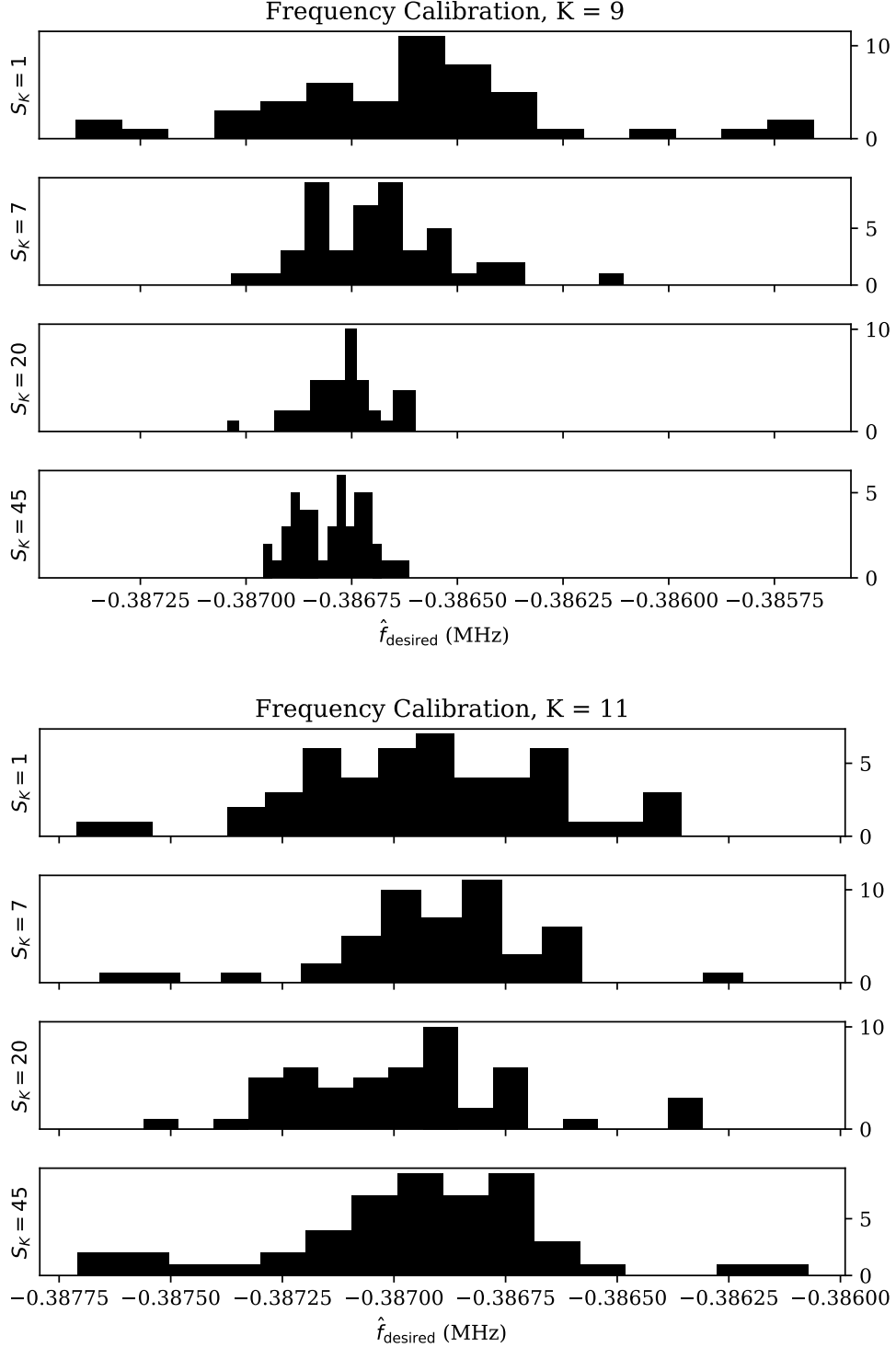


Figure A.5: Histogrammed estimates for 50 experimental trials of frequency calibration for different values of K and S_K with a $^{40}\text{Ca}^+$ qubit. Increasing S_K when K is larger than ~ 9 no longer improves the accuracy since the wait-time used at step $j = K$ of the protocol is longer than the coherence time of the qubit.

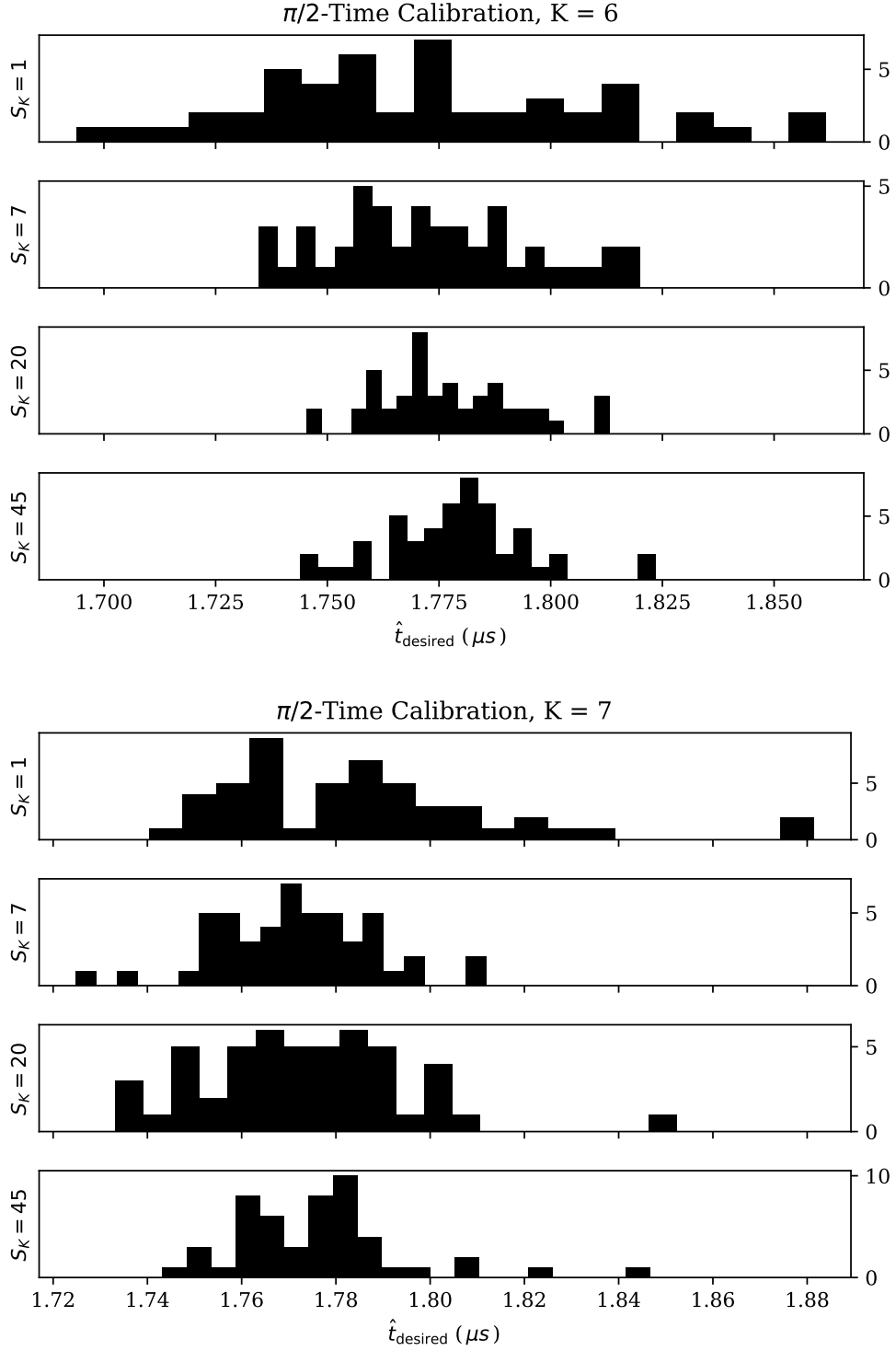


Figure A.6: Histogramed estimates for 50 experimental trials of $\pi/2$ -time calibration for different values of K and S_K with a $^{40}\text{Ca}^+$ qubit. Increasing S_K when K is larger than ~ 6 no longer improves the accuracy due to incoherent errors.

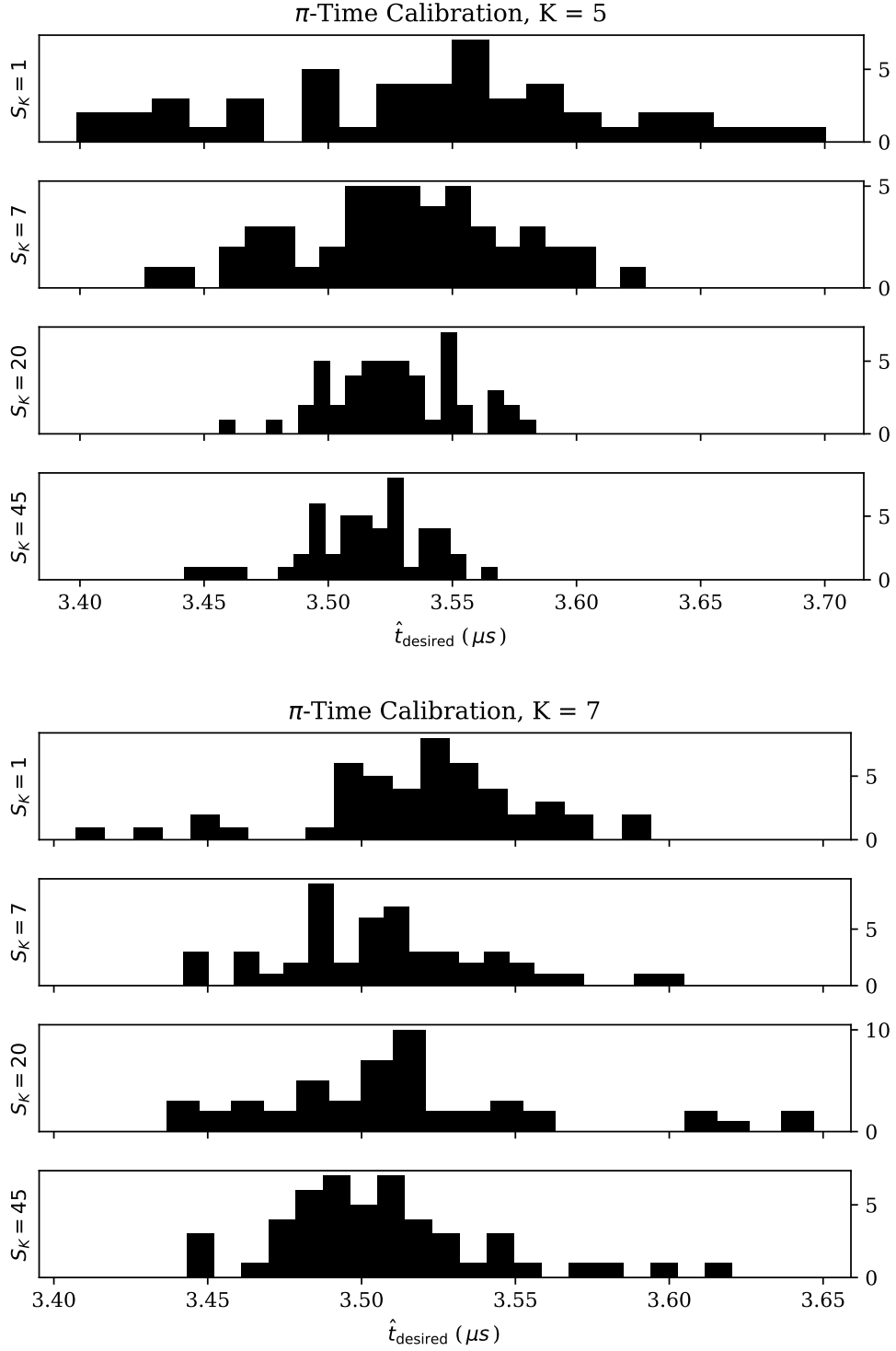


Figure A.7: Histogrammed estimates for 50 experimental trials of π -time calibration for different values of K and S_K with a $^{40}\text{Ca}^+$ qubit. Increasing S_K when K is large that ~ 5 no longer improves the accuracy due to incoherent errors.

Bibliography

- [1] R. P. Feynman, “Simulating Physics with Computers,” *International Journal of Theoretical Physics*, vol. 21, pp. 467–488, 1982.
- [2] W. G. Unruh, “Maintaining coherence in quantum computers,” *Physical Review A*, vol. 51, pp. 992–997, 1995.
- [3] P. W. Shor, “Scheme for reducing decoherence in quantum computer memory,” *Physical Review A*, vol. 52, pp. R2493–R2496, 1995.
- [4] A. M. Steane, “Error Correcting Codes in Quantum Theory,” *Physical Review Letters*, vol. 77, pp. 793–797, 1996.
- [5] A. R. Calderbank and P. W. Shor, “Good quantum error-correcting codes exist,” *Physical Review A*, vol. 54, pp. 1098–1105, 1996.
- [6] R. Laflamme, C. Miquel, J. P. Paz, and W. H. Zurek, “Perfect Quantum Error Correcting Code,” *Physical Review Letters*, vol. 77, pp. 198–201, 1996.
- [7] C. H. Bennett, D. P. DiVincenzo, J. A. Smolin, and W. K. Wootters, “Mixed-state entanglement and quantum error correction,” *Physical Review A*, vol. 54, pp. 3824–3851, 1996.
- [8] E. Knill, R. Laflamme, and W. Zurek, “Accuracy Threshold for Quantum Computation,” 1996.
- [9] D. Aharonov and M. Ben-Or, “Fault-tolerant Quantum Computation with Constant Error,” in *Proceedings of the Twenty-ninth Annual ACM Symposium on Theory of Computing*, STOC ’97, (New York, NY, USA), pp. 176–188, Association for Computing Machinery, 1997.
- [10] D. Deutsch, “Quantum Theory, the Church-Turing Principle and the Universal Quantum Computer,” *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, vol. 400, pp. 97–117, 1985.
- [11] A. Barenco, C. Bennett, R. Cleve, D. Divincenzo, N. Margolus, P. Shor, T. Sleator, J. Smolin, and H. Weinfurter, “Elementary Gates for Quantum Computation,” *Physical Review A*, vol. 52, pp. 3457–3467, 1995.
- [12] A. Barenco, “A universal two-bit gate for quantum computation,” *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, vol. 449, pp. 679–683, 1995.

- [13] D. Deutsch, A. Barenco, and A. Ekert, “Universality in quantum computation,” *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences*, vol. 449, pp. 669–677, 1995.
- [14] P. Boykin, T. Mor, M. Pulver, V. Roychowdhury, and F. Vatan, “A new universal and fault-tolerant quantum basis,” *Information Processing Letters*, vol. 75, pp. 101–107, 2000.
- [15] S. Haroche and J.-M. Raimond, *Exploring the Quantum Atoms, Cavities and Photons*. Oxford University Press, 2006.
- [16] D. P. DiVincenzo, “The Physical Implementation of Quantum Computation,” *Fortschritte der Physik*, vol. 48, pp. 771–783, 2000.
- [17] A. Sørensen and K. Mølmer, “Quantum Computation with Ions in Thermal Motion,” *Physical Review Letters*, vol. 82, pp. 1971–1974, 1999.
- [18] K. Mølmer and A. Sørensen, “Multiparticle Entanglement of Hot Trapped Ions,” *Physical Review Letters*, vol. 82, pp. 1835–1838, Mar. 1999.
- [19] A. Sørensen and K. Mølmer, “Entanglement and quantum computation with ions in thermal motion,” *Physical Review A*, vol. 62, p. 022311, 2000.
- [20] C. M. Caves, “Quantum-mechanical noise in an interferometer,” *Physical Review D*, vol. 23, pp. 1693–1708, 1981.
- [21] S. Barnett, C. Fabre, and A. Maitre, “Ultimate quantum limits for resolution of beam displacements,” *The European Physical Journal D - Atomic, Molecular, Optical and Plasma Physics*, vol. 22, pp. 513–519, 2003.
- [22] M. J. Holland and K. Burnett, “Interferometric detection of optical phase shifts at the Heisenberg limit,” *Physical Review Letters*, vol. 71, pp. 1355–1358, 1993.
- [23] D. J. Wineland, J. J. Bollinger, W. M. Itano, and D. J. Heinzen, “Squeezed atomic states and projection noise in spectroscopy,” *Physical Review A*, vol. 50, pp. 67–88, 1994.
- [24] G. E. Hinton, S. Osindero, and Y.-W. Teh, “A Fast Learning Algorithm for Deep Belief Nets,” *Neural Computation*, vol. 18, pp. 1527–1554, 2006.
- [25] S. Kimmel, G. H. Low, and T. J. Yoder, “Robust calibration of a universal single-qubit gate set via robust phase estimation,” *Physical Review A*, vol. 92, p. 062315, 2015.
- [26] B. L. Higgins, D. W. Berry, S. D. Bartlett, M. W. Mitchell, H. M. Wiseman, and G. J. Pryde, “Demonstrating Heisenberg-limited unambiguous phase estimation without adaptive measurements,” *New Journal of Physics*, vol. 11, p. 073023, 2009.
- [27] D. Leibfried, R. Blatt, C. Monroe, and D. Wineland, “Quantum dynamics of single trapped ions,” *Reviews of Modern Physics*, vol. 75, pp. 281–324, 2003.

- [28] D. Kienzler, *Quantum Harmonic Oscillator State Synthesis by Reservoir State Engineering*. PhD thesis, ETH Zürich, 2015.
- [29] L. E. de Clercq, *Transport Quantum Logic Gates for Trapped Ions*. PhD thesis, ETH Zürich, 2015.
- [30] M. Block, O. Rehm, P. Seibert, and G. Werth, “3d D-2(5/2) lifetime in laser cooled Ca⁺: Influence of cooling laser power,” *The European Physical Journal D - Atomic, Molecular, Optical and Plasma Physics*, vol. 7, pp. 461–465, 1999.
- [31] F. Schmidt-Kaler, J. Eschner, G. Morigi, C. F. Roos, D. Leibfried, A. Mundt, and R. Blatt, “Laser cooling with electromagnetically induced transparency: application to trapped samples of ions or neutral atoms,” *Applied Physics B-Lasers and Optics*, vol. 73, pp. 807–814, 2001.
- [32] Y. Lin, J. P. Gaebler, T. R. Tan, R. Bowler, J. D. Jost, D. Leibfried, and D. J. Wineland, “Sympathetic Electromagnetically-Induced-Transparency Laser Cooling of Motional Modes in an Ion Chain,” *Physical Review Letters*, vol. 110, p. 153002, 2013.
- [33] C. Monroe, D. M. Meekhof, B. E. King, S. R. Jefferts, W. M. Itano, D. J. Wineland, and P. Gould, “Resolved-Sideband Raman Cooling of a Bound Atom to the 3d Zero-Point Energy,” *Physical Review Letters*, vol. 75, pp. 4011–4014, 1995.
- [34] H. Yu Lo, *Creation of Squeezed Schrödinger’s Cat States in a Mixed-Species Ion Trap*. PhD thesis, ETH Zürich, 2015.
- [35] D. Nadlinger, “Laser Intensity Stabilization and Pulse Shaping for Trapped-Ion Experiments using Acousto-Optic Modulators,” Semester project report, ETH Zürich, 2013.
- [36] J. J. Sakurai and J. Napolitano, *Modern Quantum Mechanics*. Pearson, 2nd ed., 2011.
- [37] A. Datta and A. Shaji, “Quantum metrology without quantum entanglement,” *Modern Physics Letters B*, vol. 26, p. 1230010, 2012.
- [38] C. Flühmann, “Stabilizing lasers and magnetic fields for quantum information experiments,” Master’s thesis, ETH Zürich, 2014.
- [39] E. L. Lehmann, *Theory of Point Estimation*. Wiley, 1983.
- [40] F. James, *Statistical Methods in Experimental Physics*. World Scientific, 2nd ed., 2006.
- [41] D. Nadlinger, “Experimental Control and Benchmarking for Single-Qubit Trapped-Ion Transport Gates,” Semester project report, ETH Zürich, 2016.
- [42] A. Lebedev, “Metrology Report.” private communication. Dr. Andrey Lebedev, Institut für Theoretische Physik, ETH Zürich.

- [43] L. H. Pedersen, N. M. Møller, and K. Mølmer, “Fidelity of quantum operations,” *Physics Letters A*, vol. 367, no. 1, pp. 47–51, 2007.