



Charge noise and spin in single-electron silicon CMOS quantum dots

Cameron Spence

► To cite this version:

Cameron Spence. Charge noise and spin in single-electron silicon CMOS quantum dots. Mesoscopic Systems and Quantum Hall Effect [cond-mat.mes-hall]. Université Grenoble Alpes [2020-..], 2021. English. NNT : 2021GRALY054 . tel-03566464

HAL Id: tel-03566464

<https://theses.hal.science/tel-03566464>

Submitted on 11 Feb 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

THÈSE

Pour obtenir le grade de

DOCTEUR DE L'UNIVERSITÉ GRENOBLE ALPES

Spécialité : Physique de la Matière Condensée et du Rayonnement

Arrêté ministériel : 25 mai 2016

Présentée par

Cameron SPENCE

Thèse dirigée par **Tristan MEUNIER**, Université Grenoble Alpes
et co-encadrée par **Matias URDAMPILLETA**, Université Grenoble Alpes

préparée au sein du **Laboratoire Institut Néel**
dans l'**École Doctorale Physique**

**Bruit de charge et spin dans des boîtes
quantiques CMOS à électron unique**

**Charge noise and spin in single-electron
silicon CMOS quantum dots**

Thèse soutenue publiquement le **30 septembre 2021**,
devant le jury composé de :

Monsieur PHILIPPE LAFARGE

PROFESSEUR DES UNIVERSITÉS, UNIVERSITÉ DE PARIS,
Rapporteur

Madame EVA DUPONT-FERRIER

PROFESSEUR, Université de Sherbrooke, Examinatrice

Monsieur GERARD GHIBAUDO

DIRECTEUR DE RECHERCHE, CNRS DELEGATION ALPES,
Examineur

Monsieur JEAN-PHILIPPE POIZAT

DIRECTEUR DE RECHERCHE, CNRS DELEGATION ALPES, Président

Madame VALIA VOLIOTIS

PROFESSEUR DES UNIVERSITÉS, SORBONNE UNIVERSITÉ,
Rapporteuse



Abstract

The push towards a scalable quantum computer is entering a crucial phase, with several different solid-state qubit designs demonstrated as strong candidates for the basis of a future quantum computer. One such candidate is silicon-based quantum dot spin qubits, a relatively recent entry into the field but with potential for long coherence times and high-fidelity measurement and manipulation. Silicon also has a significant advantage in leveraging industrial expertise and compatibility with current semiconductor fabrication techniques to produce reliable, reproducible and scalable qubit designs. To develop a fully-functional quantum computer at scale, the individual qubit design should be compact with a low control overhead, and have minimal interaction with the environment to prevent loss of quantum information. One of the major environmental interactions experienced by quantum dots is the characteristic $1/f$ charge noise present in all electronic devices. Electric field fluctuations in the region of the quantum dot pose a significant challenge for single- and multi-qubit measurement, manipulation, and coherence in semiconductor spin qubits.

In this thesis, we investigate and characterize the charge noise experienced by quantum dots in CMOS silicon nanowire devices. We use frequency-domain analysis to determine the variability in the charge noise experienced by a quantum dot due to its position within the nanowire and its interaction with different interfaces, finding typical charge noise values on the order of $10 \mu\text{eV}^2/\text{Hz}$ in the many-electron regime. We analyse the energy spectrum of charge noise fluctuators in the region of a typical quantum dot and identify two frequencies which dominate the charge noise in the temperature range of interest. We also demonstrate a novel technique of measuring the charge noise experienced by a single electron, a regime which is highly applicable to future qubit development, and extract the single-electron charge noise value of $(130 \pm 60) \mu\text{eV}^2/\text{Hz}$ at a typical measurement temperature of 400 mK. The potential decoherence induced by this degree of charge noise is compared to the decoherence caused by the hyperfine field in natural silicon, and we find that the theoretical charge noise-limited T_2^* is more than two orders of magnitude longer than the nuclear spin limited T_2^* , underpinning the benefits of isotopic purification on the spin coherence time.

Secondly, we demonstrate single-shot measurement of the spin of a single electron and characterization of the spin physics in a nanowire quantum dot. We demonstrate measurement of the spin-lattice relaxation time T_1 with a state visibility greater than 90%. The spin-valley relaxation hotspot is detected via magnetic field spectroscopy, finding a

valley splitting energy of $(297 \pm 5) \mu\text{eV}$. Finally, we analyse the magnetic field anisotropy of the spin-valley mixing, demonstrating suppression of the relaxation mechanism by an order of magnitude in a field oriented along the main symmetry axis of the device. These experiments have developed our understanding of the noise environment and spin physics present in CMOS nanowire quantum dots, and form the foundations of in-depth characterization that can be applied at scale to direct future development of spin qubits in silicon.

Abstract - Français

La développement d'un ordinateur quantique évolutif entre dans une phase cruciale, avec plusieurs conceptions de qubit à semi-conducteurs différentes démontrées comme de solides candidats pour la base d'un futur ordinateur quantique. Un de ces candidats est le qubit de spin quantique à base de silicium, une entrée relativement récente dans le domaine mais avec un potentiel pour de longs temps de cohérence et une mesure et une manipulation haute fidélité. Le silicium a également un avantage significatif en tirant parti de l'expertise industrielle et de la compatibilité avec les techniques actuelles de fabrication de semi-conducteurs pour produire des conceptions de qubit fiables, reproductibles et évolutives. Pour développer un ordinateur quantique entièrement fonctionnel à grande échelle, la conception de qubit individuel doit être compacte avec une faible surcharge de contrôle et avoir une interaction minimale avec l'environnement pour éviter la perte d'informations quantiques. Une des principales interactions environnementales rencontrées par les boîtes quantiques est le bruit de charge caractéristique $1/f$ présent dans tous les appareils électroniques. Les fluctuations du champ électrique dans la région de la boîte quantique posent un défi important pour la mesure, la manipulation et la cohérence des qubits de spin à un ou plusieurs qubits.

Dans cette thèse, nous étudions et caractérisons le bruit de charge ressenti par les boîtes quantiques dans les dispositifs à nanofils de silicium CMOS. Nous utilisons l'analyse du domaine fréquentiel afin de déterminer la variabilité du bruit de charge subi par un boîte quantique en raison de sa position dans le nanofil et de son interaction avec des différentes interfaces, en trouvant des valeurs de bruit de charge typiques de l'ordre de $10 \mu\text{eV}^2/\text{Hz}$ dans le régime à plusieurs électrons. Nous analysons le spectre d'énergie des fluctuations du bruit de charge dans la région d'un boîte quantique typique et identifions deux fréquences qui dominent le bruit de charge dans la plage de température d'intérêt. Nous démontrons également une nouvelle technique de mesure du bruit de charge subi par un seul électron, un régime qui est applicable au développement futur des qubits, et extrayons la valeur du bruit de charge à un électron de $(130 \pm 60) \mu\text{eV}^2/\text{Hz}$ à une température de mesure typique de 400 mK. La décohérence potentielle induite par ce degré de bruit de charge est comparée à la décohérence provoquée par le champ hyperfin dans le silicium naturel, et on constate que la charge théorique limitée au bruit T_2^* est plus de deux ordres de grandeur plus longue que le nucléaire spin limité T_2^* , sous-tendant les avantages de la purification isotopique sur le temps de cohérence du spin.

Deuxièmement, nous démontrons la mesure dans un seul coup du spin d'un seul électron

et la caractérisation de la physique du spin dans un boîte quantique à nanofils. Nous démontrons la mesure du temps de relaxation spin-réseau T_1 avec une visibilité d'état supérieure à 90 %. Le point chaud de relaxation spin-vallée est détecté par spectroscopie de champ magnétique, trouvant une énergie de division de vallée de $(297 \pm 5) \mu\text{eV}$. Enfin, nous analysons l'anisotropie du champ magnétique du mélange spin-vallée, démontrant la suppression du mécanisme de relaxation d'un ordre de grandeur dans un champ orienté le long de l'axe de symétrie principal du dispositif. Ces expériences ont développé notre compréhension de l'environnement sonore et de la physique du spin présents dans les boîtes quantiques de nanofils CMOS, et forment les fondations d'une caractérisation approfondie qui peut être appliquée à l'échelle pour diriger le développement futur des qubits de spin dans le silicium.

Contents

1	Introduction	1
1.1	Introduction	1
1.1.1	Classical computation	1
1.1.2	Quantum computation	2
1.1.3	The qubit	2
1.1.4	The DiVincenzo criteria	3
1.1.5	Quantum error correction	6
1.2	Physical realizations of a quantum computer	7
1.2.1	Superconducting qubits	7
1.2.2	Photonic qubits	7
1.2.3	Trapped ion qubits	9
1.2.4	Other realizations	9
1.3	A spin-based silicon quantum computer	10
1.3.1	Si-MOS spin qubits	11
1.3.2	CMOS spin qubits	13
1.4	Thesis structure	16
2	Quantum Dots	17
2.1	Introduction	17
2.2	Quantum dots	17
2.2.1	Single dot energy	17
2.2.2	Coulomb blockade	20
2.2.3	Quantum tunnelling	24
2.2.4	Charge sensing	25
2.3	Spins in quantum dots	27
2.3.1	Spin states in a single dot	27
	Single electron spin states	27
	Two electron spin states	28
2.3.2	Spins in two quantum dots	29
2.3.3	Spin measurement	31
	Energy selective tunnelling	31
	Spin dependent tunnel rate readout	32
	Spin blockade	34

3	Experimental setup and characterization	37
3.1	Introduction	37
3.1.1	Cryogenics	37
	"Dipstick" immersion refrigerator	38
	Helium-3 circulation refrigerator	38
	Dilution refrigerator	38
3.1.2	Electronics	38
	Wiring	38
	Voltage supply	40
	Current readout	40
	Magnetic coil	40
3.1.3	Software	40
3.2	Characterization	41
3.2.1	Sample design	41
3.2.2	Room temperature characterization	42
3.2.3	Low-temperature characterization	45
	Many-electron regime	45
	Few-electron regime	49
4	Low frequency charge noise in FDSOI quantum dots	55
4.1	Introduction	55
4.2	Background	55
4.2.1	Charge noise models	55
4.2.2	Charge traps in silicon	56
4.2.3	Charge noise induced decoherence in quantum dots	58
4.3	Low frequency charge noise in the many electron regime	61
4.3.1	Introduction	61
4.3.2	1/f noise	61
4.3.3	Measurement method	63
4.3.4	Wavefunction manipulation	64
	Vertical manipulation	65
	Lateral manipulation	67
4.3.5	Identifying fluctuators	69
	Lorentzian spectra	70
	Temperature dependence of charge noise	71
	Fluctuator energies	74
	Spectral analysis	75
	Characteristic frequency analysis	76
4.4	Low frequency charge noise in the few-electron regime	77
4.4.1	Measurement method	78
4.4.2	Charge noise at the single electron level	81
4.4.3	Evolution of charge noise with electron number	81
4.4.4	Prediction of T_2^* in the presence of a magnetic field gradient	83
4.5	Conclusion	84

5	Characterisation of spin physics in a single CMOS quantum dot	87
5.1	Introduction	87
5.2	Spin detection and readout	87
5.2.1	Background	88
5.2.2	Locating spin	91
5.2.3	Measurement fidelity	93
	State visibility	94
	Initialization error	96
	Thermally activated tunnelling	97
	Readout fidelity	97
	Loading spectrum	97
	Conclusion	99
5.2.4	T1 measurement	99
5.3	Spin-valley interactions	101
5.3.1	Valley states in silicon	101
5.3.2	Relaxation mechanisms	102
5.3.3	Valley relaxation hotspot	103
5.3.4	Valley energy tuning	106
5.3.5	Relaxation field anisotropy	107
5.3.6	T_1 charge noise	112
5.4	Conclusion	115
	Conclusion	117
	Bibliography	121
	Appendix	137
	Appendix	137
A	Double Quantum Dots	137
A.1	Double dot energy	137
A.2	Double dot transport	140
A.3	Bias triangles	143

CHAPTER 1

Introduction

1.1 Introduction

Since the turn of the twentieth century, physics has been undergoing a so-called "quantum revolution". The discovery of the quantization of energy at atomic scales, as well as the discovery of quantum-specific phenomena such as quantum tunnelling, entanglement, and superposition of states, enabled a rapid swathe of developments in not only fundamental physics, but also in everyday technological applications. Quantum principles underpin technologies such as magnetic resonance imaging (MRI), electron microscopy, lasers, light-emitting diodes (LEDs) and, at the heart of the digital age, semiconductor transistors of ever-decreasing size.

1.1.1 Classical computation

The semiconductor transistor is arguably one of the most important inventions of the last century. Since its development in 1947 [Rio99], the transistor has become the building block of the powerhouse of the digital age: the classical computer. The rapid advancements in transistor technology have allowed computers to be scaled down from room-sized machines operating with punch cards and mechanical keys to today's ubiquitous laptops, smartphones, and high-density supercomputers. This progression has remarkably reliably followed Moore's law [Mac11], which predicts a doubling of the density of transistors roughly every two years. Such an exponential growth has led us from centimetre-sized transistors to the state-of-the-art Intel® 14 nm CMOS technology as of 2020, with 10 nm technology in development [Nat14]. CMOS technology in particular has emerged as by far the dominant process for fabricating cutting-edge transistors, demonstrating low noise and low power consumption [Voi13], and it is used in 99% of all modern electronic devices.

However, modern transistors are approaching the length scales where quantum effects are non-negligible. Nanoscale transistors begin to suffer from power leakage due to quantum tunnelling [Liu09] on a scale that exceeds the rate at which leakage heat can be removed from a device. These effects are likely to slow or stall the continuing increase in transistor density within the next decade. Fortunately, these same quantum effects that are detrimental to classical transistors can be exploited in a new type of computer: the quantum computer.

1.1.2 Quantum computation

Towards the end of the twentieth century, a field of thought developed combining the principles of classical computation and information processing, and the fundamental field of quantum mechanics. Aptly termed quantum information, this new field created a new paradigm of computation, leading to the development of specifically quantum algorithms and cryptography protocols. One of the first of these, developed in 1984 by Bennett and Brassard, is known as the "BB84" quantum key distribution scheme [Ben84]. BB84 is a provably secure theoretical key distribution protocol, which relies on the quantum properties of entanglement and projection - that is, the measurement of a system disturbs the system itself. Another similar protocol known as E91 was posited in 1991 by Artur Ekert [Eke91]. The widespread use of such a secure communication channel would naturally have significant consequences in an age of constant worldwide communication. However, potentially of even greater impact, is the potential for a computer based on quantum mechanics to perform computations that a classical computer could not efficiently simulate.

A picture of a truly quantum computer began to develop as the concept of a "qubit", analogous to the classical "bit" that forms the fundamental basis of a classical computer, was proposed. It is this that distinguishes the quantum computer from a classical computer. In a classical computer, information is manipulated classically, using the digital "bit". In a quantum computer, information is fundamentally quantum, using the analog "qubit", and exploiting uniquely quantum resources such as entanglement. In general, a quantum computer can efficiently simulate a classical computer, but a classical computer cannot efficiently simulate a quantum computer. This is the main advantage of quantum computers over their classical counterparts, and does not necessarily constitute a universal increase in speed, efficiency or computing power. Instead, there are certain problems which can be formulated such that they resemble a quantum system, allowing a quantum computer to provide a significant speed-up.

In the 1990s, the first purely quantum algorithms were developed, with theoretical proof of a "quantum speed-up" - that is, proof of superiority in terms of efficiency compared to the same task being performed on a classical computer. Some of these, such as the famous Shor's algorithm for factorizing large numbers and Grover's algorithm for searching a database, have significant implications for classical cryptography - much of which (such as the widely-used RSA algorithm) is based on the difficulty that classical computers have in solving such problems with "brute force" methods. These algorithms, along with several more developed in the following years, cemented the idea of a new quantum computing paradigm, in which such algorithms could be realized, being the future of information processing. However, it was known that practically implementing a quantum computer would be a formidable challenge in terms of both physics and engineering.

1.1.3 The qubit

The quantum bit, hereafter referred to as the qubit, is the fundamental unit of information in quantum information processing. The physical implementation of a qubit can take many forms, such as photons, trapped ions, electron and nuclear spins, or superconducting qubits. In general, a qubit can be any system defined by two basis states, representing 0 and 1, which can take a continuum of possible superpositions of these two states. In

Dirac notation, a qubit can be represented by $|\Psi\rangle = \alpha|0\rangle + \beta|1\rangle$, with the condition that $|\alpha|^2 + |\beta|^2 = 1$. The values of α and β can be referred to as the "amplitudes" of the two states $|0\rangle$ and $|1\rangle$ respectively. The classical state 0 (1) can be recovered when $\alpha = 1(0)$ and $\beta = 0(1)$, giving $|\Psi\rangle = |0\rangle$ and $|\Psi\rangle = |1\rangle$ respectively. However, the true value of the qubit lies in its ability to take a value that is neither $|0\rangle$ or $|1\rangle$, but a linear combination of the two - a so-called superposition state. An example of a superposition state is $\alpha = \beta = \frac{1}{\sqrt{2}}$, whereby the qubit has equal amplitude of both the $|0\rangle$ and $|1\rangle$ states.

A measurement of a qubit corresponds to a projection of a qubit into its basis states $|0\rangle$ and $|1\rangle$. One of these two results will always be obtained, regardless of the qubit state before measurement. The probability of obtaining the state $|0\rangle$ is given by $|\alpha|^2$, and the probability of obtaining $|1\rangle$ is given by $|\beta|^2$. This makes the choice of measurement basis meaningful, as a qubit will not generally give the same result if measured in a different basis.

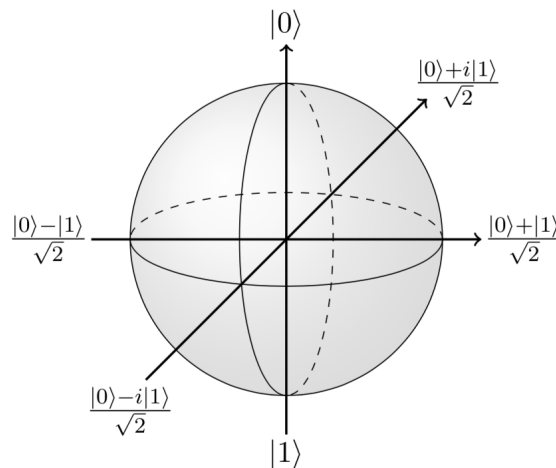


Figure 1.1: Bloch Sphere The Bloch sphere is often used to depict the state of a qubit. A qubit can take any state which lies on the surface of this sphere, which has a radius of unity. The basis states $|0\rangle$ and $|1\rangle$ lie at the north and south poles of the sphere. Therefore, a qubit pointing directly towards the north pole has the state $|0\rangle$. However, a qubit which points towards the equator of the sphere is in a superposition state containing equal amplitudes of both $|0\rangle$ and $|1\rangle$, as well as a phase, which is determined by its longitude around the sphere. Any arbitrary state a qubit can take can be represented as a vector within the Bloch sphere. [Image credit [Tow]]

A qubit is often depicted using the Bloch sphere, as seen in Fig 1.1 [Blo46]. In this picture, the north and south poles of the sphere correspond to the $|0\rangle$ and $|1\rangle$ states. The qubit can take any state which lies on the surface of the Bloch sphere. This gives rise to the picture of a complex qubit "phase", that being the "longitude" of the qubit vector in the Bloch sphere, which has little bearing on the state of a single qubit measured in the $|0\rangle, |1\rangle$, but becomes crucially important for two interacting qubits.

1.1.4 The DiVincenzo criteria

The physical realization of a quantum computer has some significant requirements, each one of which is not trivial to implement in a physical system. These were first laid out

by DiVincenzo in 2000 and have come to be known as the DiVincenzo requirements for quantum computation [DiV00].

1. A scalable physical system with well-characterized qubits

The first criterion is immediately a formidable challenge. To produce a quantum computer which can perform large-scale computations, we will require at least 10^6 qubits [Ste99]. At this scale, qubits cannot be individually tuned, and therefore reliable, reproducible qubits that can be produced at scale will be necessary. Well-characterized qubits have been developed, especially using superconducting qubits and ion traps. However, scalability of these systems can be a significant practical issue in terms of hardware, power dissipation at low temperatures, and fabrication scalability.

2. The ability to initialize the state of the qubits into a known fiducial state

The initialization of a qubit is crucial to computation. The initial state must be known for the outcome of an operation to have any meaning. In general, we aim to initialize in the $|0000\dots\rangle$ state in the $\{|0\rangle, |1\rangle\}$ basis. For many systems (such as spin-based qubits) this simply means allowing the qubit sufficient time to relax to the energetically favourable ground state. This relaxation can be accelerated through use of relaxation hotspots (specific configurations which allow rapid relaxation), or by using a system with a naturally short relaxation time. In others, the initialization must be dynamically controlled, such as by pre-polarizing a photon qubit, or using microwave pulses to induce rapid relaxation. In a real system, qubit errors mean that many repeated measurements will be necessary, meaning efficient initialization is highly beneficial [Tuo17].

3. Relevant decoherence times longer than the gate operation time

Qubits are naturally coupled to their environment. The quantum state of a qubit can be perturbed through interaction with the environment. For a single qubit state, this can lead to relaxation, for example, from the $|1\rangle$ state to the $|0\rangle$ state. For multiple qubits, however, the relative phase of the qubits can also decay due to interaction with the environment, leading to what is known as decoherence. To conduct quantum computation, the time required for a qubit to decohere must be much longer than the time required to implement computational operations on the qubit. This corresponds to a low noise environment, or a qubit implementation which is highly decoupled from the environment.

4. A universal set of quantum gates

The actual process of computation involves applying operations to a register of qubits. These operations are performed via a set of unitary transformations, referred to as quantum gates. Quantum gates are operations which rotate a single qubit, or a pair of qubits, about the Bloch sphere. An example of a simple quantum algorithm involving two-qubit gates is depicted in Fig 1.2. A universal set of gates means that any arbitrary operation can be carried out through combinations of one- or two-qubit gates. In general, one-qubit gates and two-qubit CNOT gates, or any set of gates that can reproduce these, are an adequate set of gates to enact universal operations [DiV00]. However, increasing the number of gates

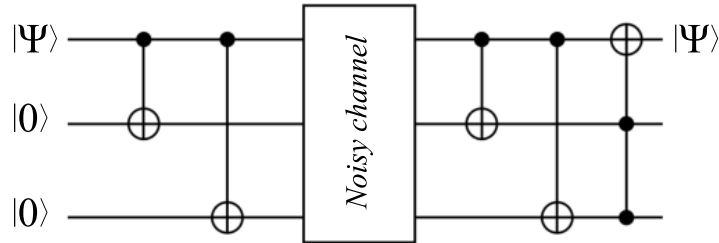


Figure 1.2: Quantum algorithm An example of a quantum algorithm. The horizontal lines represent qubits, with the x axis moving forward in time towards the right. A typical gate is the two-qubit CNOT gate, represented here by a white circle containing a "+" on the target qubit, with the control qubit indicated by a solid black circle. This specific algorithm is a simple error correction algorithm, in which a qubit and two ancilla are passed through a noisy channel and post-corrected to yield the original qubit state.

required to produce an arbitrary operation also increases the error in the final result, as each gate induces additional errors into the computation. Therefore it would be preferable to have a computational basis upon which a wide set of one- and two-qubit gates can be directly implemented.

5. The capability to measure specific qubits

The final requirement for general quantum computation is the ability to readout and measure the quantum state of individual qubits. For an ideal measurement, the qubit with state $|\Psi\rangle = \alpha|0\rangle + \beta|1\rangle$ should yield the outcome 0 with probability $|\alpha|^2$ and 1 with probability $|\beta|^2 = 1 - |\alpha|^2$, independent of any other parameters including the state of nearby qubits. An ideal measurement should also not perturb the rest of the quantum computer (beyond projecting the state of the measured qubit). Again, environmental interaction with the qubit means that 100% efficiency will never be reached in a physical system. This means that a real quantum computer will have a finite error inherent in the readout. This can be accounted for by duplicating the operation over many qubits and repeating the measurement. Such redundancy will be necessary for any physical implementation of a quantum computer.

In addition to these five criteria, there are two additional requirements, which do not directly pertain to quantum computation, but instead to the development of a connectable

set of quantum computers. A quantum computer could be built without these, but they would be necessary, for example, to transfer a qubit state over a very long distance. They are requirements for quantum communication and cryptography due to the necessity of a specifically quantum channel of communication in order to transfer quantum information.

6. *The ability to interconnect stationary and flying qubits*

Different qubit realizations have different advantages. Some, such as superconducting and quantum dot qubits, demonstrate high degrees of control and have a small physical footprint, allowing them to be scaled to large numbers. However, the interaction range of these qubits is very short, on the order of the size of a single qubit. Whilst some medium-range transport of solid-state qubits has been demonstrated (for example, using surface acoustic waves to coherently transmit electrons over several micrometres, or using spin chains to transfer a spin state through a series of coupled quantum dots), for very long range communication on the order of tens or hundreds of kilometres, it is likely that photons in fibre optic cables will be the go-to qubit. In order to transfer quantum information from a static computer across a communication channel, we require the ability to transfer information from one type of qubit to another. Such interactions can be realized via optomechanical resonators, or through photoexcitation in solid state materials...

7. *The ability to transmit flying qubits between specified locations*

The final requirement to complete a fully quantum-compatible network is the ability to transmit flying qubits over long distances whilst maintaining quantum coherence. Photon qubits in fibre optic cables will be the most likely realization of this, although decoherence over long distances may necessitate the use of "booster stations" - intermediary points along the network where operations can be performed to preserve and re-transmit the state. To enable such a network to connect multiple quantum computers, the decoherence time of the quantum state must be longer than the time required to transmit the state to another point.

1.1.5 Quantum error correction

Any real quantum computer that satisfies all of the first five of DiVincenzo's criteria will be comprised of qubits which have a finite error. This is unavoidable, and whilst individual qubit errors must be minimized, the field of quantum error correction has developed to realize a "fault tolerant" quantum computer. Quantum error correction generally involves an encoding of a "logical" or "computational" qubit onto many physical qubits to correct for bit-flip and phase-flip errors. This method requires additional quantum gates in an operation to encode the logical qubit. The threshold theorem for quantum computation is stated as follows: A circuit containing n qubits and $p(n)$ gates can be simulated with an error ε using a number of gates N on a register of qubits that produce an error with probability p , given that $p < p_{th}$ [Nie01].

$$N = O(\log^c(p(n)/\varepsilon)p(n)) \quad (1.1)$$

The quantum threshold theorem shows that if the individual gate error is a small

constant, arbitrarily long computations can be performed to arbitrarily good precision with only a minor overhead (on the order $\log(1/\epsilon)$). However, to implement this, the number of qubits required is expanded significantly. An example error correction algorithm is depicted in Fig 1.2, where three qubits, including two ancilla, are used to correct a single logical qubit. In large-scale quantum error correction, many individual physical qubits are required to encode a single logical qubit. One promising implementation of quantum error correction, known as surface codes, suggests that for a typical individual qubit error of 0.1% (a 99.9% fidelity qubit), one logical qubit will require between 1,000 to 10,000 physical qubits [Cam17]. This indicates the kind of numbers of qubits that will be necessary to create a functional quantum computer, and underpins the importance of scalability in the physical realization of a quantum computer.

1.2 Physical realizations of a quantum computer

Ultimately, there is no fundamental barrier to a fault tolerant quantum computer. The challenge to overcome is a technological one, and as research progresses with cutting-edge engineering techniques, it seems inevitable that a large-scale quantum computer will be possible. However, there are many different propositions for building a quantum computer, each of which has its advantages and disadvantages. Here we will outline some of the most developed qubit implementations.

1.2.1 Superconducting qubits

The first superconducting qubit was developed in 1999 [Nak99], not long after the first quantum computing architecture proposals. Since then, the field of superconducting quantum computing has rapidly developed, with one- and two-qubit gates demonstrated with fidelity over 99% [Bar14]. Superconducting qubits are the qubits of choice for current industrially-developed quantum computers, with Google and IBM leading the way with 53 and 65 qubit arrays respectively [Moo21]. Notably, it was the first qubit implementation in which the milestone of so-called "quantum supremacy" has been claimed to be achieved, using Google's 53-qubit processor (in 2019, [Aru19]).

There are many different designs of superconducting qubits, including charge qubits, flux qubits, and phase qubits [Hua20]. They are versatile, and can operate in different modes depending on the requirements of the system. Additionally, they are built using microfabrication processes, meaning several qubits can be placed on a single chip, offering good short-term scalability. Coupling between superconducting qubits is relatively simple, and usually implemented via capacitive or inductive coupling. Finally, superconducting qubits can be controlled via microwave excitation, which can be reliably implemented using current technology. With this in mind, it is easy to see why the field has had significant success over the past decade.

1.2.2 Photonic qubits

Photonic qubits have been part of the quantum computing conversation since the earliest days for two main reasons [Slu19]. Firstly is the simplicity in which high-fidelity single-qubit operations can be performed with minimal environmental interaction. Secondly is the necessity of transporting quantum information over large distances, for which photonic

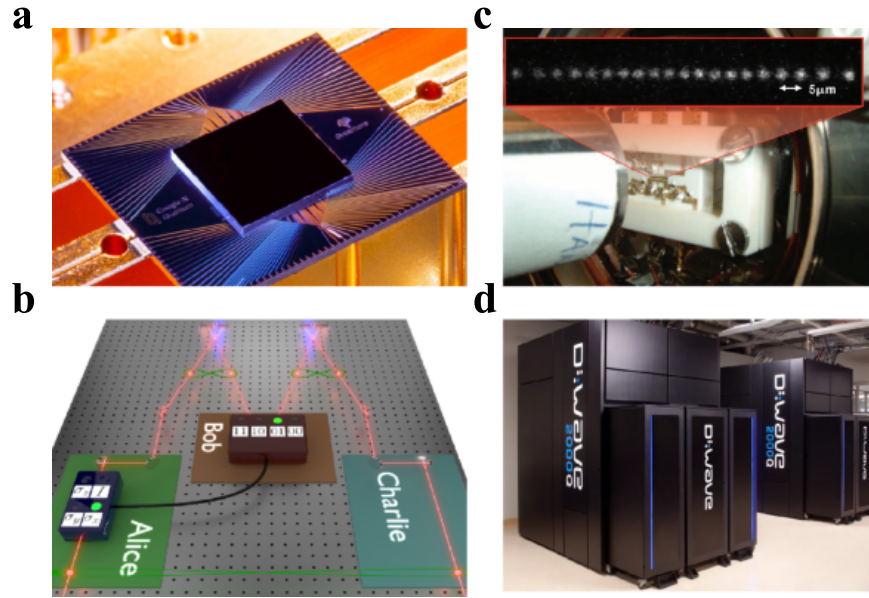


Figure 1.3: Qubit realizations a) The Google Sycamore quantum chip, used to demonstrate quantum supremacy in 2019 and currently one of the most advanced quantum computers. [Image credit [Aru19]]. b) Photonic quantum computers make use of interferometry and single-photon emission and detection to implement flying qubits, [Image credit [Fla18]] c) Trapped ion quantum computers use magnetic fields to trap single ions, which can be manipulated with lasers, [Image credit [Mon13]] d) The D-Wave 2000Q quantum computer contains 2048 qubits, which are used to implement quantum annealing - a sub [Image credit ©D-Wave Systems Inc.]

qubits are an obvious candidate. Photons are minimally-interacting systems, so that they can transport quantum information over very long distances without decoherence.

In their current state, photonic qubits fulfil five of the seven DiVincenzo requirements, partly due to their natural advantage in the two communication-based criteria. Difficulties arise when considering the last two criteria, however: implementing a universal set of quantum gates, and qubit-specific measurement. By their nature, photons have minimal interaction. This gives them long decoherence times due to low interaction with the environment.

Photonic qubits have already entered the market, with applications in cryptography and random number generation. Commercial quantum key distribution systems are commercially available, with companies such as MagiQ Technologies, QNu Labs, Quintessence Labs, SeQureNet and ID Quantique offering quantum key distribution systems based on

quantum optics. Quantum random number generators are commercially available using quantum optics, such as the Quantis product line from ID Quantique. With their prevalence in commercial quantum applications and potential for enacting flying qubits, photonic quantum computation is one of the deepest areas of research in the field.

1.2.3 Trapped ion qubits

No discussion of quantum computing implementations would be complete without mentioning qubits formed of cold trapped ions. Some of the first propositions of implementing quantum computations, predating even the Kane proposal, involve the use of trapped ions as qubits [Cir95]. Single trapped ions have been confined in radiofrequency traps since the 1980s, and displayed promising properties for quantum computation. The ions have long lifetimes within the traps, long coherence times of their internal basis states (up to 50 s [Bru19]), and can be manipulated with laser pulses and laser cooling. Additionally, two-qubit gates were rapidly demonstrated using the hyperfine coupling in a Beryllium ion [Mon95].

Since these first experiments, several of the DiVincenzo criteria have been satisfied for trapped ion qubits. Single-qubit gates with errors below 10^{-4} have been demonstrated on single ions [Bro11], as well as two-qubit gates shown to be able to entangle multiple ions [Ben08]. High-fidelity state preparation and measurement have also been demonstrated, above the limit for fault-tolerant quantum computation [Mye08]. Trapped ions fulfil the first five DiVincenzo criteria with remarkably high fidelity, and have even been shown to be able to implement Shor's algorithm [Mon16].

1.2.4 Other realizations

In addition to the above qubit designs, there are many other types of qubits or quantum computers, most of which fulfil a specific purpose or play a particular role. Many of these have exciting prospects for quantum computation which are more specialized than, for example, superconducting qubits.

One approach to quantum computation, especially that involving quantum simulation of molecules, is liquid nuclear magnetic resonance (NMR) quantum computation. Liquid NMR is based on molecules in solution that contain atomic nuclei which act as spin-1/2 systems, each of which has a characteristic energy and interaction with its neighbours [Ger98]. The qubit is formed by the spin-1/2-like nuclei, and the qubit-qubit interactions are mediated by spin-spin interactions at inter-atomic bonds. Liquid NMR is performed on a large ensemble of molecules, which has both advantages (such as being able to work at higher temperature) and disadvantages (unable to address specific qubits or sets of qubits, introduction of decoherence). Shor's algorithm was demonstrated on a 7-qubit ensemble liquid NMR quantum computer in 2001 [Van01]. Solid-state NMR operates on similar principles, but with solid state samples, allowing lower temperatures, individual qubit addressability, and a higher degree of control [Jon01].

Nitrogen-vacancy (N-V) centers in diamond are another qubit implementation of interest due to its atom-like properties and solid state environment [Chi13]. They contain spin degrees of freedom which can be addressed via optical transitions. Additionally, stimulated single photon emission has been demonstrated in N-V centers, making them of particular

interest when considering the sixth DiVincenzo requirement of interacting solid state qubits and flying qubits [Jes17].

Finally, a prominent type of quantum computer is the adiabatic quantum computer, particularly that based on quantum annealing. This is the basis of the D-Wave Systems commercial quantum computers. It is not a universal quantum computer, but specialized in solving minimization problems - that is, problems which involve finding a global minimum in a finite parameter space. It cannot, for example, solve Shor's algorithm, as it does not fall into this class of problems. However, quantum entanglement has been demonstrated on their quantum annealers [Lan14], and a speedup over classical methods such as simulated annealing and quantum Monte Carlo simulation was shown in 2015 [Den16]. Their largest quantum computers have over five thousand qubits operating with more than a million Josephson junctions, outpacing the current universal quantum computer implementations.

1.3 A spin-based silicon quantum computer

Spins in silicon have been central to the quantum computing conversation since the first proposals. The archetypal model for a quantum computer, proposed by Kane in 1998 [Kan98], involved the use of nuclear spins in phosphorous atoms embedded in a silicon lattice, controlled by so-called "J" and "A" gates, which tune the coupling between qubits and the qubit state itself respectively. It proposed to leverage the quiet nuclear spin field to have a minimal-decoherence substrate, and the hyperfine interaction between the electronic and nuclear spin to couple qubits.

Also in 1998, Loss and DiVincenzo proposed a quantum computer using electron spins in quantum dots [Los98]. Quantum dots in semiconductor materials have emerged as a method of tightly confining electrons [Cro97; Gam96; Kum90; Liv96; Swi98]. Loss and DiVincenzo proposed to use quantum dots to trap single electrons in a two-dimensional array. The electrons would then be controlled through manipulation of the wavefunction by gate voltages to tune the coupling and chemical potentials - the "A" and "J" gates.

Silicon - and semiconductor materials in general - was chosen for these early proposals due to its negligible hyperfine field after purification, as well as the history of silicon in industrial transistor fabrication. The electron and nuclear spins are both strong candidates for qubits, forming a natural two-level system which can be accessed and manipulated using electric and magnetic fields.

Qubits have been enacted in many types of semiconductor quantum dots, from InAs nanowires to GaAs heterostructures. These substrates are excellent test beds, and have been used to demonstrate complex electron manipulation, such as coherent electron spin transfer using surface acoustic waves [Jad21] and shuttling of electrons in a two-dimensional array of quantum dots [Mor18]. Silicon spin qubits in particular have been the subject of rapid advancements in recent years [Hut16; Pla12; Tyr12; Vel17; Yon17]. Their potential for scalable development via industrial fabrication processes is especially exciting, and a great deal of effort is now focussed towards demonstrating that silicon spin qubits - and CMOS qubits in particular - can fulfil the DiVincenzo requirements. Due to the simple initialization of spins (through fast ground state relaxation) and the long coherence times characteristic in silicon, the main three objectives to be demonstrated are addressable,

high fidelity readout; single-qubit manipulation; and two-qubit manipulation.

1.3.1 Si-MOS spin qubits

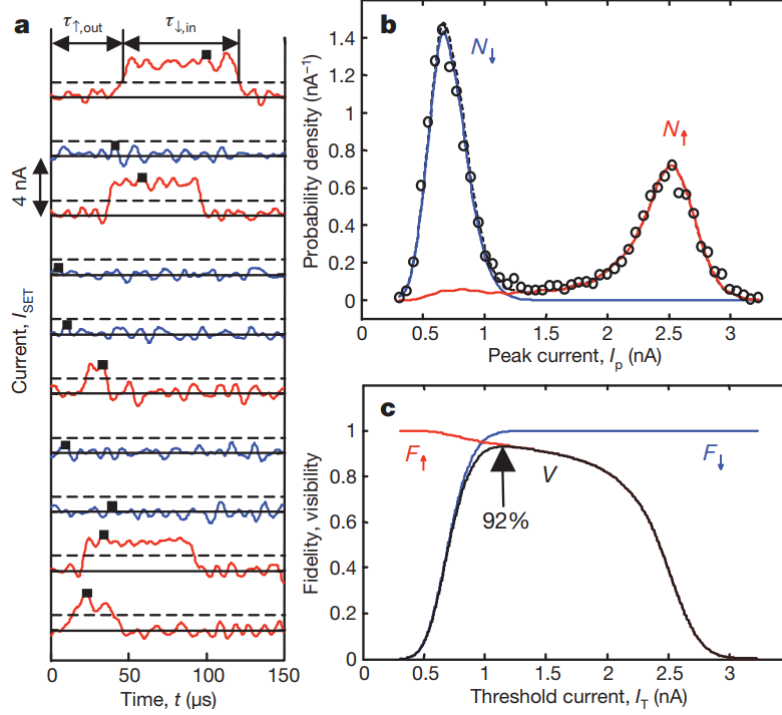


Figure 1.4: High-fidelity readout The energy-selective readout used to detect the spin of a single electron in a phosphorous donor is depicted. a) Single-shot readout traces distinguishing between spin-up (red) and spin-down (blue) measurements. If the current jumps beyond the threshold (black dashed line) within the measurement time, the spin is detected to be up, otherwise it is down. b) Histogram of the current distribution in the spin-up (red) and spin-down (blue) measurements. The overlap of the two peaks gives the readout error - this represents a false detection (detecting spin-up as spin-down, and vice versa). The more separated the peaks are, the higher fidelity the measurement is. c) As the threshold current (black dashed line in a)) is varied, the fidelity of measuring spin-down (spin-up) is increased (decreased). The visibility $V = F_{\uparrow}F_{\downarrow}$ is plotted in black. At the optimum threshold current, the maximum visibility is 92%. Adapted from [Mor10].

The first objective of addressable, high-fidelity readout was achieved in 2010 [Mor10]. Energy-selective readout using an on-chip charge sensor was used to detect a single electron spin, with single-state fidelity approaching 99%, and state visibility around 92%. The readout is depicted in Fig 1.4. Such single-shot spin detection is an important step towards addressable spin readout. Demonstration of this in principle suggests that any arbitrary qubit in a quantum computer could be read out in a single measurement, with the only geometric requirements being a nearby charge sensor and an electron reservoir. Energy-selective readout has proven to be a reliable and high-fidelity choice for spin readout.

High-fidelity single qubit gates have been demonstrated with 99.9% fidelity [Yon17]. The gate is implemented via EDSR spin control with microwave pulses on an electrostatic gate,

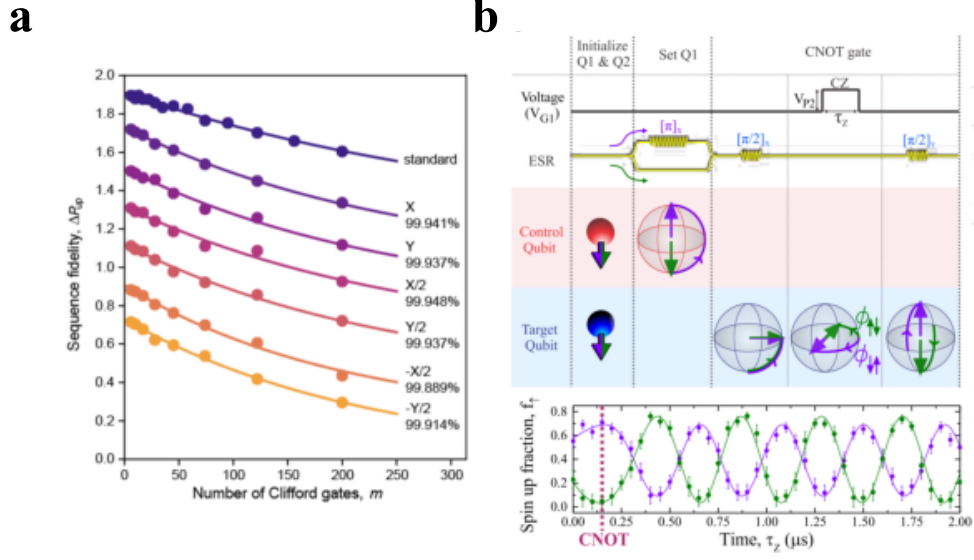


Figure 1.5: Single- and two-qubit gates Demonstrations of one- and two-qubit gates in silicon quantum dots. a) Fidelity analysis of a one-qubit gate in a silicon quantum dot [Yon17]. The gate fidelity for N iterations of different Clifford gates. The Clifford gates correspond to the Pauli transformations and the identity transformation, comprising a complete set of single-qubit transformations which, in conjunction with a two-qubit CNOT gate, could be used to implement any arbitrary quantum gate. The fidelity of each gate is above or close to 99.9%, above the 1% error threshold for error-corrected fault tolerant quantum computation. b) Schematic and readout of a two-qubit CNOT gate [Vel15]. A target qubit undergoes a π rotation conditional on the state of the control qubit. Due to the exchange interaction, the frequency of the microwave pulse required is different depending on the state of the control qubit. It is set to one of these state-specific frequencies, then pulsed when the control qubit is in a known initial state. A CNOT gate with a fidelity greater than 99% was demonstrated with this method. Adapted from [Vel15; Yon17].

manipulating an electron spin in a magnetic field gradient and detected with energy-selective readout. The fidelity is quantified via repeated EDSR pulses of varying length, with the resulting spin-up population measured and compared to the expected population. It is benchmarked by applying Clifford gates, common single-spin operations which correspond to the Pauli matrices; that is, π -rotations about the cardinal axes of the Bloch sphere. The fidelity as a function of the number of applied gates is depicted in Fig 1.7, with a fidelity of more than 99.9% found for almost all types of single-qubit gate.

Two-qubit gates are the most complex to implement, requiring excellent control over individual qubits and the inter-qubit coupling. However, they are necessary to be able to implement a universal set of quantum gates. It is sufficient to be able to demonstrate a

CNOT gate (controlled-NOT, whereby one qubit undergoes a π rotation dependent on the state of the second), as any arbitrary gate can be implemented as a combination of CNOT and single-qubit gates. A two-qubit CNOT gate has been demonstrated in silicon [Vel15]. Two quantum dot qubits are coupled via the exchange interaction, which is controlled with the detuning. A schematic of the CNOT procedure is laid out in Fig ???. The averaged state readout depicted in (b) demonstrates coherent interaction during the CNOT gate. The fidelity of the two-qubit gate was determined to be above 99%, which is compatible with error correction codes which allow fault-tolerant computing with errors as high as 1%.

With these three main objectives demonstrated for Si-MOS qubits, and the inherent benefits of silicon, the future looks promising for development of a universal quantum computer with this qubit realization. The success of Si-MOS qubits has motivated research into a similar kind of qubits with formidable potential for scalability: CMOS spin qubits.

1.3.2 CMOS spin qubits

CMOS fabrication processes are used to fabricate high-density processors containing billions of transistors. These processes could be leveraged to allow mass fabrication of CMOS spin qubits, combining the promising properties and demonstrated capabilities of Si-MOS qubits, and the scalability and reproducibility of CMOS fabrication. Despite being a relatively young field, significant progress has been made towards demonstrating the required objectives. The MOS-Qubit project, consisting of a collaboration between the nanoelectronics giant CEA-LETI and various academic institutions, is a European project to create a high-fidelity qubit on a 300 mm CMOS platform. Whilst the ultimate objective of a fully-functional qubit with two-qubit gates has not yet been achieved, several advances have been made in this direction.

Maurand et al in 2016 demonstrated qubit functionality in a hole spin qubit using a CMOS device [Mau16]. It was demonstrated that quantum dots could be formed in the corner of a depleted silicon channel using only a single control gate, with two quantum dots in series formed. Spin readout was demonstrated via Pauli spin blockade, with the channel biased so that holes could flow in one direction only. Importantly, single-spin manipulation was demonstrated by measuring the spin as a function of the power and burst time of a microwave pulse applied to one of the gates. The crucial elements of spin readout and single spin manipulation were demonstrated, and these results - as well as the promising results from Si-MOS qubits - spurred further research into these CMOS-fabricated devices.

One of the drawbacks - which can also be an advantage - of electrons in silicon is their low spin-orbit coupling. This means that electronic spins can only weakly be controlled by electric fields, presenting a challenge to performing EDSR manipulation. However, it also means that a major source of spin relaxation is minimized for electrons. EDSR was however demonstrated for electrons, by making use of spin-valley mixing to enhance the spin-orbit coupling [Cor18]. This allows use of low power microwave excitation which would normally not be sufficient to drive conventional ESR. EDSR is a promising tool for selective spin manipulation, as it can be applied on a local gate rather than using a microwave stripline for magnetic field control.

Single-shot spin readout is necessary to be able to detect the spin of specific qubits - one of the main requirements of the DiVincenzo criteria. RF reflectometry has emerged as

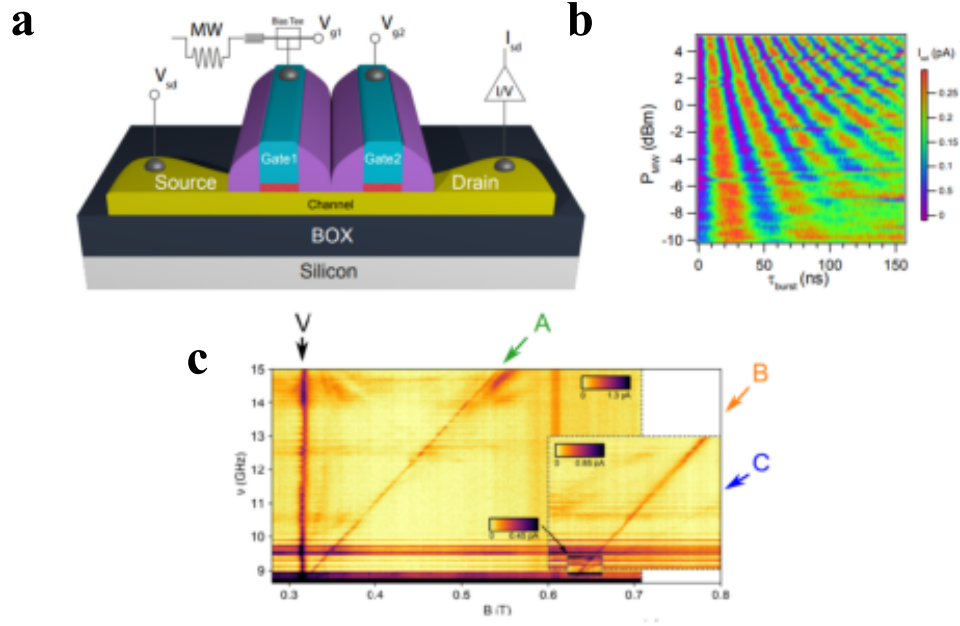


Figure 1.6: CMOS qubit manipulation Demonstration of single-qubit manipulation in CMOS silicon quantum dots. a) Schematic of a typical two-gate CMOS-fabricated device. The depleted channel (yellow) is bounded by two reservoirs, source and drain. The spin is read out via Pauli spin blockade using two quantum dots in series below the two gates, with a dot containing a ground-state spin under gate 2 blocking transport of a ground-state spin from the dot under gate 1. b) Rabi oscillations as a function of microwave power and burst time. The microwave burst is applied on gate 1, rotating the blocked spin into a state dependent on the duration and power of the burst. c) Demonstration of EDSR control of an electron spin in a similar device to that depicted in a). The EDSR is mediated via the spin-valley interaction, and controlled with a frequency matching E_Z , which is variant with magnetic field. The resonance lines A,B and C correspond to different resonant spin-flip transitions. Adapted from [Cor18; Mau16].

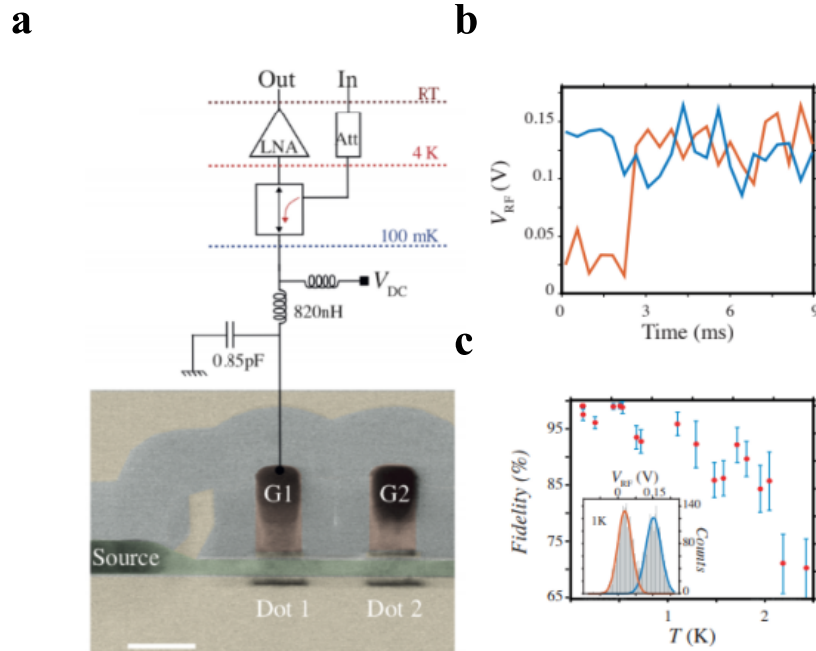


Figure 1.7: Single-shot readout Single-shot readout demonstrated in a CMOS device using RF reflectometry. a) Schematic of the device and RF circuit. A tank circuit is connected to the quantum dot under $G1$, with the reflected signal used to probe the charge state of the dot under $G2$. b) Single-shot readout of the charge state of $G2$. The distinctive spin-down (orange) and spin-up (blue) traces are used to distinguish between the two states. c) Fidelity analysis of this readout method. The fidelity is degraded with temperature due to thermal charge transitions, but remains above 90% up to a temperature of 1 K. At low temperature, the readout fidelity approaches 99%. Adapted from [Urd19].

a method of implementing high-fidelity single-shot readout integrated within the device [Urd19]. This is a method of charge sensing which uses a reflectometry circuit connected to a gate or reservoir nearby the probed system. A resonance in the reflected signal is shifted in frequency when the capacitive coupling between the gate (or reservoir) and the probed system is changed by a charge entering or leaving the probed quantum dot. This integrated charge sensing was used with an electron latching mechanism to detect the spin of a double quantum dot in single-shot. The readout fidelity was shown to remain above 90% up to 1 K, indicating that spin readout can be possible at relatively high temperatures.

With single-spin manipulation and high fidelity single-shot readout achieved in CMOS quantum dots, only a few milestones remain to be demonstrated to have functional, coupled spin qubits. However, this field remains very young, and despite the current promising

results, CMOS quantum dots have not been well-characterized, and their properties are not fully understood. In particular, experimental results do not always line up with simulations of the quantum dot behaviour, and the reasons why are currently unknown. A more complete characterization of the properties of such CMOS devices is necessary to be able to proceed towards functional qubits, and this is the problem that this thesis aims to tackle.

1.4 Thesis structure

Characterization of CMOS devices at room temperature is a routine part of fabrication. However, low-temperature characterization of the quantum dots is not yet routine, and is necessary to inform simulations and future development in device design, as well as obtaining data on the variability in low-temperature characteristics across different devices. This thesis aims to demonstrate in-depth low-temperature characterization of quantum dots in a CMOS device in terms of the local electric field fluctuations and disorder, as well as the spin and valley physics which are highly influenced by the local device characteristics.

In **Chapter 2** the theory and fundamentals behind the formation of single and double quantum dots is outlined. The electrostatic model and the energetic structure of the levels within the dot are detailed. We discuss the principles behind charge sensing, and the spin structure within the dot, with several spin-to-charge conversion mechanisms presented.

In **Chapter 3** we outline the experimental setup used to characterize a CMOS device. We present the device design and structure, and discuss how it can be characterized at room and low temperature. Finally, we demonstrate how various parameters of a quantum dot can be extracted using current-voltage measurements, with example measurements from real devices.

In **Chapter 4** we outline how charge noise affects a qubit defined in a semiconductor device, the physical source of the noise, and how a $1/f$ noise spectrum can be interpreted for semiconductor quantum dots. We then present experimental results, demonstrating clear $1/f$ charge noise experienced by a quantum dot in a silicon CMOS nanowire device. We demonstrate how the charge noise varies as the quantum dot is manipulated within the device and experiences different spatially separated sources of charge noise. We also explore the sources of charge noise and measure the energies of two-level fluctuators believed to be the dominant cause within the measurement regime. Finally, we present a novel method of measurement of charge noise at the single electron level, and use it to probe the charge noise experienced by the first three electrons in a quantum dot.

In **Chapter 5** we present a method of detection and addressable single-shot measurement of a single electron spin, which could be used to characterise large arrays of silicon CMOS nanowire quantum dots. We first present the readout method used and fidelity analysis, and demonstrate the measurement of the spin relaxation time T_1 with a spin state visibility greater than 90%. The relaxation time is analysed as a function of the magnetic field, and used to detect the spin-valley relaxation hotspot at $(297 \pm 5) \mu\text{eV}$. We demonstrate control over the valley splitting via electric field tuning from $297 \mu\text{eV}$ to $260 \mu\text{eV}$, detected via measurement of the relaxation hotspot. Finally, we measure the magnetic field anisotropy of the spin-valley mixing, demonstrating suppression of the relaxation mechanism in a field oriented along the main symmetry axis of the device.

CHAPTER 2

Quantum Dots

2.1 Introduction

A promising candidate for a physical implementation of a scalable qubit is the quantum dot. A quantum dot is a sharp potential well, defined on the nanoscale by some mechanism of potential confinement; this can be through the use of electrostatic gates, or by using a naturally-occurring potential well, such as a single dopant.

Here we review the dynamics of a quantum dot in a semiconductor, defined via electrostatic gates acting on a 2D electron gas. First we present the theory behind the formation and operation of a quantum dot, then discuss the use of a quantum dot as an electrometer for probing the single-electron regime. Finally we present the use of single and two electron spins as a two-level system appropriate for use as a qubit.

2.2 Quantum dots

2.2.1 Single dot energy

A quantum dot is a sharp potential well that is able to trap and hold a number of electrons (or holes). Quantum dots in semiconductors can be formed by accumulating electrons into what is known as a 2D electron gas, or 2DEG, at a physical interface in a semiconductor device through attraction by a positive electric field. The dot is then defined in two dimensions via electrostatic fields generated by metallic gates, isolating a small region of this 2DEG.

The size of the dot is typically a few to tens of nanometers, to be on the order of the fermi wavelength of the electrons. In isolation, the confinement in all dimensions results in a quantized energy spectrum, and therefore a quantum dot holds an integer number of electrons that occupy the available energy states below the Fermi energy. The number of free electrons in the dot can be changed by manipulating the confinement potential.

As electrons are loaded onto the quantum dot, they exert a repulsive Coulomb force upon other electrons in the dot. In order to add an additional electron to the dot, this repulsive force must be overcome. The energy associated with this is called the charging energy, and is a purely classical effect. The charging energy can be derived simply from electrostatics. If we consider a quantum dot as one terminal of a capacitor, with the other

being an electron reservoir, then the charge on the quantum dot will be given by $Q = CV$, with C the capacitance and V the potential difference between the dot and the reservoir. The energy on the quantum dot is therefore:

$$E = \frac{1}{2} \frac{Q^2}{C} = \frac{(Ne)^2}{2C} \quad (2.1)$$

Where N is the number of electrons trapped on the quantum dot. For the case of $N = 1$:

$$E = \frac{e^2}{2C} \quad (2.2)$$

This is called the Coulomb charging energy, or simply charging energy, and is an important energy scale for quantum dots. For semiconductor quantum dots on the order of 100 nm in size, the charging energy is typically on the order of a few meV.

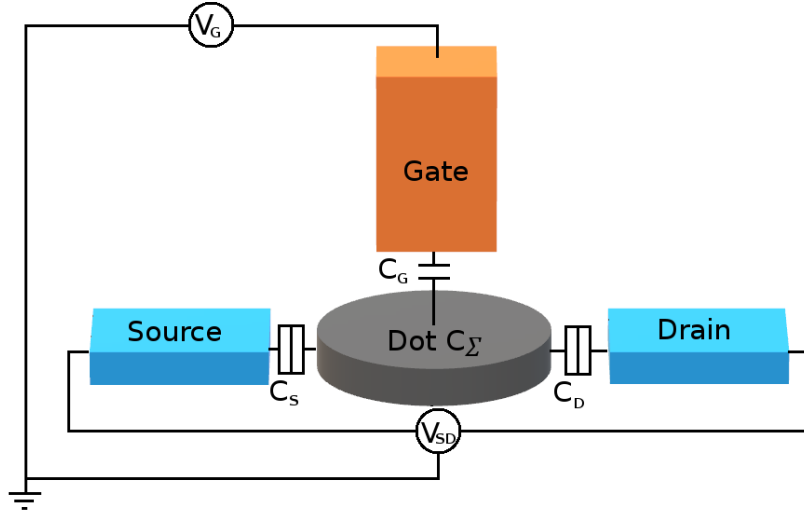


Figure 2.1: Quantum Dot model Schematic of the electrode arrangement for a single quantum dot. The dot has a self-capacitance $C_\Sigma = C_G + C_S + C_D$. It is capacitively coupled to a gate with an applied potential V_G with a capacitance C_G . The dot is also capacitively coupled to two electron reservoirs, labelled here Source and Drain, with capacitance C_S and C_D respectively. A bias potential V_{SD} may be applied across the reservoirs.

We consider a quantum dot using the constant interaction model, which assumes that the Coulomb interaction between electrons within the dot is independent of the number of electrons N . In the CI model, the quantum dot is an island tunnel coupled to two electron reservoirs, labelled here Source and Drain (see Fig. 2.1), and capacitively coupled to a plunger gate with an applied potential V_G that directly controls the chemical potential of the dot.

The tunnel barriers separating the dot from the reservoirs are modelled as leaky capacitors that can allow charges to flow from reservoir to dot and vice versa. A leaky capacitor can be considered as an ideal capacitor in parallel with a resistor. When a voltage is applied

across the capacitor and resistor, a current can flow across the resistor. This system is called a tunnel junction or tunnel barrier, and is used to describe the coupling between the quantum dot and the source and drain reservoirs.

The electrostatic energy of the dot is given by the energy stored inside an ideal capacitor:

$$U = \frac{Q^2}{2C_\Sigma} \quad (2.3)$$

Where $C_\Sigma = C_G + C_S + C_D$, the sum of the relevant capacitance contributions, and Q is the charge inside the dot. The charge inside the dot is the sum of the number of free electrons on the quantum dot N multiplied by the electronic charge e , and the charge induced by each contact, $C_i V_i$ where C_i is the capacitance and V_i the voltage on the contact i .

$$Q = -N|e| + C_G V_G + C_S V_S + C_D V_D \quad (2.4)$$

For now we can consider the source and drain to be grounded such that $V_S = V_D = 0$, and therefore the charge on the quantum dot is controlled primarily by the gate voltage V_G . This gives the electrostatic energy of the quantum dot to be:

$$U(N) = \frac{(-N|e| + C_G V_G)^2}{2C_\Sigma} \quad (2.5)$$

Quantum dots are often referred to as "artificial atoms", as they exhibit similar properties [Kou01]. They are defined by a tightly confined positive potential, which in atoms is the nucleus, and for quantum dots is gate-defined. Both contain an integer number of electrons. In atoms, the number of electrons can only vary slightly via ionization, whereas in quantum dots, we can tune the electron number over a wide range. And, similar to an atom, a quantum dot contains discrete energy levels than an electron can occupy, which can be considered equivalent to atomic orbital energy levels. Due to the Pauli exclusion principle, each orbital level can only contain two electrons, each of which must have opposite spin.

This manifests as an extra energy term E_{orb} , describing the single-particle orbital energy that must be paid when loading an extra electron into the dot. The splitting of these orbital levels in a 2D circular quantum dot of diameter L in a 2D electron gas is closely approximated by the solution to the Schrödinger equation for a particle in a two dimensional box:

$$\begin{aligned} \Delta E_{orb} &= E_{N+1} - E_N \\ &= \frac{1}{m_e} \left(\frac{\hbar\pi}{L} \right)^2 \end{aligned} \quad (2.6)$$

Where m_e is the electron mass. Due to the "filling" of orbital levels with opposite spin electrons, not every electron will have to pay this additional orbital energy. Thus we include a generic $E_{N-1,N}$, which is the specific energy requirement to add the N^{th} electron. In some cases $E_{N-1,N}$ will be 0.

$$U(N) = \frac{(-N|e| + C_G V_G)^2}{2C_\Sigma} + \sum_{N=0}^N E_{N-1,N} \quad (2.7)$$

To avoid complication via varying numbers of electrons in the dot, it is preferable to use the electrochemical potential over the energy of the dot. This is defined as $\mu_{dot}(N)$, the energy required to add the N th electron to the island:

$$\mu_{dot}(N) = U(N) - U(N-1) \quad (2.8)$$

$$\mu_{dot}(N) = U(N) - U(N-1) = (N - \frac{1}{2})E_C + \frac{C_G V_G}{C_\Sigma} + E_{N-1,N} \quad (2.9)$$

Where $E_C = \frac{e^2}{C_\Sigma}$ is the charging energy of the system and $\frac{C_G}{C_\Sigma}$ is the conversion factor for gate voltage to energy. This parameter is also known as the lever arm or alpha factor. The alpha factor is always positive, and typically varies across large ranges of gate voltage due to the increase in quantum dot size and resulting changes to the capacitance matrix. The charging energy is the difference between electrostatic energies only, not including the single particle energies.

From here we can define the addition energy E_{add} , which is simply the energy required to add one additional electron to the island:

$$E_{add}(N) = E_C + E_{N-1,N}(N) \quad (2.10)$$

The addition energies as a function of N form an addition energy spectrum. When a quantum dot is filled with N electrons, no more electrons can enter the dot until the addition energy has been paid.

2.2.2 Coulomb blockade

The state with N electrons has an electrochemical potential $\mu_{dot}(N)$. We define the potential of the source and drain as μ_S and μ_D respectively, as shown in Fig. 2.2. At $V_{SD} = 0$, $\mu_S = \mu_D$. By applying a finite V_{SD} , we open a potential window where $\mu_S = \mu_D + eV_{SD}$. Now if $\mu_{dot}(N)$ is positioned between the source and drain potentials, fulfilling $\mu_S \geq \mu_{dot}(N) \geq \mu_D$, then electrons are able to tunnel into and out of the quantum dot. When there is no available potential level within this window, transport is blocked, as all available energy states are filled, and the system is considered to be in Coulomb blockade.

In Fig. 2.2a, the quantum states in the dot are filled up to $\mu_{dot}(N)$. The addition of one extra electron would require shifting the potential of the dot to $\mu_{dot}(N+1)$, which lies above the bias window. As such, electron tunnelling into the dot is forbidden, and the device is in Coulomb blockade. In Fig. 2.2b, transport into the dot is allowed, as the potential level $\mu_{dot}(N)$ lies within the bias window. The dot is filled with $N-1$ electrons, but tunnelling in and out is freely allowed through the available level, causing the number of electrons on the dot at any one time to fluctuate between $N-1$ and N .

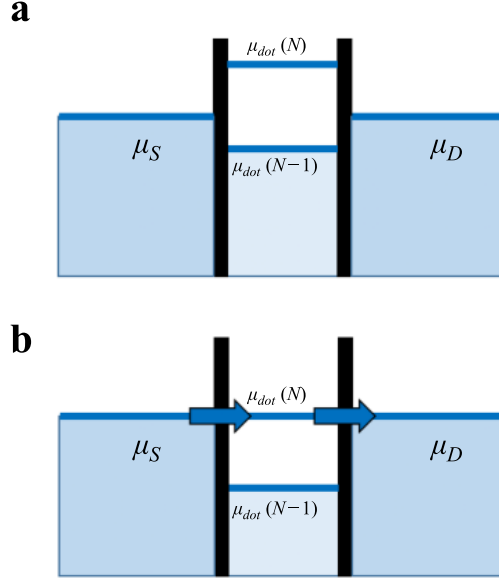


Figure 2.2: Coulomb blockade model Energy diagram for the quantum dot seen in Fig. 2.1. The electron reservoirs source and drain here have electrochemical potentials μ_S and μ_D respectively. These are related via the external voltage $\mu_S - \mu_D = eV_{SD}$. In (a), the transport is blocked and the dot contains a fixed number of electrons. In (b), the gate voltage is tuned into a configuration whereby current can flow through the dot via the state with potential $\mu_{dot}(N)$, leading to a peak in the conductance of the dot.

For a small $V_{SD} \simeq 0$ (referred to as the linear transport regime), we therefore have a peak in the conductance whenever the condition $\mu_S \geq \mu_{dot}(N) \geq \mu_D$ is satisfied. The gate voltage can be continuously changed, leading to peaks and blockade in the conductance through the dot, as shown in Fig. 2.3. The charge on the quantum dot is fixed when the system is in blockade, and when set to a peak, the number of electrons fluctuates between $N - 1$ and N .

The potential level $\mu_{dot}(N)$ is physically manipulated by changing V_G , which induces a potential shift at the quantum dot equal to $V_G \alpha_G$. At $V_{SD} \simeq 0$, it follows that the gate voltage required to position $\mu_{dot}(N)$ in conduction is at $\mu_S = \mu_{dot}(N) = \mu_D = 0$ (we define $\mu_D = 0$), which for the N th electron is given by:

$$V_G^{(N)} = \frac{1}{e\alpha_G} [E_{N-1,N} + (N - \frac{1}{2})E_C - e\alpha_G V_G] \quad (2.11)$$

The conductance of the dot can be determined for varying V_G by measuring the current

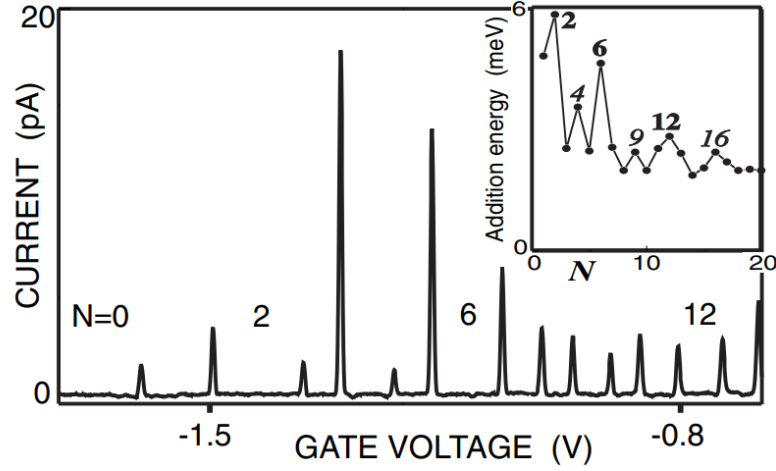


Figure 2.3: Coulomb peaks A typical coulomb blockade spectrum in the few-electron regime. Here, peaks can be seen at a spacing corresponding to potential levels residing within the bias window. The width of each peak is determined by the electron temperature and the applied bias, with the bias dominating at low temperature. The height is independent of temperature, and depends on the coupling of the quantum dot to the reservoirs, which can be strongly dependent on the coupled states. Inset: The addition energy spectrum, extracted from the peak spacing and charging energy, is plotted as a function of electron number. [Adapted from [Kou01], original data from [Tar96]]

at the drain contact. This will result in a series of conductance peaks, known as Coulomb peaks, as indicated in Fig. 2.3. Consecutive peaks are separated by a voltage proportional to the addition energy E_{add} . In the inset of Fig. 2.3, the addition energy spectrum, extracted from the peak spacing through subtraction of the charging energy, is plotted. Atoms have a three-dimensional shell and subshell structure that has orbital transitions at $N = 2, 10, 18, \dots$. The shells and subshells correspond to the lowest energy states available to a given number of electrons in the atom. A quantum dot has a different structure due to its 2D symmetry, and therefore has the shell structure of a 2D harmonic potential. This leads to shell filling at values of $N = 2, 6, 12, \dots$, which can be seen in the addition energy spectrum.

To properly analyse the spectrum, it is necessary to extract the lever arm of the tuning gate. The easiest way to determine this is through measurement of coulomb blockade diamonds by varying the bias V_{SD} as a function of gate voltage. Increasing V_{SD} opens the bias window further, narrowing the energy range in which the quantum dot is in blockade. This manifests as a widening of the coulomb peak. The bias can be increased until it exceeds E_{add} , in which case a potential level within the quantum dot is always positioned within the bias window, allowing conduction at any V_G above the first electron state.

This principle is demonstrated in Fig. 2.4. The resulting structure is often referred to as a Coulomb diamond due to its characteristic shape, where the shaded regions indicate where the quantum dot is in Coulomb blockade. Where the edges of the diamonds intersect at $V_{SD} = 0$, the linear-response coulomb peaks are seen. For simplicity, we consider the

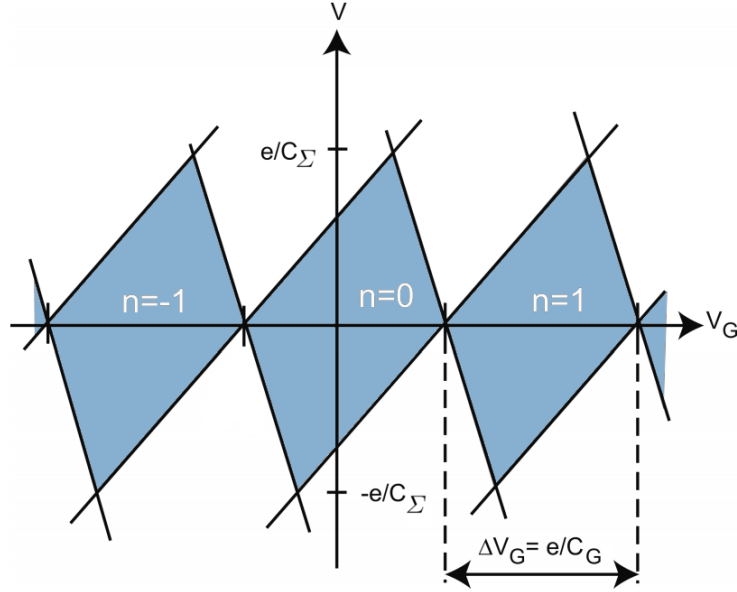


Figure 2.4: Coulomb diamonds A coulomb map of a single-electron transistor. Within the shaded diamonds the dot is in coulomb blockade, and current cannot flow. Outside of the diamonds, the coulomb blockade is overcome and current can flow from source to drain. The slopes of the boundaries are dependent on the coupling to each reservoir, as described in Eqn. 2.16 and Eqn. 2.17. [Adapted from [Hei03]]

bias V_{SD} to be applied symmetrically to the source and drain, such that $\mu_S = \frac{eV_{SD}}{2}$ and $\mu_D = \frac{eV_{SD}}{2}$. Then the stability requirements at positive bias voltage become:

$$\mu_{dot}(N) < \frac{eV_{bias}}{2} \quad (2.12)$$

$$\mu_{dot}(N+1) > \frac{eV_{bias}}{2} \quad (2.13)$$

and for negative bias:

$$\mu_{dot}(N) < \frac{eV_{bias}}{2} \quad (2.14)$$

$$\mu_{dot}(N+1) > \frac{eV_{bias}}{2} \quad (2.15)$$

The blockade region as a function of the bias voltage is defined by two lines that correspond to the energy conditions being met to have a potential level within the bias window. The relative coupling to the reservoirs defines the shape of the blockade region. α_S and α_D are the alpha factors of the source and drain contacts. In the case of a symmetric geometry, $\alpha_S = \alpha_D$, and the equations 2.16 and 2.17 have exactly opposite slope, yielding symmetric coulomb diamonds. The gradient of each slope is determined by the alpha

factor, where $\alpha_i = \frac{C_i}{C_\Sigma}$ for $i = \{S, D\}$. As such, if the quantum dot is not symmetrically coupled to the source and drain, the diamonds can resemble parallelograms instead, and the asymmetry is increased the larger the difference between C_D and C_S . Assuming that the current through the dot is measured at the drain, the gradient of the positive slope is given by:

$$\frac{dV_{SD}}{dV_G} = \frac{C_G}{C_G + C_S} \quad (2.16)$$

And the negative slope is:

$$\frac{dV_{SD}}{dV_G} = -\frac{C_G}{C_D} \quad (2.17)$$

The lines cross at the point $V_{SD} = \frac{1}{e}(E_{N,N+1} + E_C)$, meaning that the lever arm can be extracted through comparison of the extent of the diamonds to the separation of the coulomb peaks at zero bias. The distance between two coulomb peaks, indicated in Fig. 2.4, is given by:

$$\Delta V_G(N, N+1) = \frac{1}{e\alpha_G}(E_{N,N+1} + E_C) \quad (2.18)$$

Taking $\Delta V_{SD} = \frac{1}{e}(E_{N,N+1} + E_C)$ at the top of the coulomb diamond, the lever arm can therefore be extracted:

$$\alpha_G = \frac{\Delta V_{SD}}{\Delta V_G} \quad (2.19)$$

Typical values for the alpha factor are in the region 0.1 – 0.3eV/V.

2.2.3 Quantum tunnelling

A quantum dot connected to a reservoir is defined by one or more tunnel barriers. These are regions of resistive material which present a potential barrier to electrons attempting to flow into or out of the quantum dot. Classical current flow can be attained by supplying a bias voltage which can overcome the barrier potential, or by reducing the potential of the barrier. However, when considering individual electrons and barriers with high potential, quantum tunnelling must be considered. In the classical picture, when a kinetic particle (such as an electron) encounters an obstacle, it cannot pass it unless it is supplied enough energy to overcome it. However, quantum tunnelling is a phenomenon whereby an electron has a finite probability to pass through a narrow barrier that it does not have the potential energy to overcome. The probability for an electron to tunnel through a potential barrier depends exponentially on the width and the height of the barrier. This is a quantum effect that can occur even at high temperatures, and can be a significant contribution to the leakage current through a nanoscale transistor [nakhmedov2005quantum].

Tunnelling is generally described by the tunnel rate parameter, Γ , with units Hz. The tunnel rate is generally dependent on the height and width of the potential barrier, as well as the energy of the electrons in the reservoir/dot ($k_B T$). The transmission probability of

a single particle with energy E to travel through a barrier of potential U and length L is given by:

$$P_{\text{transmission}} \approx 16 \frac{E}{U} \left(1 - \frac{E}{U}\right) e^{-2\beta L} \quad (2.20)$$

Where $\beta = \frac{\sqrt{2m(U-E)}}{\hbar}$. Typical tunnel rates in the quantum dots studied here range from a few MHz (singly detectable electron tunnelling events) to hundreds of GHz (continuous current), and can be measured (for multiple tunnel barriers in series) by $I_{SD} = Ie$, or by the frequency of single-shot electron tunnelling events. In order to detect a single electron tunnelling event, we require charge sensing techniques which are sensitive to single charge shifts.

2.2.4 Charge sensing

Measurement of a quantum dot via direct electron tunnelling is one way of characterising a quantum dot system. However, in the few-electron regime, it is often not possible to measure the loading of the first few electrons directly due to weak tunnelling current. In this situation, it is generally favourable to employ a charge sensor instead. A charge sensor, as the name implies, is a method of detecting the change in electrostatic energy of a system without directly affecting it. In typical semiconductor quantum dots, a typical charge sensor is the quantum point contact, or QPC. However, it is also possible to use a quantum dot as a charge sensor.

The general principle of a charge sensor is to detect a change in signal due to a shift in the charge of its environment. For a quantum dot charge sensor, this manifests as a shift in the electrochemical potential of the quantum dot via simple coulomb interaction.

In order to use a quantum dot as a charge sensor, we can consider a simple case of a single quantum dot tunnel coupled to two electron reservoirs and capacitively coupled to an arbitrary system which will be probed (see Fig. 2.5a). This is identical to the single dot case, except that we include an additional term to account for the charge occupancy of the probed system, C_p , which has an associated potential V_p :

$$\mu_{\text{dot}}(N) = U(N) - U(N-1) = (N - \frac{1}{2})E_C + \frac{C_G V_G}{C_\Sigma} + \frac{C_p V_p}{C_\Sigma} + E_{N-1,N} \quad (2.21)$$

Here $\frac{C_p}{C_\Sigma}$ is α_p , or the lever arm for the probed system, a measure of the sensitivity of the quantum dot to the probed system. A similar correction must be applied to the formula for the gate voltage position of a coulomb peak that marks the transition from $N-1$ electrons to N electrons:

$$V_G^{(N)} = \frac{1}{e\alpha_G} [E_{N-1,N} + (N - \frac{1}{2})E_C - e\alpha_G V_G - e\alpha_p V_p] \quad (2.22)$$

From here it can be seen that when the probed system is charged or discharged, and thus its voltage changed by an amount ΔV_p , that the position of the coulomb peak in the

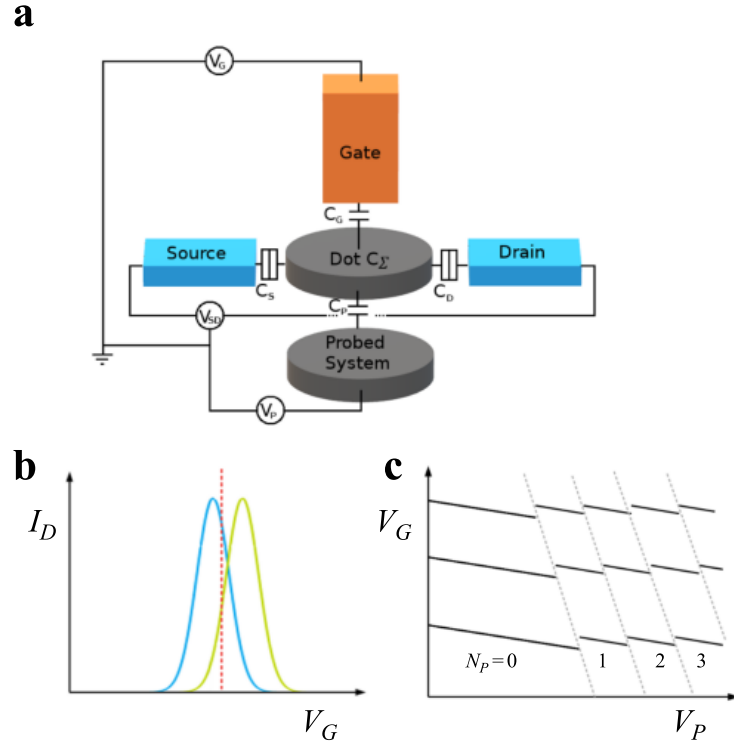


Figure 2.5: Charge detector A quantum dot configured for use as a charge detector. a) The model for a charge detector quantum dot is identical to the single dot case, with an adjacent probed system that is coupled to the quantum dot via a capacitive coupling characterised by C_P . The probed system has a voltage V_P . b) A change in the potential energy of the probed system manifests as a shift in the gate voltage required to see a coulomb peak for the detector dot. Here, the blue curve is the coulomb peak at $N_P = 0$. When an electron is loaded into the probed system, the potential shift induced in the detector dot shifts the coulomb peak to the green curve for $N_P = 1$. By measuring at the position indicated by the dashed line, it is possible to detect whether the system is in the state $N_P = 0$ or $N_P = 1$. c) A stability diagram sweeping V_P and V_G . The loading of additional electrons into the probed system is seen as abrupt line breaks in the coulomb peak lines as a function of V_G .

V_G space will be shifted by an amount equal to $\frac{\alpha_p V_p}{\alpha_G}$.

This movement of the coulomb peak can be detected by fixing V_G at a point of high sensitivity to gate movement, such as that indicated in Fig. 2.5b. The conductance through the sensor dot is measured. As the charge on the probed system is changed, the current through the sensor dot is correspondingly shifted. This can be used to detect a charging event that may not be visible in a direct-current measurement. In Fig. 2.5c, this principle is demonstrated. Here, the breaks in the coulomb peak lines correspond to the "jump" experienced by the sensor dot due to the charging of the probed system. This method can be used to detect the first charging event for a quantum dot with a high degree of certainty.

2.3 Spins in quantum dots

When considering the viability of a quantum dot for use as a qubit in a quantum computer, the most fundamental question to ask is: how do we form our two-level system? A clear candidate naturally arises in the quantum dot system in the form of electron spin. Electrons are fermions and therefore a spin-1/2 particle, meaning they can take one of two states (commonly termed spin-up, $\uparrow$, and spin-down, $\downarrow$). These spin states are energy degenerate except under the influence of a magnetic field, and the Pauli exclusion principle precludes any two electrons from occupying the same spin and orbital state. This creates an ideal two-level system, with the computational states represented by $|\uparrow\rangle$ and $|\downarrow\rangle$, in which it is possible to create superposition states and entangled states. Here we outline the spin physics in a quantum dot for both the single dot and double dot case, involving either one or two electrons, and discuss some of the readout methods used to detect the spin state of an electron.

2.3.1 Spin states in a single dot

Single electron spin states

The simplest system involving spin is the situation with a single electron in a single quantum dot. This electron can have one of two spin orientations, either spin-up ($|\uparrow\rangle$) or down ($|\downarrow\rangle$). Each of these states has an associated energy, which we denote $E(\uparrow, 0)$ and $E(\downarrow, 0)$. The first excited state is the next lowest orbital, where the spin states will have energy $E(\uparrow, 1)$ and $E(\downarrow, 1)$. Then the electrochemical potentials are:

$$\mu(0 \Rightarrow \uparrow, 0) = E(\uparrow, 0) \quad (2.23)$$

$$\mu(0 \Rightarrow \downarrow, 0) = E(\downarrow, 0) = E(\uparrow, 0) + \Delta E_Z \quad (2.24)$$

$$\mu(0 \Rightarrow \uparrow, 1) = E(\uparrow, 1) = E(\uparrow, 0) + \Delta E_{orb} \quad (2.25)$$

$$\mu(0 \Rightarrow \downarrow, 1) = E(\downarrow, 1) = E(\uparrow, 0) + \Delta E_{orb} + \Delta E_Z \quad (2.26)$$

Here ΔE_{orb} is the orbital energy level spacing and ΔE_Z is the Zeeman splitting, which is given by $\Delta E_Z = g\mu_B B$. In the absence of a magnetic field, the up and down spin states are degenerate in energy. When a magnetic field is applied, however, all energy levels are separated into Zeeman doublets, lifting the spin degeneracy.

Two electron spin states

If we now consider the case in which a single quantum dot contains two electrons, it is necessary to consider that the two electrons will be indistinguishable and must both occupy an orbital state and a spin state. Due to the Pauli exclusion principle, which prevents two electrons from occupying the same orbital and spin state, either the orbital state or the spin state must be asymmetric, resulting in a limited set of states that the system can take. The two-spin states in which the total spin number $S = 0$ are the singlet states with an antisymmetric spin part, whilst those with a total spin number of $S = 1$ are the triplet states with a symmetric spin part.

For simplicity, we consider the lowest two energy states in the quantum dot, these being the first and second electron orbitals, which we denote as the ground (first) orbital $|G\rangle$ and the excited (second) orbital $|E\rangle$. Under the condition of zero magnetic field, the lowest energy state of the system is a singlet state with the two electrons each in the lowest orbital:

$$|S\rangle = |GG\rangle \times \frac{(|\uparrow\downarrow\rangle - |\downarrow\uparrow\rangle)}{\sqrt{2}} \quad (2.27)$$

The next lowest energy two-electron states require an electron in the excited state. These are the triplet states, which are degenerate at zero magnetic field, and can be written as follows:

$$|T_+\rangle = \frac{|GE\rangle - |EG\rangle}{\sqrt{2}} \times |\uparrow\uparrow\rangle \quad (2.28)$$

$$|T_0\rangle = \frac{|GE\rangle - |EG\rangle}{\sqrt{2}} \times \frac{(|\uparrow\downarrow\rangle + |\downarrow\uparrow\rangle)}{\sqrt{2}} \quad (2.29)$$

$$|T_-\rangle = \frac{|GE\rangle - |EG\rangle}{\sqrt{2}} \times |\downarrow\downarrow\rangle \quad (2.30)$$

To complete the set of possible two-electron states in the lowest two orbitals, we have two additional singlet states, whereby either one or both of the electrons occupy the second (excited) orbital:

$$|S_1\rangle = \frac{|GE\rangle + |EG\rangle}{\sqrt{2}} \times \frac{(|\uparrow\downarrow\rangle - |\downarrow\uparrow\rangle)}{\sqrt{2}} \quad (2.31)$$

$$|S_2\rangle = |EE\rangle \times \frac{(|\uparrow\downarrow\rangle - |\downarrow\uparrow\rangle)}{\sqrt{2}} \quad (2.32)$$

These singlet states have a higher energy at zero magnetic field than the triplet states due to the symmetric orbital part. The associated energies for the singlet and triplet states can be directly extracted from the sum of the relevant single-particle energies. For the four lowest energy states:

$$U_S = 2E_{\uparrow,0} + \Delta E_Z + E_C \quad (2.33)$$

$$U_{T_+} = 2E_{\uparrow,0} + E_{ST} + E_C \quad (2.34)$$

$$U_{T_0} = 2E_{\uparrow,0} + E_{ST} + \Delta E_Z + E_C \quad (2.35)$$

$$U_{T_-} = 2E_{\uparrow,0} + E_{ST} + 2\Delta E_Z + E_C \quad (2.36)$$

Where E_{ST} is the energy difference between the singlet and the lowest energy triplet state. At zero magnetic field, $\Delta E_Z = 0$, and the three triplet states are degenerate in energy.

2.3.2 Spins in two quantum dots

We have thus far considered the case of a single spin in a single quantum dot. However, to access some configurations required for procedures such as Pauli spin blockade, we require the double dot configuration. Here we apply spin physics to the system outlined in section A.1.

When containing a single electron, the physics are identical to the single dot situation. At zero magnetic field, the spin states are degenerate; when a magnetic field is applied, then the spin-up and spin-down states are separated by the Zeeman energy, E_Z . Adding a second electron to the same dot produces the single-dot singlet and triplet states. If the right dot is populated:

$$|S(0,2)\rangle = \frac{(|\uparrow_2\downarrow_2\rangle - |\downarrow_2\uparrow_2\rangle)}{\sqrt{2}} \quad (2.37)$$

$$|T_+(0,2)\rangle = |\uparrow_2\uparrow_2\rangle \quad (2.38)$$

$$|T_0(0,2)\rangle = \frac{(|\uparrow_2\downarrow_2\rangle + |\downarrow_2\uparrow_2\rangle)}{\sqrt{2}} \quad (2.39)$$

$$|T_-(0,2)\rangle = |\downarrow_2\downarrow_2\rangle \quad (2.40)$$

At zero magnetic field, the triplets are separated from the singlet by E_{ST} . Now if we have one electron in each dot, we also obtain the singlet triplet basis, with electrons in separate dots:

$$|S(1,1)\rangle = \frac{(|\uparrow_1\downarrow_2\rangle - |\downarrow_1\uparrow_2\rangle)}{\sqrt{2}} \quad (2.41)$$

$$|T_+(1,1)\rangle = |\uparrow_1\uparrow_2\rangle \quad (2.42)$$

$$|T_0(1,1)\rangle = \frac{(|\uparrow_1\downarrow_2\rangle + |\downarrow_1\uparrow_2\rangle)}{\sqrt{2}} \quad (2.43)$$

$$|T_-(1,1)\rangle = |\downarrow_1\downarrow_2\rangle \quad (2.44)$$

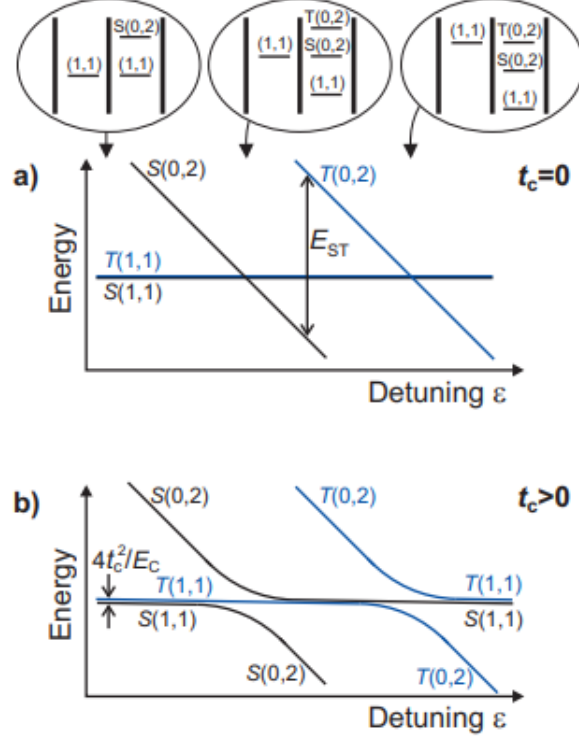


Figure 2.6: Double dot spin states The energy diagram for the two dot singlet/triplet spin states as a function of the detuning between the dots, ε . a) This is the case for two completely decoupled dots, in which the spin states are dependent purely on the single dot energies. The $|S(1,1)\rangle$ and $|T(1,1)\rangle$ states are degenerate, and the $|S(0,2)\rangle$ and $|T(0,2)\rangle$ states are separated by the singlet-triplet splitting, E_{ST} . b) When a finite tunnel coupling t_c is applied between the two dots, the $(1,1)$ and $(0,2)$ state hybridize, leading to the characteristic anti-crossings seen in the detuning energy diagram. When the single dot energy states are aligned, the singlet and triplet states are split by the exchange energy $J = \frac{4t_c^2}{E_C}$. [Adapted from [Han07]]

The energy difference between the lowest energy singlet and triplet states when the single-dot levels are aligned is given by $J = \frac{4t_c^2}{E_C}$. This energy difference, J , is referred to as the exchange energy, and strongly depends on the detuning between the dots. At $t_c = 0$, there is no coupling between the dots, and the $|S(1,1)\rangle$ and $|T(1,1)\rangle$ states are degenerate in energy. This situation is demonstrated in Fig. 2.6a.

At finite tunnel coupling, the exchange interaction is enabled. Due to finite exchange energy J , the $(1,1)$ and the $(0,2)$ charge states hybridize, meaning there is no sharp transition from one to another; instead, there is a smooth transition from $(1,1)$ to $(0,2)$, as indicated in Fig. 2.6b. This can be thought of intuitively as an overlap of the states from one dot to another as the detuning is changed. At the detuning value where there is an avoided crossing between the singlet states, the charge state for the singlet becomes $(|(1,1)\rangle + |(0,2)\rangle)/\sqrt{2}$, meaning that there is an equal probability to find the system in the charge state $(1,1)$ or $(0,2)$.

Since tunnelling does not generally involve a spin-flip, there is coupling only between the same spin states: $|S(1,1)\rangle \leftrightarrow |S(0,2)\rangle$ and $|T(1,1)\rangle \leftrightarrow |T(0,2)\rangle$. The singlet and triplet transitions are separated in energy, and it is this that enables the Pauli spin blockade readout method detailed in section 2.3.3. Crucially for this readout, there is a range of detuning where the lowest energy state for a singlet is the (0,2) charge state and the lowest energy state for a triplet is (1,1), enabling non-destructive spin-to-charge conversion.

2.3.3 Spin measurement

Measurement and resolution of a single-electron spin is challenging due to the small associated energy scales. In order to probe the spin state of an electron, we convert the spin information to charge information. This method is called spin-to-charge conversion. It is comparatively simple to measure the charge state of a quantum dot, and this can be used to probe the spin. There are three main ways to convert spin to charge, all of which rely on the difference in how spin configurations tunnel from (or into) a quantum dot.

Here we present energy selective and tunnel-rate selective tunnelling. Both methods can be used to probe either a single or two-electron spin state, and can be implemented using a single quantum dot coupled to a reservoir with an adjacent charge detector.

Energy selective tunnelling

Energy-selective tunnelling is a method that exploits the difference in energy between two spin states. For a single spin, this will be the energy difference between $|\uparrow\rangle$ and $|\downarrow\rangle$. For a two-spin state, we use $|S\rangle$ and $|T\rangle$, where $|T\rangle$ denotes the generic triplet state that comprises the three triplet states $|T_+\rangle$, $|T_0\rangle$ and $|T_-\rangle$, degenerate at zero magnetic field. For simplicity, the lowest energy state will be referred to as $|G\rangle$ and the higher energy state as $|E\rangle$, corresponding to the system ground state and excited state.

To perform energy-selective tunnelling, we require a quantum dot coupled to a reservoir, and an adjacent charge detector which is configured to measure the charge state of the dot. The charge state of the dot can be modulated through changing the chemical potential of the dot relative to the Fermi energy of the reservoir. We consider the case in which the potential of the dot is positioned such that the Fermi energy lies between the states $|G\rangle$ and $|E\rangle$. Since $|G\rangle$ is the lowest energy state, it lies below the Fermi energy, and $|E\rangle$ lies above the Fermi energy. As such, transport from the dot to the reservoir is blocked if the system is in the $|G\rangle$ state, but it is allowed if the system is in the $|E\rangle$ state. Then, an electron will tunnel back into the quantum dot, as there is an energy level available below the Fermi energy.

This tunnel-out-tunnel-in event can be observed as a charge change via the charge detector, as the dot temporarily loses the electron to the reservoir, as shown in Fig. 2.7. Curve B demonstrates the signal seen when the system is in the ground state, i.e. no tunnelling event is observed and the charge in the quantum dot does not change. Curve A demonstrates a single tunnelling event, followed by a new electron tunnelling into the dot. Through analysis of the charge time trace, the spin state of the electron at the start of the measurement can be deduced in a single-shot manner.

For this mechanism to work, the separation in energy of the two states $|G\rangle$ and $|E\rangle$ must be large compared to the thermal energy of the charge carriers. This puts a lower

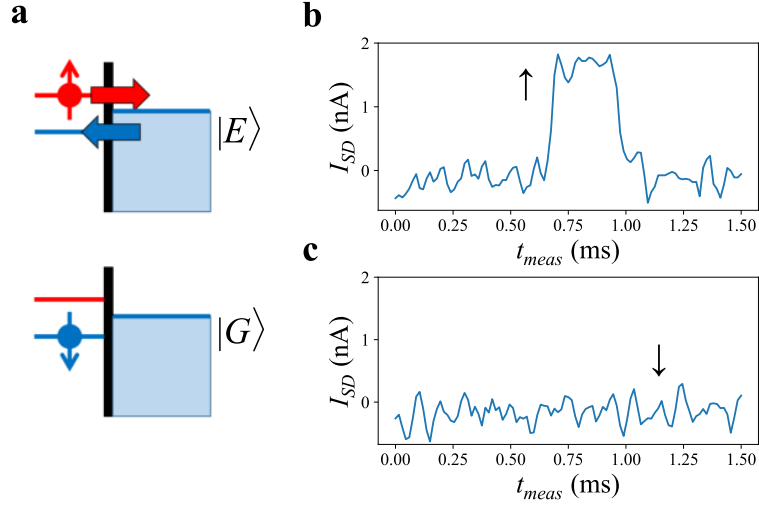


Figure 2.7: Energy selective readout Schematic of the energy selective spin readout mechanism. a) In the event that the electron in the dot is in the excited state $|E\rangle$, it is able to tunnel out of the dot and into the reservoir. After some time, a new electron will tunnel back into the newly vacated dot into the ground state $|G\rangle$. If the electron is already in the $|G\rangle$, tunnelling into the reservoir is blocked. b) & c) Time-traces of the current through a nearby QPC or other charge detector. b) is the signature seen by the detector for the situation with initial state $|E\rangle$, with the tunnelling out event and the repopulation event visible. c) is the signature seen for the situation with initial state $|G\rangle$, with no tunnelling events observed. The red dashed line indicates a threshold that can be selected to distinguish between measurement of an $|E\rangle$ state and a $|G\rangle$ state.

bound on the energy scale that is necessary to use this type of measurement in a system at a given temperature. For single spin states, for example, this therefore requires a high magnetic field ($g\mu_B B \gg k_B T$). It is also necessary to have a high tunnelling rate compared to the relaxation time from $|E\rangle$ to $|G\rangle$, so that the electron tunnels out of the dot before it relaxes to the ground state. However, a low tunnelling rate compared to the measurement bandwidth is required to be able to detect the tunnelling event, leading to a lower and upper bound for the coupling to the reservoir. This method of spin measurement is extremely sensitive to the charge environment of the quantum dot due to the precise positioning required to keep the $|G\rangle$ and $|E\rangle$ states either side of the Fermi energy.

Spin dependent tunnel rate readout

Tunnel-rate selective measurement is a viable alternative to energy selective measurement in cases where the system has a spin-dependent tunnel rate. The system is the same as in the energy selective case, configured to a state where both the $|G\rangle$ and $|E\rangle$ states for electron number N are above the fermi energy of the reservoir. The electron will tunnel off the dot and into the reservoir from either state. However, if the tunnelling time is much longer for the ground state than the excited state ($\Gamma_E \gg \Gamma_G$), there will be a period of time in which it is highly likely for the electron to have already tunnelled into

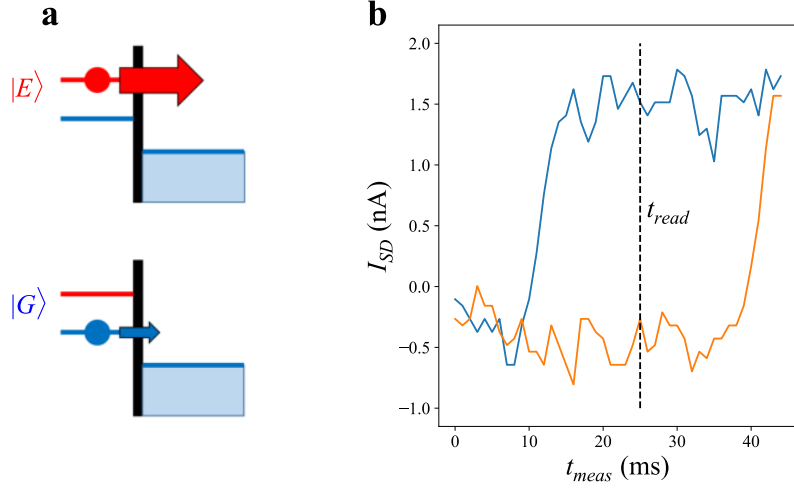


Figure 2.8: Tunnel rate readout The mechanism for a tunnel rate dependent readout. a) Here two states $|E\rangle$ and $|G\rangle$ are lifted above the energy of the reservoir. Each has an associated tunnelling rate Γ_E and Γ_G . Tunnel rate dependent readout is possible when $\Gamma_E \gg \Gamma_G$. b) Time trace readout of a nearby QPC or charge detector. The excited state $|E\rangle$ tunnels quickly into the reservoir, leading to a quick shift in QPC current to the unloaded state. The ground state $|G\rangle$ takes longer due to the lower coupling to the reservoir. A boundary in time t_{read} can be chosen to select between the two states, indicated by the dashed line. If a high current value is measured here, the dot can be said to have been in the $|E\rangle$ state initially, and conversely for a low current value the dot is measured in the $|G\rangle$ state. [Current-time trace adapted from [Ber15]]

the reservoir if the system was in the ground state $|G\rangle$, but not if it was initially in the excited state $|E\rangle$. When measured at this position, measurement of a charge state of N would indicate that the electron is in the excited state, whilst measurement of a charge state $N - 1$ would indicate that the electron was in the ground state. This protocol is indicated in Fig. 2.8. This technique is less sensitive to charge noise, since the positioning of the energy levels does not need to be as precise as in the energy selective case. Typical charge noise signatures will affect the entire energy landscape, not a single tunnel path [Jun04], and therefore are less impactful on this kind of measurement. The measurement visibility is instead largely dependent on the difference in tunnel rate.

In the two-electron spin case, there is often a difference in tunnel rate between the singlet and the triplet. In the singlet, the electrons occupy the same orbital, whilst in the triplet case, one electron is in a higher orbital. In the few-electron regime, there can be a significant difference in the coupling to the reservoir for the higher orbital due to the difference in wavefunction close to the edges of the confining potential. However, this difference can be small, and is strongly dependent on the configuration of the dot.

The same technique can be used for single spin through exploiting a difference in tunnel rate between spin-up and spin-down, which has been measured but currently has no physical explanation. The difference is small, and as such the measurement has a very low fidelity

compared to the singlet-triplet case. It also requires very high magnetic fields, in which case it is likely that energy-selective readout would have a far greater fidelity and would be more achievable. In practice, this technique would likely not be used for single-spin measurement.

Spin blockade

Both the energy selective and tunnel rate selective spin measurements are destructive, in the sense that in the process of measurement you lose the measured electron to the reservoir and therefore lose the spin information. This can be acceptable to determine the result of an operation, but is not sufficient to perform further operations; it would be necessary to re-initialise the spin in its measured state each time. To have a measurement (projection) operation that does not affect the spin, we require a protocol that preserves the spin state, a so-called quantum nondestructive measurement.

The technique described here is called Spin Blockade readout. Significantly in comparison to the other two techniques, it requires a two-electron system in a double dot consisting of Dot 1 and Dot 2. The initial state is the (1,1) configuration, i.e. a single electron in each dot. A charge detector is needed to sense the charge state of either dot.

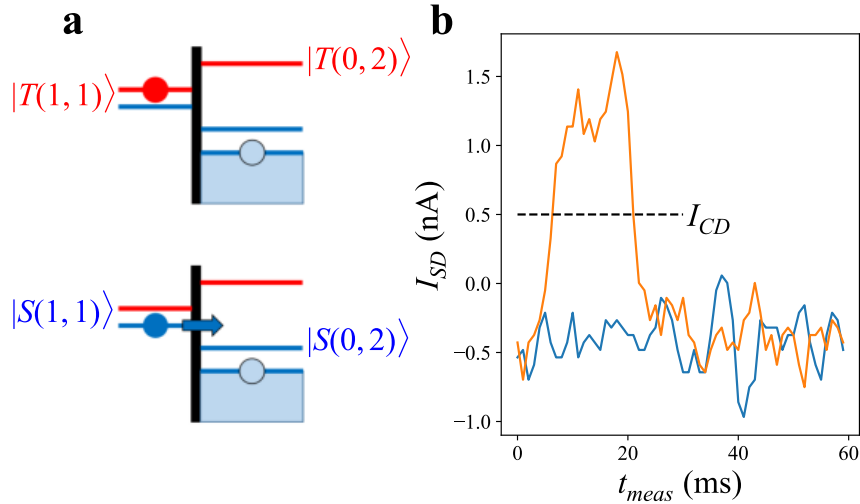


Figure 2.9: Spin blockade readout The mechanism for spin blockade readout. a) The energy level of the left dot is positioned such that the $|T(1,1)\rangle$ and $|S(1,1)\rangle$ states lie between the energy levels necessary to form $|S(0,2)\rangle$ and $|T(0,2)\rangle$. The $|T(1,1)\rangle$ state is blocked from tunnelling into the $|T(0,2)\rangle$ state, whilst the $|S(1,1)\rangle$ is able to tunnel into the $|S(0,2)\rangle$. b) Time trace readout of a nearby charge detector. In the event that the system is in a singlet state, the system will immediately move to the $|S(0,2)\rangle$ state, indicated in blue. However, if the system is initially in the triplet state, it is blocked from tunnelling until a spin-flip event occurs, which can take several milliseconds, as indicated in orange. A threshold can be defined to distinguish the two measurement results as indicated by the dashed line.

The system is then tuned to a state in which an electron in dot 1 will selectively tunnel from Dot 1 to Dot 2 depending on its spin state. This is done by detuning the dots until the energy level of the electron in Dot 1 is positioned between the energy states $|S(0,2)\rangle$ and

$|T(0,2)\rangle$ (see Fig. 2.9a). If the spins are antiparallel, then the system is in a singlet state, and the electron can tunnel from Dot 1 to Dot 2, i.e. from the $|S(1,1)\rangle$ state to $|S(0,2)\rangle$. However if the spins are parallel, then the system is in a triplet state, and the electron is blocked from tunnelling, remaining in $|T(1,1)\rangle$. The system will remain in $|T(1,1)\rangle$ until a spin-flip occurs in Dot 1 (as seen in Fig. 2.9b), which can be a slow mechanism compared to the manipulation time. Provided the transfer rate between the dots and the measurement speed are significantly faster than the spin-flip mechanism, it is possible to distinguish the two states simply through measurement of the charge state of either dot. If the system is in a singlet state, a change in charge state from (1,1) to (0,2) will be seen. If the system is in a triplet state, no change will be seen.

Crucially, the spin state of the system is preserved. This allows successive single-shot spin readout of the same electron to prove that it is a nondestructive measurement. This then allows for subsequent operations on the same state, preventing the need for re-initialization.

Pauli spin blockade can be measured experimentally in a double dot system via direct current transport measurements. If we consider the triple point between states (0,1), (0,2), and (1,1), then at negative bias the electron transfer sequence is $(0,1) \rightarrow (0,2) \rightarrow (1,1) \rightarrow (0,1)$. There is always an electron in the right dot, and this electron can be either spin-up or spin-down. In either case, in the transition from $(0,1) \rightarrow (0,2)$ the dot can only accept an electron with the opposite spin from the leads. Since the leads act as an electron reservoir, there is no problem with requiring a spin-down electron, and the cycle can proceed as normal.

However, if we instead apply a positive bias, we now have the sequence $(0,1) \rightarrow (1,1) \rightarrow (0,2) \rightarrow (0,1)$. In this case, an electron of any spin state is allowed to tunnel in to Dot 1 from the reservoir. If this electron has opposite spin to that in Dot 2, then it is able to tunnel into Dot 2, and out into the drain. However, if the electron has the same spin as that in Dot 2, then this tunnelling process is blocked, and we obtain Pauli spin blockade, where no current is allowed until a spin-flip occurs (and this can be on a long time scale relative to the measurement time). This can be directly measured by looking at bias triangles for $N = 1\%2 \rightarrow N = 2\%2$ transitions. When the system is blockaded, the bias triangles will show current at positive bias, but be rectified at negative bias (or vice versa).

CHAPTER 3

Experimental setup and characterization

3.1 Introduction

The control and characterization of a silicon CMOS quantum dot requires many steps to achieve. In addition to the hardware requirements, each individual device is unique, and must be characterized and tuned to assess its viability as a spin qubit. One of the consequences of the top-down approach taken by the CMOS silicon quantum dot community is that mass characterization and comparison of device designs, dimensions, and operating parameters is necessary to guide future development. Here, we are not constructing the perfect qubit from scratch, but instead using - and working within the limitations of - existing industrial processes to create nanodevices which are viable spin qubits. As such, rapid characterization and parameterization of devices is necessary, as is having figures of merit which can be compared across devices to determine reliability and reproducibility.

Additionally, on a smaller scale, in-depth parameterization of a single device is needed to understand its

In this chapter, the experimental setup required to cool a sample to the requisite temperatures is outlined, as well as the hardware used to control the electric fields in the device and detect current signals. We present the type of device measured in this thesis, including the fabrication process flow and the types of characterizations performed at room temperature and low temperature, and how a device can be tuned into the appropriate measurement regime. Finally, the methods for extracting various important parameters of the device and quantum dots are presented, including extraction of the charging energy, gate lever arms and electron temperature.

3.1.1 Cryogenics

Quantum physics operates on very low energy scales. At room temperatures, the thermal energy is far larger than the kind of energies we intend to probe. The most significant energy is the charging energy of our quantum dot. This gives the separation between coulomb peaks, and we require $k_B T \ll E_C$ to be able to observe quantum phenomena. As such, we require the device to be cooled to cryogenic temperatures, as typical E_C values (a few meV) correspond to a temperature around 10 K. This necessitates the use of a cryogenic refrigerator which can provide the cooling power necessary to operate the

device in this regime. For the low-temperature measurements outlined in this thesis, three different cryogenic refrigerators were used.

"Dipstick" immersion refrigerator

For rapid characterisation of many devices at low temperature, we require a refrigerator which has a short cycle time. The "dipstick" immersion fridge works well for this purpose. It consists of a cryogenic stage with 40 DC lines which can be immersed in a bath of liquid ^4He . The cryogenic stage is well-thermalized via exchange gas, bringing the sample to 4 K. This fridge can be cooled in less than an hour, allowing for rapid measurements of many devices. Most devices were first tested and characterised in this fridge first, before being measured at lower temperature.

Helium-3 circulation refrigerator

At 4 K, thermal broadening can still render quantum dot features blurred, obscuring detail of the quantum dot structure. Measurements such as single-shot charge sensing can remain difficult at 4 K. To investigate promising devices further, we use a ^3He circulation refrigerator. Low-pressure ^3He is pre-cooled via a pumped 1 K pot, and circulated to the device at the cold finger in the cryogenic stage of the fridge. The cold finger is kept under vacuum to minimize thermal coupling to the ^4He bath. The ^3He is pumped with a primary and turbo pump to obtain the minimum temperature of approximately 400 mK. This refrigerator was used for the majority of the charge noise measurements and characterization of the main device.

Dilution refrigerator

In order to probe the spin physics of electrons within these devices, we require temperatures $k_B T \ll E_Z$. For spin measurements, we moved the device to a Kelvinox[®] MX $^3\text{He}/^4\text{He}$ dilution refrigerator. This fridge offers a lower base temperature compared to the ^3He refrigerator, around 70 mK (with a corresponding electron temperature of around 120 mK). It makes use of a dilution process involving the phase transition of the two helium isotopes. The principle of experimental operation is similar to that in the He-3 refrigerator. The cryogenic principles for dilution refrigerators are well documented and will not be covered in detail here.

3.1.2 Electronics

In order to control the device in each refrigerator, we require one DC line for each gate, reservoir contact, and the top and back gates. All voltage supply, pulse generation and readout is performed via room-temperature electronics.

Wiring

Each refrigerator used here has approximately 40 thermally-anchored DC cables to carry signal from the room temperature electronics to the cryogenic stage of the fridge. The majority of the DC lines used were copper-nickel alloy Constantan wires, which display similar characteristics from high to low temperature. The bandwidth of these lines is approximately 10 MHz, which is sufficient for control and readout on a typical time scale of tens of microseconds.

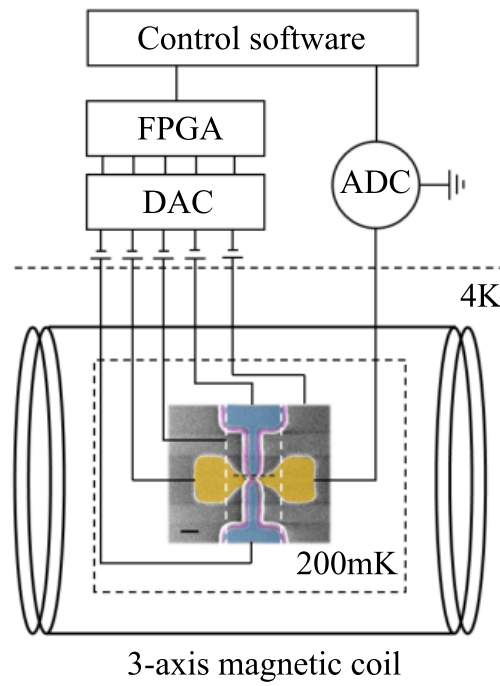


Figure 3.1: Experiment schematic Depiction of the measurement setup. The experiment is controlled digitally, with instructions sent to the field-programmable gate array (FPGA), which controls the digital-to-analog voltage supply. DC cables carry the voltage from room temperature to the device at a nominal temperature of 200 mK (real temperature ranges from 100 – 200 mK). A superconducting coil, immersed in the ^4He bath at 4 K, is used to control the magnetic field along the three cardinal axes. The current through the device is measured using a room temperature analog-to-digital converter in combination with a current-voltage converter and amplifier with a gain of $1 \times 10^9 \text{ V/A}$. In the ^3He circulation refrigerator, the base temperature is 400 mK instead, and the magnetic coil is monodirectional. The magnetic coil and 200 mK stage are not present in the "Dipstick" refrigerator.

Voltage supply

The state of the device is manipulated by supplying voltages to the gates. To set the voltage state, we use a set of purpose-built low-noise digital-to-analog converter (DAC) boards. They can sweep at a rate of $2.5 \text{ V}/\mu\text{s}$ with a noise of $25 \text{ nV}/\sqrt{\text{Hz}}$. The range of each board is $\pm 5 \text{ V}$, with a resolution of $150 \mu\text{V}$. For finer voltage resolution, an additional DAC output is used to provide a supplementary voltage through a voltage adder-divider. The secondary output has a resistance of $5 \text{ k}\Omega$, dividing the DAC output by 50. This gives a voltage resolution of $3 \mu\text{V}$.

The DAC outputs are controlled by a National Instruments® sbRIO-9208 field programmable gate array (FPGA). The FPGA receives the ramp or pulse instructions and stores them temporarily in memory before passing the pulse shape to the DAC, allowing voltage ramps and pulses on a faster timescale than the communication between the control software and the FPGA. The minimum setting time is $16 \mu\text{s}$, which gives the minimum step time in an experiment. In practice, this is faster than our typical integration time, and is shorter than the time scales probed here (typically hundreds of μs up to hundreds of ms).

Current readout

Measurement of the output current is done via a National Instruments® USB-6229 analog-to-digital converter. This is combined with a purpose-built current-voltage converter and amplifier, with a gain between $1E6$ and $1 \times 10^9 \text{ V/A}$. In combination the readout bandwidth is around 100 kHz at a gain of $1 \times 10^7 \text{ V/A}$. Digital triggering allows for precise timing of pulse sequences down to a few tens of μs .

Magnetic coil

The magnetic coil used in the majority of the magnetic field experiments is a three-directional electromagnet controlled by an Oxford Instruments® Mercury IPS power supply. The z -coil can supply a field of up to $\pm 6 \text{ T}$ perpendicular to the plane of the device. The x and y coils are limited to $\pm 3 \text{ T}$ in the plane of the device. Each coil can be swept at a rate of 0.2 T/min and are controlled digitally to synchronize the field with measurement of the device.

3.1.3 Software

The electronics were interfaced with digital control via an in-house Python and LabVIEW hybrid control software. Experiment "batch" files are generated via Python scripts, which then pass a set of instructions to the LabVIEW interface. Simple voltage sweeps are configured using a "ramp" method, whilst more complex experiments can be designed with a pulse map. The LabVIEW interface provides real-time data updates, logging, and experiment queuing. Data generated by an experiment is stored in a compressed .h5 file format.

For data treatment, custom Python scripts were written for analysis and plotting. The PyPlot library was extensively used to generate the graphs and figures used in this thesis. The .h5 format data files were addressed using the h5py library. For data treatment, the NumPy and SciPy libraries provide tools for manipulation and analysis of data; in particular, the `scipy.signal.periodogram` function was used to calculate the power spectral

density of charge noise traces.

3.2 Characterization

The measurements made in this thesis were made in the context of the quantum silicon project, spearheaded by several groups of researchers around Grenoble. The samples were provided by our industrial partner CEA-LETI. One of the main goals of this project is to create not only a reliable qubit in silicon, but one that is reproducible and scalable using standard industrial practices. As such, parametric testing of devices at room temperature and correlated measurements at low temperature are necessary to characterise the properties of devices produced at scale.

3.2.1 Sample design

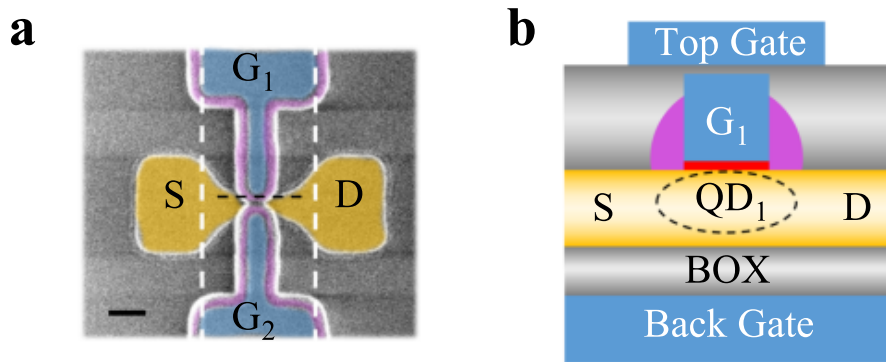


Figure 3.2: Device design a) SEM image of a typical two-gate face-to-face device after deposition of the polysilicon gates and spacers. The solid black bar (bottom left) is a 100 nm reference scale. The nanowire (yellow) is formed of etched intrinsic silicon and has a width $W_{approx} 90$ nm in the channel region. The polysilicon gates (blue) are deposited overlapping the nanowire as a single gate, and then cut horizontally to form two face-to-face gates. They have a length of 50 nm and a contact area of 50×20 nm². 35 nm SiN spacers are deposited around the edges of the gates (pink) to define the tunnel barriers in the device and protect the channel from implantation. After encapsulation, a metallic top gate is placed 400 nm above the sample in the region outlined by the white dashed lines. b) A cut along the black dashed line in a) (not to scale). The wafer consists of bulk silicon, which can be polarized using an applied voltage and an LED to generate photocarriers, and is operated as a back gate. It is insulated from the nanowire by a 145 nm buried oxide (BOX). The nanowire is depicted in yellow, and consists of a 16 nm high channel with source and drain defined by implantation. The approximate location of the quantum dot is indicated, as it is formed by accumulation under the plunger gate, G_1 . The polysilicon gate is insulated from the channel by 6 nm of silicon oxide, and an additional 5 nm of TiN forms the interface (in red). The channel, gate, and spacers are encapsulated, and a metallic top gate placed 400 nm above the sample.

The device used here was an industrially-fabricated CMOS silicon nanowire design. It

was designed and fabricated using an industry-standard 300 mm wafer CMOS process at CEA-LETI. The design consists of an intrinsic silicon nanowire of width $W \approx 90$ nm, with electron reservoirs at either end formed by ion implantation. Two accumulation gates are placed on top of the wire to form two parallel quantum dots. Above the channel and accumulation gates, a metal gate is placed to act as a top gate.

The fabrication process involves a 300 mm SOI wafer consisting mostly of intrinsic silicon. A 145 nm buried oxide (BOX) insulates the bulk silicon wafer from the active layer. The active layer is 16 nm in depth and comprises intrinsic silicon. It is possible to isotopically purify this active layer (the wafers tested here did not undergo this process). The nanowire is formed via electron beam lithography for precise definition, and etched down to a width of $W \approx 90$ nm.

Next, electrostatic gates are defined to cover the nanowire. The stack is insulated from the channel by 6 nm silicon oxide, and the interface is formed by 5 nm of TiN placed by atomic layer deposition. Then 50 nm of polysilicon forms the bulk of the resistive gate material. The gate stack is protected in the region of the channel by a hard mask of SiN and SiO₂, which protects the gates during etching. The hard mask is removed after the gate definition. The device after hard mask removal is pictured in Fig 3.2, consisting of the etched nanowire and gates. The gates in this device have a length of $L_G = 50$ nm and a vertical spacing of $S_V = 50$ nm, leaving an overlap area of 50×20 nm². Spacers consisting of 35 nm SiN are deposited on the channel surrounding the gates. These shield the channel from ion implantation, which is used to define the reservoirs. The device is then contacted through silicon vias, and encapsulated. 400 nm above the device, a metallic top gate is placed to allow an additional degree of control.

An important point to note is the 5 nm of TiN is not present on all devices fabricated on these wafers due to a problem in the fabrication process. This has led to some deviations from expected properties in devices missing TiN. However, it was found that the presence or lack of TiN could be detected by examining the room temperature current-voltage characteristics and comparing the threshold voltage V_{th} . Specifically, devices with $V_{th} < 0.2$ V are missing TiN, whilst those which have a $V_{th} \approx 0.2$ V have the TiN. Based on the room temperature characteristics, we believe that the device measured here contains the TiN, which may account for some differences between this and other similar devices.

3.2.2 Room temperature characterization

Parametric testing is performed via a four- or six-tip probe station at CEA-LETI on the wafer scale. Mass batch characterization of many devices is possible in a relatively short time, giving statistics on device behaviour. For each device, the $I_{SD} - V_G$ characteristics of each gate are measured, along with the bias dependence through the channel. An $I_{SD} - V_G$ curve is measured by applying a voltage on the source contact (typically in the range of a few mV or tens of mV) and measuring the current through the sample at the drain contact, I_{SD} , as a function of the gate voltage, V_G . A typical $I_{SD} - V_G$ measurement is shown in Fig 3.3a. From such measurements, we can extract the threshold voltage V_{th} , the sub-threshold slope, and the channel resistance. At this stage, statistics on the population of non-functional or defective devices can be obtained. We can also extract the threshold voltage, the drain-induced barrier lowering, and subthreshold slope, all of which

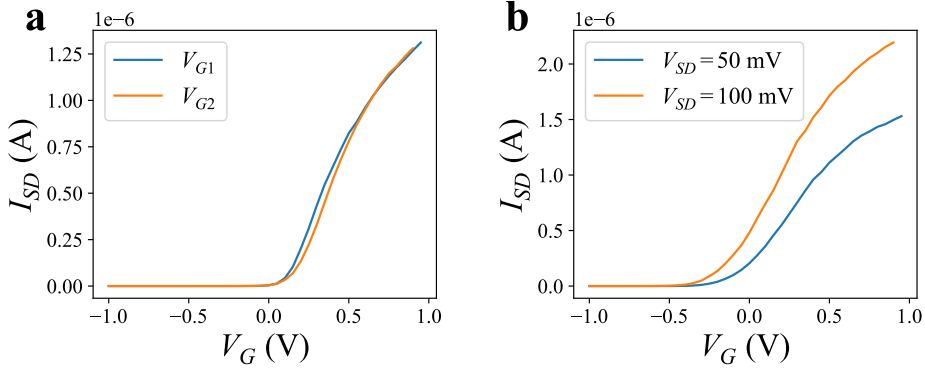


Figure 3.3: Room temperature characterization a) $I_{SD} - V_G$ characteristic of a split-gate type device at room temperature. Each curve is measured by sweeping one of the gates whilst keeping the other closed. The two gates display similar characteristics, with a voltage threshold $V_{th,G1} \approx V_{th,G2} \approx 0.18$ V. b) Drain-induced barrier lowering (DIBL) effect measured in a single gate transistor type device. The gate is swept for two different bias voltages ($V_{SD} = 50$ mV and 100 mV). Short-channel effects cause the threshold voltage to shift with applied bias. For a $V_{th} = -0.19$ V at $V_{SD} = 50$ mV, increasing the bias to $V_{SD} = 100$ mV lowers the threshold voltage to $V_{th} = -0.28$ V.

are parameters which can be used to compare devices.

Threshold voltage

The threshold voltage, V_{th} , is defined as the minimum voltage at which there is a conductance path from source to drain. It is one of the fundamental parameters used to characterize MOSFET devices [OC02]. We use the constant-current (CC) method to assess the V_{th} of a device. In the CC method, an arbitrary threshold source-drain current value is chosen, with a bias voltage $|V_{bias}| < 100$ mV applied. During room-temperature parametric testing, $I_{SD} - V_G$ curves are generally acquired at $V_{bias} = 50$ mV. A typical $I_{SD} - V_G$ characteristic is plotted in Fig 3.3a.

The threshold current value is dependent on the geometry of the device. A wider (wider channel) device requires a comparatively lower gate voltage to attain the same current level, and the inverse is true for a longer (wider gate) device. As such, the drain current is defined proportional to the width and inversely proportional to the gate length: $I_{th} = (W/L_G) \times 10^{-7}$ [Tsu99]. The threshold voltage V_{th} corresponds to the voltage at which $I_{SD}(V_{th}) = I_{th}$. It is a relatively simple method of extracting the threshold voltage, and can be used to characterize large numbers of devices in a relatively short time. It also allows for comparison of the voltage threshold across different device geometries. However, it is also highly dependent on the value of the bias voltage. The applied bias voltage is therefore fixed to allow comparison of the voltage threshold between devices.

Other methods may be used to define the bias threshold, such as the linear extrapolation method, the ratio method, or the second-derivative method. The linear extrapolation

method makes use of the linear part of the $I_{SD} - V_G$ curve: a linear fit is applied at the point of maximum transconductance (dI_{SD}/dV_G), and V_{th} is defined as the point where this linear fit intercepts $I_{SD} = 0$. This method can be degraded by mobility fluctuations and parasitic resistances due to their distortion of the $I_{SD} - V_G$ characteristics. A similar method makes use of the ratio between the current and the transconductance to obtain a linear fit instead. The ratio method takes the intercept of a linear fit to the curve $I_{SD}/g^{1/2}$, where g is the transconductance; V_{th} is where a linear fit to this curve reaches $I_{SD}/g^{1/2} = 0$. This method was shown to be independent of mobility effects [Ghi88], as well as parasitic resistances [Fik95]. Finally, in the second-derivative method, V_{th} is defined as the gate voltage at which the second derivative of the current, d^2I_{SD}/dV_G^2 , is maximum. This can be combined with the constant-current method [Zho99] to allow the definition of the threshold current I_{th} to be more rigorous, whilst maintaining simplicity for large-scale device characterization.

Drain-induced barrier lowering

The devices measured here are on the order of a few tens of nanometres in length. They therefore experience short-channel effects, such as the drain-induced barrier lowering (DIBL) effect. This is a phenomenon whereby the short channel allows a strong source (or drain) voltage to reduce the gate voltage requirement to turn the device on. It is defined as:

$$\text{DIBL} = -\frac{V_{th}^{high} - V_{th}^{low}}{V_{bias}^{high} - V_{bias}^{low}} \quad (3.1)$$

Where V_{th}^{high} is the threshold voltage at a high bias voltage V_{bias}^{high} , usually 1 V. V_{th}^{low} is the threshold voltage at a low bias voltage V_{bias}^{low} , which is usually 50 mV.

The DIBL effect is caused by an extension of the reservoir potential into the channel, closer to the gate, as the bias voltage is increased. This shortens the depletion region controlled by the gate, and therefore lowers the gate potential required to reach conductance, as shown in Fig 3.3b. However, it also has an impact on the subthreshold region, and the subthreshold slope can be reduced. This can be interpreted as the gate having a lesser effect on the conductance of the channel, and a greater change in the gate voltage required to have the same effect as it has at low bias.

Subthreshold slope

Despite the definition of the threshold voltage as the voltage required to turn on a device, there is no sharp cutoff in conductance at room temperature, as seen in Fig 3.3a. Even below the threshold voltage, there is a parasitic leakage current. Leakage is a quantum effect, consisting of a combination of thermionic emission and tunnelling effects, whereby electrons are able to tunnel through even a high tunnel barrier potential with a finite probability [Jun06]. As the device dimensions decrease, the tunnel barrier provided by the channel length also decreases, and the leakage current can become a significant factor, accounting for up to 50% of power consumption in nanoelectronic circuits. The

subthreshold region can be characterised by the subthreshold slope, which is the gradient of the slope of an $I_{SD} - V_G$ curve below the threshold voltage. The $I_{SD} - V_G$ characteristic curve exhibits linear behaviour in log scale in the subthreshold region. The gradient of the linear section of this curve, in mv/decade, gives the subthreshold slope. Since the subthreshold slope characterizes the behaviour of the device in a regime where quantum effects are relevant, it may prove to be a crucial parameter for obtaining information on quantum properties of devices at room temperature. At present, we do not have a rapid method of characterizing quantum effects in a device at room temperature, and cryogenic testing is necessary, increasing the cycle time for development and preventing large-scale testing. In order to parameterize devices in terms of quantum effects, we bring them to low temperature to measure the charging energy, gate and reservoir capacitances, lever arm, and electron temperature.

3.2.3 Low-temperature characterization

Low temperature characterization was carried out generally on the "dipstick" immersion fridge at 4 K, with some measurements also made at 400 mK on the ^3He circulation fridge. Initial characterization of a device typically involves analysis of the quantum dot structure through I-V curves and stability diagrams. A typical stability diagram of a 1S-like device is shown in Fig 3.4. Coulomb peaks due to the quantum dots are visible as conductance peaks coupled primarily to one of the two gates. I-V characteristics can reveal the coulomb blockade spectrum, but do not capture the full picture of the device structure. In particular, additional dots (or dopants in the channel) can appear as additional conductance peaks or potential shifts (similar to that seen in charge sensing) coupled to both gates. In general, we define two regimes for analysis of a device. The first, in which coulomb peaks are visible, is the many-electron regime. In some types of device, such as single-gate and pump devices, this is the only accessible regime. The first coulomb peak is usually visible at $N \approx 10$, but this can vary depending on the size of the dot, gate dimensions, and coupling to the reservoirs. When there is an adjacent quantum dot which can be operated as an SET, the few-electron regime becomes available, as shown in Fig 3.4. Here, charge sensing allows detection of the addition of electrons down to the first. The analysis of the quantum dot can be more precise, as single charge tunnelling events can be detected.

Many-electron regime

Charging energy

The charging energy is an important parameter of a quantum dot. It is a qualitative measure of the size of the dot, and defines the energy scale for processes involving the dot. It can also provide information about the type and structure of the dot. Examples include a metallic quantum dot, which will have a flat charging energy spectrum at all gate voltages (typically low); dopants, which will generally only have one to two available energy levels and a high charging energy; and multiple dots, which can be indicated by periodic variations in charging energy.

The charging energy measures the energy required to add an additional electron to a quantum dot. It is detailed in section 2.2.1 in a general context. In order to measure

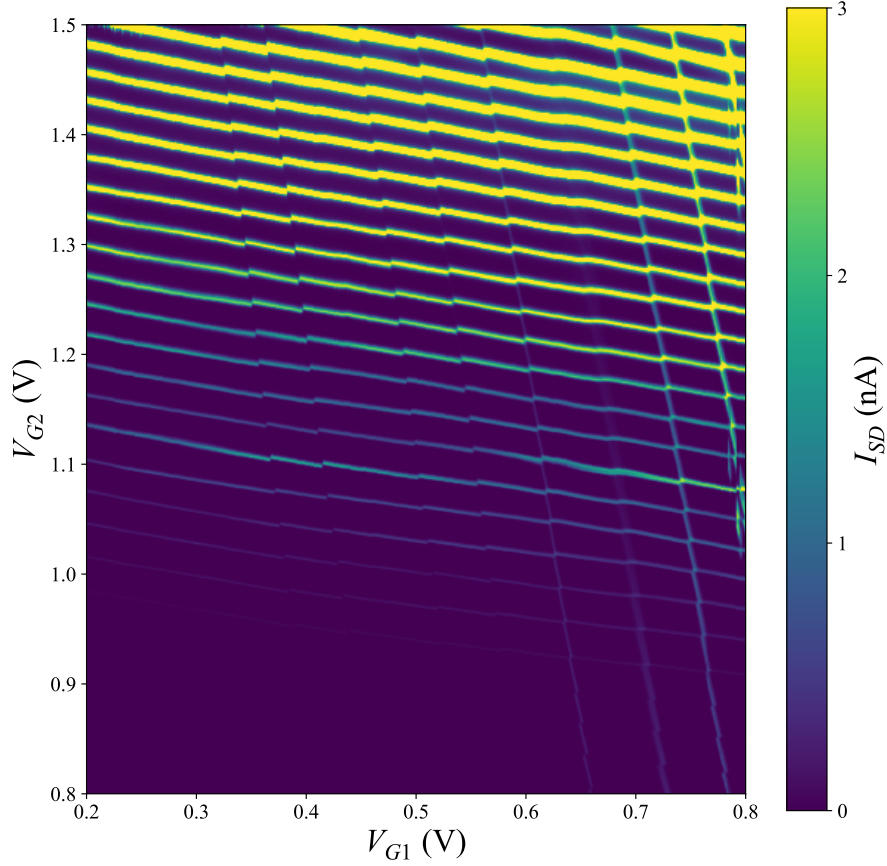


Figure 3.4: Low-temperature stability diagram Stability diagram depicting the electronic structure of the quantum dots below the gates G_1 and G_2 at V_{BG} and $V_{TG} = 0$ V. Transport through a coulomb peak is indicated by a line of high current in the stability diagram. Between coulomb peaks, the dot is blockaded and current cannot flow. At low V_{G1} , no coulomb peaks are visible, but jumps in the coulomb spectrum of the dot under G_2 are visible. These are caused by electrons loading into the dot under G_1 . Below $V_{G1} = 0.35$ V, no more such transitions are visible, allowing identification of the first loaded electron. Above $V_{G1} = 0.7$ V, coulomb peaks become visible for $N_e > 5$. The gates in this device display asymmetric behaviour at low temperature, with the first visible peak in the other dot located at $V_{G2} = 1$ V.

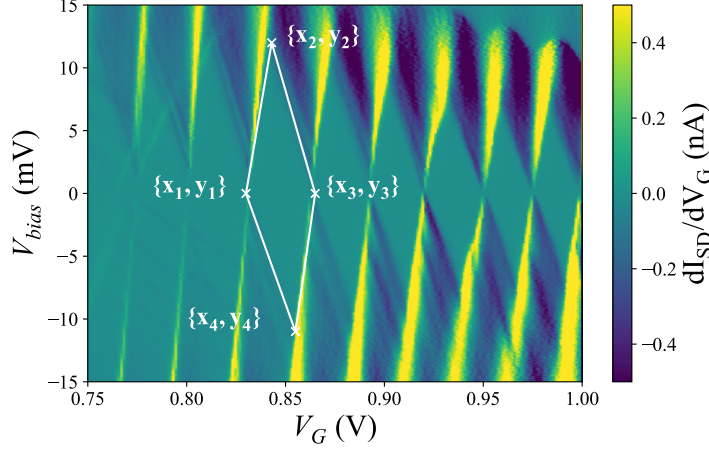


Figure 3.5: Coulomb diamonds A typical coulomb diamond map close to the transition from the few-electron to the many-electron regime. In the few-electron regime, the coulomb diamonds often do not close as $V_{bias} \rightarrow 0$ mV. Coulomb peaks begin to be visible at low V_{bias} in the many-electron regime, towards the right of the plot. The extrema points of the first visible closed coulomb diamond are plotted. These can be used to extract the configuration-specific charging energy E_C , self-capacitance C_Σ , contact capacitances C_G , C_S and C_D , and the gate lever arm, α_G , as described in the text.

the charging energy, we measure how much energy we need to supply the quantum dot to load a new electron. We have two methods of doing this: by varying the gate voltage V_G , or by increasing the source-drain bias V_{bias} . If the lever arm of the gate is not known, bias spectroscopy is another method of overcoming coulomb blockade. The source-drain bias may be increased to a point whereby the bias window is large enough that electrons with high energy can overcome coulomb repulsion and tunnel into the dot. To probe the charging energy and capacitances of the quantum dot, a coulomb diamond map is the quickest and simplest method to obtain this information. $I_{SD} - V_G$ characteristics are obtained for varying V_{bias} , as shown in Fig 3.6a. The four extrema of a coulomb diamond are indicated. Each has an x and y value, corresponding to their V_G and V_{bias} value in V, giving a set $\{x_1, x_2, x_3, x_4\}$ and $\{y_1, y_2, y_3, y_4\}$, with 1,2,3,4 corresponding to the points clockwise starting at the most negative V_G . The charging energy can be found from the average of the ΔV_{bias} of the two extrema:

$$E_C = \frac{|y_2| + |y_4|}{2} \quad (3.2)$$

A typical value for quantum dots in this type of nanowire device ranges from one to a few meV. The self-capacitance of the dot, C_Σ , is:

$$C_{\Sigma} = \frac{e}{E_C} \quad (3.3)$$

Capacitances

A quantum dot is capacitively coupled not only to its own plunger gate, but to neighbouring gates and to the reservoirs. From a coulomb diamond map, we can extract the relevant capacitances. Assuming the case where a dot is controlled by a single gate V_G and capacitively coupled to the source and drain via tunnel barriers, the gate capacitance C_G , the source capacitance C_S and the drain capacitance C_D can be reliably extracted with a four-point method. Then the capacitances (and charging energy) can be calculated as follows:

$$C_G = \frac{e}{x_3 - x_1} \quad (3.4)$$

$$C_S = \frac{2C_G}{\left| \frac{y_2 - y_1}{x_2 - x_1} \right| + \left| \frac{y_3 - y_4}{x_3 - x_4} \right|} - C_G \quad (3.5)$$

$$C_D = \frac{2C_G}{\left| \frac{y_4 - y_1}{x_4 - x_1} \right| + \left| \frac{y_3 - y_2}{x_3 - x_2} \right|} \quad (3.6)$$

The shape of a coulomb diamond gives some clues as to the relative coupling to the source and drain. If it is horizontally symmetric, the source and drain capacitances are equal, whilst the more asymmetric it is, the greater the difference between them. The capacitance of additional gates can be measured in the same way for a single dot. V_G should be set so that the dot is in coulomb blockade or on a coulomb peak, then the appropriate gate can be swept against the bias voltage. Four-point measurement of coulomb diamonds, given a wide enough gate sweep, can yield the lever arm to any arbitrary gate. It is important to note that the lever arm measured is only valid for a small deviation from the measurement position. A significant change in voltage, especially on the plunger gate, can distort the dot such that its self-capacitance increases or decreases significantly, thereby decreasing or increasing a corresponding lever arm. Similarly, distortion of the quantum dot can change the distance to a contact; this is especially relevant for the reservoirs, as the dot expands laterally, and the back gate, as the dot expands vertically in the channel. It is therefore necessary to re-measure the relevant lever arms whenever the dot is in a new configuration.

In large devices with many gates and quantum dots, the picture can become complex. Even with nearest-neighbour assumptions, the complexity can increase rapidly. As such, in devices with many dots, or multiple gates controlling a single dot, a capacitance matrix is often calculated dynamically for a quantum dot. Such a capacitance matrix can make control of a dot easier, by automating the application of compensating voltages on adjacent gates to maintain the environment of a dot whilst, for example, only increasing the on-site potential. This is effective for large devices which require fine manipulation in a small regime. However, the distortion of the dot when investigating effects that span occupancies from one electron to many tens of electrons makes such a technique infeasible even in a

small (single- or few-gate) device, and so it was not used for characterization here.

Lever arm

An important parameter to understand how the dot is manipulated by applied voltages is the so-called lever arm, or alpha factor. It is a conversion factor to translate the voltage applied on a gate to the difference in energy felt by the quantum dot. The lever arm of a contact on the quantum dot is given simply by the ratio between the relevant capacitance and the self-capacitance:

$$\alpha_i = \frac{C_i}{C_\Sigma} \quad (3.7)$$

However, both the self-capacitance and the gate capacitance can change as a function of the size and number of electrons in the dot, and coulomb diamonds are not easily measurable to low electron number. As such, extracting the lever arm of a quantum dot in the few-electron regime can be more complex.

Few-electron regime

First electron

The addition of electrons into a quantum dot is indicated by the coulomb peak spectrum. When an energy level of the dot lies in resonance with the bias window, current is allowed to flow through the dot, and a coulomb peak is observed. If the potential of the gate controlling the dot is increased further, the energy level is filled and the dot enters coulomb blockade, characterised by the current rectification between peaks. As a peak is crossed towards more positive gate voltage, we therefore know that it contains $N + 1$ electrons. However, the first coulomb peak is not in general seen at the transition of the first electron. When the quantum dot is at a low potential and contains few electrons, the coupling to the reservoirs can be weak, with the transport through an available energy level being low enough that it is below the noise floor of our current detector. As such, DC measurement of the first electron is not possible. Alternatively, we can make use of an adjacent quantum dot as a charge sensor to detect the transition of electrons into and out of the dot. Fig 3.4 demonstrates charge sensing of multiple consecutive electron transitions. As the energy levels of the dot are brought to the reservoir potential, the current through the dot is not detectable - but the change in charge state of the dot is visible as a capacitive shift in the state of the sensor dot. The magnitude of this shift is always the same, meaning we can conclude they all come from the same dot undergoing a charge state shift of $\pm 1e$.

By reducing the voltage on the gate controlling the dot, we can effectively empty the dot of electrons until we reach the last electron transition. Below this, there are no more electrons entering or leaving the dot, so the potential of the first electron can be quickly found. The potential at which the first electron can enter the dot can be heavily dependent on the configuration of the other gates, including that of the sensor dot. In some devices, it was found that the coupling between the sensor and the dot was stronger than that between the source/drain contacts and the dot, and the sensor was therefore acting as a

reservoir to load and unload electrons.

Addition energy

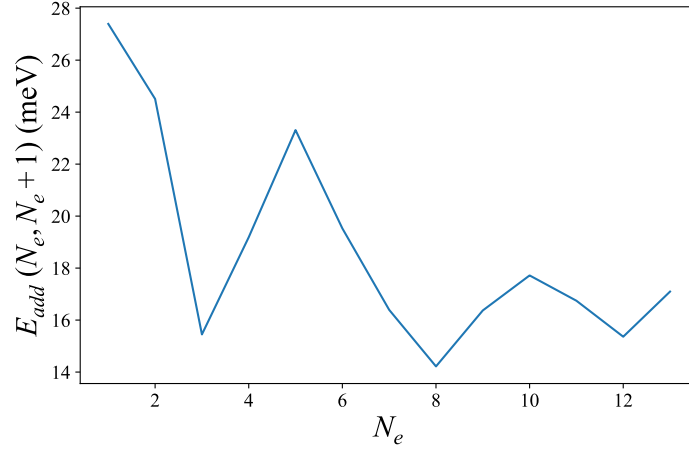


Figure 3.6: Addition energy Addition energy spectrum of the first 12 electrons in a typical quantum dot. The addition energy exhibits a general decline as the spacing between energy levels becomes closer as the quantum dot is more occupied. A significant peak is seen at $N = 5 \rightarrow 6$. Such a peak would normally be expected at $N = 4 \rightarrow 5$, characteristic of a quantum dot in silicon which has four states in the lowest orbital (two spin-degenerate and two valley-degenerate), however quantum dots can contain local potential minima which disrupt the orbital structure of the dot and shift the addition energy spectrum.

The charging energy of a quantum dot is a purely classical parameter, arising from the coulomb repulsion of the negatively-charged electrons in the quantum dot. However, a quantum dot is defined by a potential well. Under confinement, energy levels within a potential are separated in energy as electrons occupy different orbital states, similar to electron orbitals in an atom. Due to this, quantum dots are sometimes called "artificial atoms". This arises as an additional energy contribution on top of the charging energy required to add an additional electron:

$$E_{add} = E_C + E_{N-1,N}(N) \quad (3.8)$$

Where $E_{N-1,N}(N)$ is the energy difference between the orbital states of electrons $N - 1$ and N . This can be zero; for example, in the absence of a magnetic field, each state is twofold degenerate in energy, and two electrons with the same energy can consecutively occupy the dot, separated only by the charging energy. However, when loading an electron into the next orbital state, the addition energy can be measured. By probing the voltage required to add an additional electron for $N = 1, 2, 3, \dots$, we can construct an addition energy

spectrum. Such a spectrum is plotted in Fig 3.6b.

The addition energy for an ideal quantum dot is well-known, and is defined by the atomic orbitals in the electronic shell model. However, the physical realisation of a quantum dot can complicate the picture. The presence, especially in silicon, of valley states, can yield additional electron degeneracies, lifting the orbital transitions to higher electron numbers. Additionally, local disorder at the site of a quantum dot can significantly perturb the potential in the few-electron regime. In this regime, the dot is poorly defined by the applied gate potential, and can be distorted by local potential minima. This can disrupt the shell filling spectrum, and yield an addition energy spectrum which does not align with the "artificial atom" interpretation. For the quantum dot shown here, a significant peak is seen at $N_e = 5$, meaning the addition energy of the fifth electron $E_{add}(5 \rightarrow 6)$ is significantly higher. This is an indication that there are multiple potential minima in this quantum dot in the few-electron regime, and the addition energy spectrum is imperfect.

Lever arm

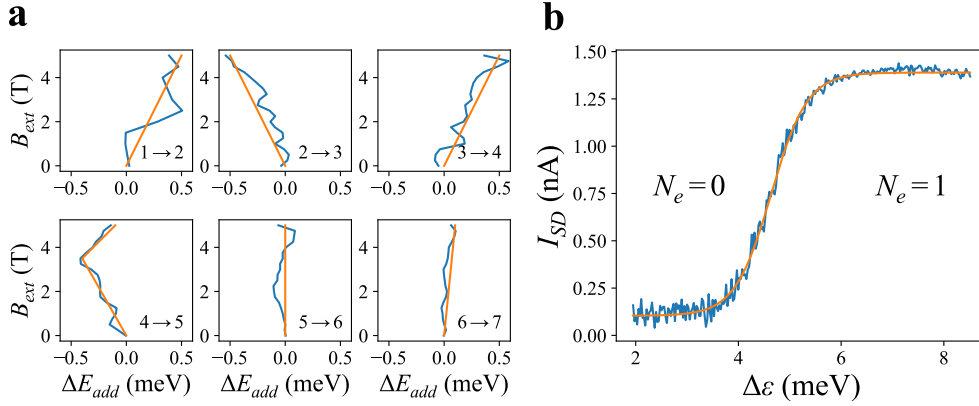


Figure 3.7: Lever arm calculation a) Magnetospectroscopy of the addition energy of the N^{th} electron. The change in the addition energy is plotted as a function of magnetic field for electron transitions starting at $N = 1 \rightarrow 2$, up to $N = 6 \rightarrow 7$. It can be seen that the change in addition energy switches from positive to negative for odd and even N , due to the relative spin parity. The linear fit is given by $\Delta E_{add} = \Delta V_{G,add} \alpha_G = g \mu_B B$. Since g and μ_B are known, this allows to extract α_G in the few-electron regime. Above $N = 4 \rightarrow 5$, the addition energy remains constant. This may be indicative of multiple degenerate states preventing spin-selective tunnelling, and is an indication that spin measurement would not be possible at higher electron numbers in this quantum dot. At $N = 4 \rightarrow 5$, the addition energy experiences a characteristic "kink", indicating a crossing with a higher energy level. This occurs at 3.5 T, and may correspond to the valley state later measured at 297 μeV . b) A fit of the average state of a quantum dot across the transition of the first electron. A current value of $I_{SD} \approx 0.1 \text{ nA}$ indicates that the quantum dot is empty, whilst a current value of $I_{SD} \approx 1.3 \text{ nA}$ indicates it contains one electron. The width of this plot is determined by the electron temperature. It is fitted with a Fermi distribution (orange) to extract the electron temperature (if α_G is known), or the lever arm (if the electron temperature is known).

Extracting the lever arm at low electron number is not possible with coulomb diamonds, and therefore two alternative methods can be used. The first makes use of the known

change in the addition energy spectrum with an applied magnetic field due to the increased Zeeman splitting, and is called the magnetospectroscopy method. The second relies on knowledge (or a good estimate) of the electron temperature at the quantum dot. Then a fit of a fermi distribution to an electron transition can yield the lever arm directly by relating the width of the transition in voltage to the expected thermal broadening.

Magnetospectroscopy

The addition energy of an electron is dependent on the magnetic field. If we consider a twofold degenerate energy state, in the absence of a magnetic field, the addition energy from the first to the second electron will be simply $E_{1 \rightarrow 2} = E_C$. However, in the presence of a finite field B , the two energy states become separated by the difference in energy between the spin states, $E_Z = g\mu_B B$. The addition energy then becomes $E_{1 \rightarrow 2} = E_C + E_Z(B)$.

The same occurs at the next energy level, $E_{3 \rightarrow 4} = E_C + E_Z(B)$. A corresponding reduction of the energy required to load the third electron follows, $E_{2 \rightarrow 3} = E_C - E_Z(B)$. In the ideal case, this will repeat up to high electron numbers, with odd-to-even electron number transitions having a positive E_Z shift and even-to-odd transitions having a negative E_Z shift. The magnitude of this shift is a fixed value $g\mu_B B$, where $g \approx 2$ in silicon and $\mu_B = 5.788 \text{ eV/T}$ is the Bohr magneton. If the magnetic field strength is known, then the shift of the addition energy ΔE_{add} for any B can be calculated. We can then relate the ΔE_{add} to the measured shift of the transition position in gate voltage, $\Delta E_{add} = \Delta V_G \alpha_G$. This can be used to extract the lever arm $\alpha_G = \Delta E_{add} / \Delta V_G$.

ΔE_{add} as a function of the magnetic field is shown in Fig 3.7. From these measurements, the lever arm of the gate at the first electron was estimated to be $\alpha_G = 0.27 \text{ eV/V}$.

This magnetospectroscopy can also be used as a "first test" to assess if a device is spin-compatible. There are many reasons why spin measurement may not be possible in a device; for example, if electrons are being loaded or read out via a spin-blockaded dopant or interstitial dot, or if the potential landscape of the quantum dot is rough and contains many potential minima, preventing sharp confinement of an electron. However, tuning a device into a regime where spin is visible is not trivial and can be a lengthy process. Magnetospectroscopy can be a good first test to ensure that the device is spin-compatible. If such an addition energy shift spectrum is seen, it is a good first indication that the spin can be characterized.

Electron temperature

To characterize the electron spin, it is necessary to measure at low temperature, so that $k_B T \ll E_Z$. The temperature measured at the fridge is not necessarily the same as the electron temperature at the sample, and this discrepancy can be greater at lower temperatures [check and include fridge specs - cooling power at 4k, 200mk]. To measure the electron temperature, we can measure the thermal broadening of the transition.

To do so, we measure the average current through the sensor dot as a function of the plunger gate voltage V_G across the transition of the first electron. A two-level current-time trace is obtained for each V_G , with $I_{SD,0}$ indicating that the probed dot has $N_e = 0$, and $I_{SD,1}$ indicating that it has $N_e = 1$. The average state of the probed dot is measured by

taking the mean current during a fixed integration time. The mean current is plotted as a function of the detuning V_G in Fig 3.7b. This is fitted with a fermi distribution of the form:

$$I_{SD} = I_{SD,0} + \frac{I_{SD,1}}{1 + e^{\frac{\alpha_G(V_G - V_0)}{k_B T}}} \quad (3.9)$$

Here, V_0 is the value of V_G where the transition occurs. If the lever arm is known, the electron temperature can be determined from this fit. Similarly, if the electron temperature is known, this can be used to measure the lever arm.

CHAPTER 4

Low frequency charge noise in FDSOI quantum dots

4.1 Introduction

$1/f$ noise, also called "pink noise" or "flicker noise", is one of the most ubiquitous noise types to appear in physical systems. It manifests as a noise signal with a power spectral density that follows a $1/f$ relationship. It is not only present in all metallic and semiconductor devices, but is also seen in large-scale physical [Mat86], biological [Gil95], and even economic systems [Wes89]. Whilst it is unclear why $1/f$ noise is so prevalent, and remarkably consistent across many different fields, the sources of $1/f$ noise are seemingly uncorrelated and system dependent.

$1/f$ noise has been thoroughly studied for the field of electronics, and it is typically believed that the main source of the noise is low-frequency fluctuations in the material of the device. This type of noise is often referred to as charge noise, and has been shown to be a limiting factor for quantum devices [Yon17].

Here we outline how charge noise affects a qubit defined in a semiconductor device, the physical source of the noise, and how the $1/f$ noise spectrum can be interpreted for semiconductor quantum dots. We then present experimental results, demonstrating clear $1/f$ charge noise experienced by a quantum dot in a silicon CMOS nanowire device. We demonstrate how the charge noise varies as the quantum dot is manipulated within the device and experiences different spatially separated sources of charge noise. We also explore the sources of charge noise and measure the energies of two-level fluctuators believed to be the dominant cause within the measurement regime. Finally, we present a novel method of measurement of charge noise at the single electron level, and use it to probe the charge noise experienced by the first three electrons in a quantum dot.

4.2 Background

4.2.1 Charge noise models

In quantum dots, charge noise manifests as a variation in the chemical potential due to changes in the electrostatic environment of the dot, which can come from several sources, including voltage noise on the electrostatic gates and charge fluctuations in potential wells nearby. Such noise is observed in many different implementations of quantum dot including

GaAs/AlGaAs heterostructures [Jun04], self-assembled InGaAs quantum dots [Hau14], and in Si/SiGe quantum dots [Con19]. Charge noise generally manifests as low-frequency noise that follows a $1/f$ behaviour.

$1/f$ noise has been a significant field of study in MOS devices since their conception. Throughout the latter half of the 20th century, two competing fields of thought developed to explain the $1/f$ noise behaviour observed in MOS structures, termed "mobility fluctuations" and "number fluctuations", or the $\Delta\mu$ and the ΔN models respectively [Van94]. These arose from the attempt to find a generally accepted model that encompassed the noise seen in all MOS devices. The $\Delta\mu$ model, based on the work of F. N. Hooge rejects the hypothesis that $1/f$ noise arises from surface effects, and instead relates the noise magnitude to the density of mobile charge carriers in the semiconductor bulk [Hoo94]. This model was shown to describe the noise observed in p-type devices well, which are in general low-noise compared to n-type devices [Van94]. However, the ΔN model was developed to explain the higher noise observed in n-type devices [McW57]. It posits that the noise magnitude is proportional to the density of charge traps in the oxide at the channel surface. Arguments have been made in either direction - for example, it was shown that in both n- and p-type devices the $1/f$ noise holds when the interface is highly charged and the conduction channel is far from the oxide [Li94]. Similarly, the ΔN model was used to describe the noise experienced by both n- and p-type devices at subthreshold gate voltages [Jak98].

Around the start of the 21st century, several unified models were postulated to cover all noise regimes in n- and p-type devices. Reduction in the scale of CMOS devices to ones with a gate area containing few or single electrons were able to demonstrate discrete modulation of the current in the transistor channel, which may be modelled as the random telegraph signal of a two-level charge fluctuator [Lun02]. This led to the resolution of this discussion to an extent in favour of the ΔN model due to its dominance at low number of effective traps, and we will generally consider $1/f$ noise in this context, as it best explains the noise observed at low device dimensions in n-type devices such as those studied.

4.2.2 Charge traps in silicon

Since the year 2000, there has been rapid development in the characterisation of charge noise in CMOS devices for a variety of applications. In particular, CMOS photosensors [Kon14], biosensors [Cre13] and quantum dots [Pet99] have all been studied extensively to determine the origins and effect of the prevalent $1/f$ noise. With the reduction of device dimensions below 100 nm, $1/f$ noise became a critical factor in device operations [Lee06]. Around the same time, support for the solid state spin qubit was growing, with coherent manipulation of semiconductor quantum dots demonstrated in both electrons [Pet05] and holes [Bru09]. The negative effects of charge noise on spin qubits were heavily studied [Cul09], and it was found that $1/f$ noise has a detrimental effect on spin coherence time, becoming more prevalent as the quantum dots became smaller in dimension.

Charge fluctuations in silicon arise from the charging and discharging of charge traps, which are isolated potential wells in the silicon lattice caused by the presence of lattice defects. These charge traps are considered bistable; that is to say, they have only two stable states, charged with a single electron and discharged. They are often referred to as two-level systems, or TLS.

In the McWhorter model the charge fluctuations in a semiconductor arise from the exchange of electrons between the surface layer of the semiconductor and the traps in the oxide layer at the interface. These are prominent at Si-SiO₂ interfaces, where dangling and distorted bonds act as charge trap sites. They are the result of lattice mismatch between the Si and SiO₂, and are generally unavoidable to a certain extent when fabricating the kind of devices used. A mismatch in the lattice results in a distribution of unpaired bonds at the boundary of the silicon lattice.

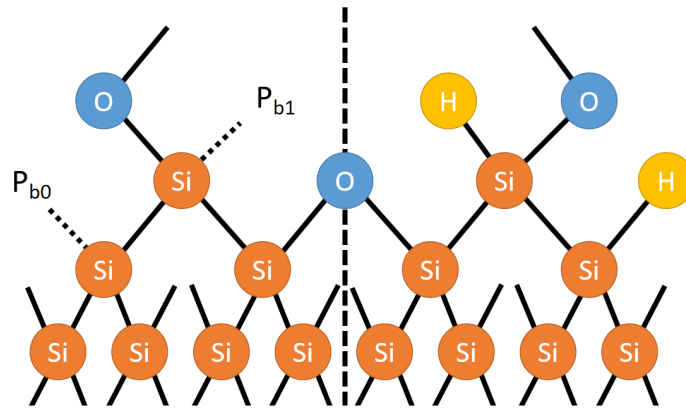


Figure 4.1: Dangling bonds example A schematic of the Si/SiO₂ interface. Left panel: two types of dangling bond commonly found at the interface due to lattice mismatch. The P_{b0} dangling bond is found at a silicon nucleus which is bonded to three adjacent silicon nuclei. The P_{b1} dangling bond is found at a silicon nucleus bonded to two adjacent silicon nuclei and one oxygen nucleus. Right panel: mirror of the left panel with dangling bonds neutralised via hydrogen passivation.

Some examples of this type of defect that are present in standard CMOS silicon wafers are the P_{b0} and P_{b1} centers, which are dangling bonds on a silicon atom at the interface. For the P_{b0} center, the Silicon atom is bonded to only three adjacent silicon atoms in the crystal lattice. For the P_{b1} center, the atom is instead bonded to two silicon atoms and one oxygen as part of the SiO₂ layer. The energy of both kinds of trap is within the silicon band gap, meaning that the charge configuration of the trap is dependent on the Fermi energy. When the trap is electrically active, the charge state of the trap can fluctuate between charged and discharged, giving rise to the two level system.

Both types of trap can lie either high or low in the band gap, meaning that they function as either donors or acceptors of electrons. The two types can be distinguished by their relative distributions of energies through the band gap [Len01]. However, they are functionally similar when considering their effect on the charge noise present in the system.

The density of the traps can be reduced via hydrogen passivation, which neutralises the unpaired electrons by forming Si-H bonds. However, even after passivation, there will be a finite number of trap sites that can contribute to the charge noise. The estimated density of charge traps at the interface D_{it} after oxidation is estimated to be approximately $D_{it} \simeq 10^{12} \text{cm}^{-2} \text{eV}^{-1}$ [Edw91]. This can be further reduced to approximately $10^{10} \text{cm}^{-2} \text{eV}^{-1}$

via hydrogen passivation. This would be considered an extremely clean interface, and approaches the trap density limit attainable with modern technology. If we assume a quantum dot of dimension 100 nm, this leads to approximately one defect in the region of the quantum dot in the ideal situation. However, due to the potentially lower quality oxide interface, we expect $\sim 10 - 100$ defects in the region of a realistic quantum dot.

The charging and discharging of these traps requires an electron to tunnel from the Fermi sea to the trap. This tunnelling has a characteristic tunnelling rate, γ_i , which has an exponential dependence on the distance from the trap to the electron source. For a two-level fluctuator, this tunnelling rate is the dominant contributor to the switching frequency. This implies that the distribution of switching frequencies is exponentially wide due to the variation in the tunnelling distance for different charge traps. For an almost continuous distribution of the tunnelling distance and barrier height across many two level systems, the distribution of switching frequencies will be proportional to $1/f$.

4.3 Low frequency charge noise in the many electron regime

4.3.1 Introduction

Here, we investigate the effects of charge noise in the simplest CMOS device that can be used as a qubit: a single "qubit" quantum dot with a nearby charge sensor. We measure the charge noise in the many-electron regime for different configurations of quantum dot to explore the effect of lateral and vertical movement within the channel. Next we conduct temperature spectroscopy to determine the source of the charge noise on the quantum dot. Finally, we demonstrate a novel method of charge noise measurement at the single-electron level, and extract the charge noise for the first few electrons in a quantum dot.

4.3.2 $1/f$ noise

Almost all noise sources in any electronic device follow a $1/f$ dependence at low frequency. This dependence arises from a coupling to a distribution of two-level systems with switching frequencies that span a broad frequency range. For a single two level fluctuator with a characteristic switching rate γ_i , then the spectral density as a function of frequency takes on a lorentzian form [Pal14]:

$$S_i(\omega) = f_i^2 L_{\gamma_i}(\omega) \quad (4.1)$$

Where f_i is the electric field influence induced by the TLS at the detector, $\omega = 2\pi f$ is the angular frequency and $L_{\gamma_i}(\omega)$ is:

$$L_{\gamma_i}(\omega) = \frac{1}{\pi} \frac{\gamma_i}{\omega^2 + \gamma_i^2} \quad (4.2)$$

A lorentzian spectrum can be an indicator of a single TLS in close proximity to the detector dominating the charge noise from other fluctuators. The McWhorter model assumes that the switching rate of the TLS is thermally activated [McW57]. Then the switching time, $\tau_i = \frac{1}{\gamma_i}$, is given by $\tau = \tau_0 e^{E_i/k_B T}$, where E_i is the activation energy of the fluctuator, τ_0 is the characteristic attempt time of the two level system, and T is the

temperature. The spectrum of a single two level system can therefore be characterized in terms of its activation energy:

$$S_i(\omega, T) = \frac{\tau_0 e^{E_i/k_B T}}{\omega^2 \tau_0^2 e^{2E_i/k_B T} + 1} \quad (4.3)$$

However, the density of charge traps indicates that it is probable that rather than one single strongly coupled TLS, we will have a distribution of many TLS in the region of the detector that all contribute to the noise spectrum. The McWhorter model assumes multiple fluctuators with a continuous distribution of activation energies $D(E)$.

$$S_\varepsilon(\omega, T) = \int \frac{\tau_0 e^{E/k_B T}}{\omega^2 \tau_0^2 e^{2E/k_B T} + 1} D(E) dE \quad (4.4)$$

If the distribution of activation energies $D(E)$ is constant, then the noise spectrum integrated across all fluctuators is proportional to $k_B T/f$. This is the typical $1/f$ noise spectrum observed in almost every physical system. Due to the assumptions made, it can be noted that a close agreement with the $1/f$ spectrum is an indicator of an approximately continuous distribution of two-level systems.

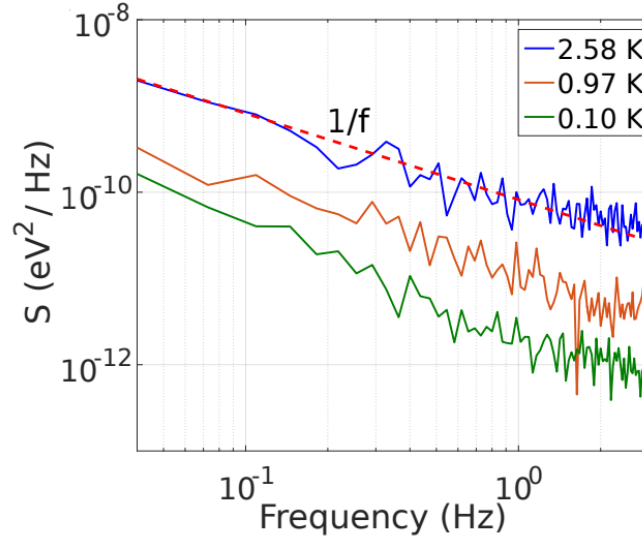


Figure 4.2: $1/f$ charge noise in an Si-MOS quantum dot Charge noise S_ε plotted as a function of frequency over a range from 10 mHz to 5 Hz for three different values of temperature. A line proportional to $1/f$ is overlaid. The charge noise spectrum can be seen to follow an approximately $1/f$ behaviour. Adapted from [Pet18].

Clear $1/f$ behaviour has been observed in Si-MOS quantum dots in the frequency range 10 mHz – 3 Hz [Pet18]. A linear temperature dependence of the charge noise at 1 Hz is seen over the range 0.1 – 4K, varying from 4 to 144 $\mu\text{eV}^2/\text{Hz}$.

$1/f$ noise has been demonstrated in Si/SiGe quantum dots [Con19] for quantum dots fabricated with varying gate oxide thickness. Charge noise values at 1 Hz were found for

these devices in the range 0.59 to 3.50 $\mu\text{eV}^2/\text{Hz}$ at 50 mK, which are among some of the lowest charge noise values reported in literature for devices of this type. A significant dependence on the thickness of the oxide was also noted, with an increase in charge noise by a factor 2 observed on devices with 46 nm thickness oxide compared to those with no gate oxide.

A low disorder MOS quantum dot was demonstrated with a low density of charge traps at cryogenic temperatures [Kim19]. This device was fabricated using high-energy processes such as electron beam lithography, demonstrating that such fabrication methods do not prevent reaching the low disorder regime. In particular, atomic layer deposition aluminium oxide (ALD) was used to protect exposed silicon oxide from contamination. The charge noise measured was around 10 $\mu\text{eV}^2/\text{Hz}$ at 300 mK, 1 Hz, which is comparable to the 46 nm oxide devices measured in [Con19].

4.3.3 Measurement method

Charge noise is manifested in a quantum dot as a modulation of the dot's chemical potential. This is similar to the case where a quantum dot is used as a charge detector. However, instead of a simple two-state current readout, the quantum dot may be coupled to many different fluctuators in its environment. This is observed in a current-time trace measured on the slope of a coulomb peak. A small fluctuation in the chemical potential of the quantum dot, $\delta\varepsilon$, induces a small current fluctuation due to the movement of a coulomb peak, δI . Therefore:

$$\delta I = \frac{dI}{dV_G} \frac{\delta\varepsilon}{\alpha} \quad (4.5)$$

Where α is the lever arm of the quantum dot, and $\frac{dI}{dV_G}$ is the slope of the coulomb peak at the measurement position. Figure 4.4a shows the current modulation as a function of time in a quantum dot system tuned to the slope of a coulomb peak. The position of maximum gradient is used to obtain the maximum sensitivity to charge fluctuations on the quantum dot site.

$$S_\varepsilon = \frac{\alpha^2 S_I}{|dI/dV_P|^2} \quad (4.6)$$

To obtain the noise spectrum in the frequency domain, the time trace is converted via fast fourier transform to a power spectrum S_I . The noise spectrum is then renormalised to obtain the noise spectrum experienced by the electrochemical potential of the quantum dot using Eqn 4.8. The resulting spectrum corresponding to the time trace seen in Figure 4.4a is plotted in Figure 4.4b. This charge noise spectrum, S_ε^2 , takes units of $\mu\text{eV}^2/\text{Hz}$. As such, charge noise values are typically quoted in these units.

Two example spectra are shown in Figure 4.4b and c. Figure 4.4b is an example of a characteristic $1/f$ spectrum. This may be fitted with a curve using $S_\varepsilon(f) = A/f^\gamma$. Here, γ is the gradient of the curve in log scale, and A is a scaling factor. However, Figure 4.4c demonstrates the other type of spectrum commonly observed. Spectra of this type are best described using a combination of the $1/f$ curve and an additional Lorentzian curve, as described in Eqn 4.9.

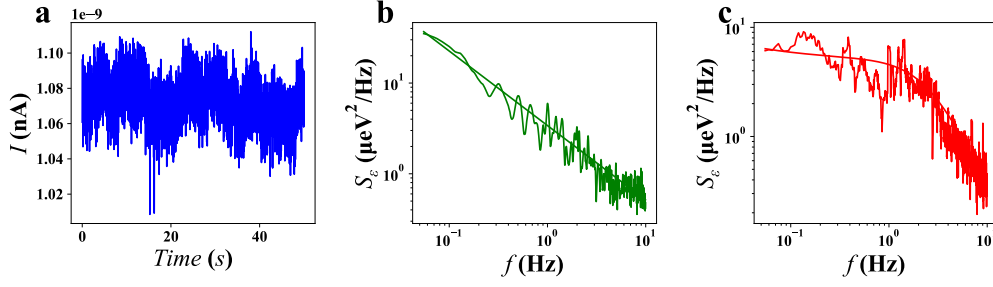


Figure 4.3: Charge Noise of a typical quantum dot. a) A current-time trace measured on the slope of a coulomb peak for a standard quantum dot configuration. The variations in the current level are due to charge noise-induced fluctuations of the quantum dot's electrochemical potential. b) A time trace similar to that seen in a) converted via fourier transform into a power density spectrum, and renormalized using Eqn 4.8. The fit is proportional to $1/f^\gamma$, with γ close to 1. c) A separate time trace which exhibits the characteristics of a single dominant fluctuator, as indicated by the Lorentzian shape. The fit parameters give the estimated characteristic frequency of this fluctuator to be 2.6 Hz.

$$S_\epsilon(f) = \frac{A}{f^\gamma} + \frac{B}{f^2/f_0^2 + 1} \quad (4.7)$$

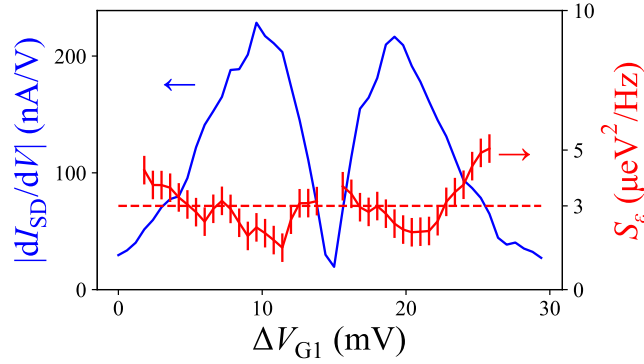


Figure 4.4: Charge noise variance across a coulomb peak. The absolute dI/dV_G , used to renormalize the charge noise, is shown in blue. The resulting charge noise at 1 Hz is shown in red. From this, the charge noise at this coulomb peak is estimated to be $3 \mu\text{V}^2/\text{Hz}$.

To obtain a precise charge noise value for a specific configuration, such measurements are made whilst varying the plunger gate voltage across a coulomb peak. This is demonstrated in Fig 4.5. The fact that the charge noise is (relatively) demonstrably constant across the coulomb peak is a good indication that the renormalization to account for the sensitivity works well. At the limits (where the dI/dV tends to zero) the sensitivity is arbitrarily low and the renormalization unreasonably large, so limits in dI/dV are set, beyond which

the charge noise results are discarded. The charge noise value for a specific configuration is then taken to be the mean of the selected data, with the error given by the standard deviation.

4.3.4 Wavefunction manipulation

The properties of a quantum dot at the single-electron level is highly dependent on the confinement of the electronic wavefunction. A quantum dot is formed through accumulation at the gate-channel interface. This interface provides the vertical confinement necessary to lift the valley spin degeneracy, which is crucial for creating a viable spin qubit. However, interface trap sites are expected to dominate low-frequency charge noise. Local manipulation of the wavefunction to minimize the effect of charge noise may therefore be necessary to operate a low-noise spin qubit. The wavefunction may be distorted via control of the electric fields within the channel. For this purpose, there are four control gates that may be used.

The plunger gate for the dot, denoted V_{B1} , is used to control the on-site chemical potential. This gate is not used to distort the wavefunction, but instead to compensate changes in the on-site potential to preserve the state of the quantum dot. The back gate, V_{BG} , acts as a universal control over the channel. By applying a positive voltage on the back gate, the quantum dot can be attracted more strongly to the back interface. With the plunger gate providing a compensation on the on-site potential, this back gate can be used to move the quantum dot vertically within the channel. Additionally, a top gate, V_{TG} is used to control the lateral extension of the wavefunction.

Vertical manipulation

Within the channel of the nanowire device, there are multiple interfaces formed at different stages of the fabrication process. The back interface is formed via growth of the polysilicon channel on the silicon oxide substrate (often referred to as the buried oxide, or BOX). Conversely, the upper and side interfaces in the region of the quantum dots are formed by deposition of the metallic gates at a later step in the process. They are separated from the channel by a native oxide of a thickness of a few nanometres.

Outside of the gate region, the metallic layer and oxide is etched away, and the interface is instead formed by the silicon nitride spacers which are deposited around the gates and protect the depletion region between the quantum dot and the reservoirs. The reservoirs are formed by ion implantation from above. The metallic gates and spacers nominally protect the channel from this implantation. However, high energy ions can penetrate the protective layer and damage the oxide. Such ions have fixed configurations in which they are active, and the quantum dot may be tuned away from such a configuration to avoid their direct effect. However, their implantation can cause damage to the silicon oxide, which manifests as charge traps, often in the form of dangling bonds. As such, it is expected that the upper interface has a higher charge trap density than the lower interface.

By probing the change in the charge noise experienced by the quantum dot as it moves vertically in the channel, we aim to identify the dominant source, and determine to what degree it is possible to screen the effects of charge noise through vertical distortion. The number of electrons in the dot is kept constant through on-site potential compensation

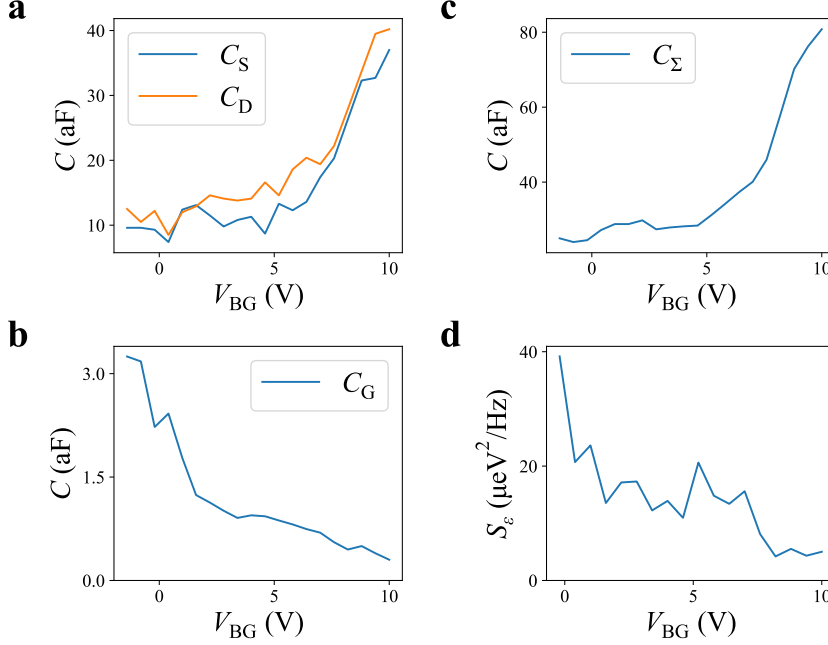


Figure 4.5: Vertical manipulation of the quantum dot in the channel. a) The capacitance between the quantum dot and the source and drain reservoirs, labelled C_S and C_D respectively. This capacitance is extracted from coulomb diamond measurements in each configuration for which the charge noise was measured. b) The capacitance between the quantum dot and the plunger gate C_G as a function of V_{BG} . The capacitance is strongly linked to the distance of the quantum dot from the gate interface. This implies a significant vertical movement within the channel, changing the capacitance to the gate by more than a factor 4. c) The self-capacitance of the quantum dot, C_Σ , as a function of the V_{BG} . C_Σ can be interpreted as a measure of the quantum dot's spatial extension. d) The charge noise measured at the quantum dot for equivalent configurations at values of V_{BG} between 0 to 10 V. Each point represents the average charge noise measured in at least twenty identical experiments.

with the plunger gate. The charging energy of the dot may be studied to determine the relevant capacitances as discussed in Chapter 2. From these capacitances, we are able to study how the charge noise responds to the manipulation of the dot.

Fig 4.6a demonstrates that the capacitance to source and drain, C_S and C_D , remains approximately constant below $V_{BG} = 5$ V, and increases sharply after 5 V. This implies that the dot only begins to undergo lateral extension after $V_{BG} = 5$ V. Fig 4.6b indicates the opposite trend for the capacitance to the plunger gate: below $V_{BG} = 5$ V, the capacitance decreases, and then remains relatively stable above 5 V. Finally, Fig 4.6c shows the self-capacitance of the dot, C_Σ , approximately follows the same trend as C_S and C_D , as the self-capacitance is dominated by these two. The self-capacitance is a qualitative measure of the size of the quantum dot, and we can therefore conclude that the dot size

is approximately constant in the first 5 V. By combining these three factors, we build a picture of the movement of the quantum dot. During the first 5 V of manipulation, the dot does not change size or lateral extension, but does move vertically in the channel, as indicated by the change in V_G . After $V_{BG} = 5$ V, the dot position remains constant, and we can conclude that it is then confined against the back interface. Increasing V_{BG} increases this confinement, and additionally the spatial extension, leading to the increase in C_S , C_D , and consequently C_Σ .

With this in mind, we can apply this movement of the dot to our study of the charge noise. The change in the power spectral density S_ϵ measured at 1 Hz is plotted as a function of V_{BG} in Fig 4.6. It can be immediately seen to match the change in C_G well, with a sharp decrease in charge noise observed below 5 V and a reduced decrease observed at higher V_{BG} . Indeed, no clear correlation with the increase in C_S and C_D is observed. Thus we conclude that the reduction in charge noise from $40 \mu\text{eV}^2/\text{Hz}$ to $4 \mu\text{eV}^2/\text{Hz}$ is due to the displacement of the dot vertically in the channel, and not the change in lateral extension. We interpret this as being due to the nominally cleaner back interface having fewer charge traps, and therefore a lower magnitude of charge noise, whilst the noisy upper interface has a reduced effect due to the reduced capacitive coupling. Due to the strong correlation between the charge noise and C_G , we conclude that it is likely that the charge noise at $V_{BG} = 0$ V is dominated by charge traps in the gate stack, which experience a similar decrease in capacitance to the dot as it moves in the channel.

However, this conclusion does not reveal any information about the effect of the lateral extension of the quantum dot. To investigate this, we therefore need to change the shape of the dot without changing its position in the channel.

Lateral manipulation

The source and drain are created by implantation of donor ions into the silicon lattice. If not fully annealed, these can also act as charge traps. Such traps are expected to have a high characteristic frequency, due to their position within the channel and close to the boundary of the electron reservoir. The density of these traps is expected to increase towards the reservoir boundary. In order to distinguish the charge noise screening effect due to the change in interface from a changed coupling to the reservoirs, we now control the lateral extension of the wavefunction without inducing vertical motion, to study the effect of the reservoirs.

In order to achieve this, we apply a positive voltage on the top gate, V_{TG} . Due to the metallic gate screening the on-site potential of the quantum dot, this positive electric field is preferentially applied at the sides of the gate. The effect of this is twofold: firstly, it is expected to extend the quantum dot wavefunction beyond the edges of the gate, without significantly affecting its on-site potential. This has been observed through stability diagram measurements at varying V_{TG} , whereby the potential of the quantum dot is only weakly modified, but the coupling to the reservoirs (measured via the tunnel rate through the dot) can be changed by more than factor 3 for a V_{TG} range of a few volts. Second, this positive electric field will attract the electron reservoirs, extending them closer to the quantum dot. This has the effect of increasing the coupling of the quantum dot to the reservoirs without affecting its vertical orientation.

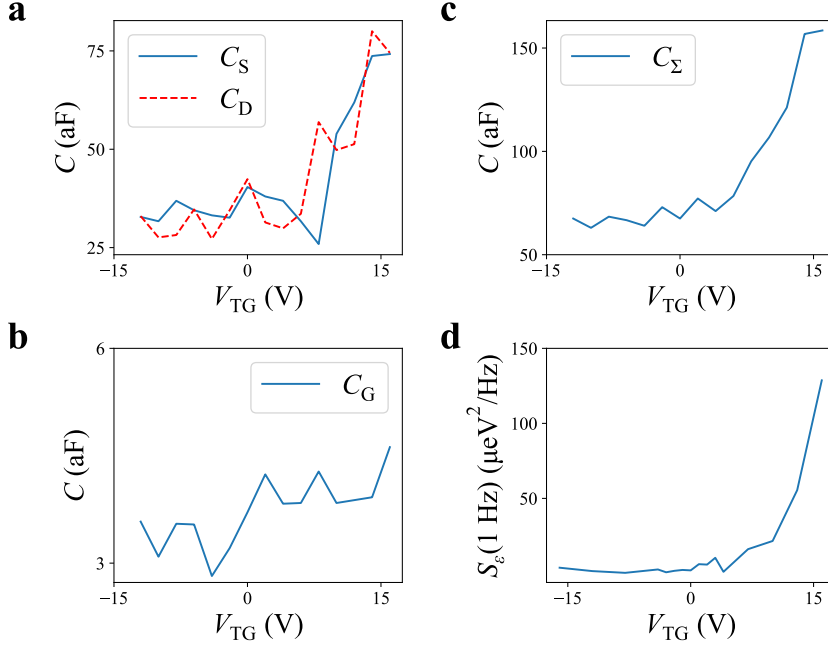


Figure 4.6: Lateral manipulation of the quantum dot in the channel. a) The capacitance between the quantum dot and the source and drain reservoirs, labelled C_S and C_D respectively, as a function of V_{TG} . b) The capacitance between the quantum dot and the plunger gate C_G as a function of V_{TG} . c) The self-capacitance of the quantum dot, C_Σ , as a function of V_{TG} . d) The charge noise measured at the quantum dot for equivalent configurations at values of V_{TG} between -15 to $+15$ V. Each point represents the average charge noise measured in at least twenty identical experiments.

We tune the quantum dot to a low-noise regime, that is, $V_{BG} = 5$ V, with an identical number of electrons loaded as in the previous measurement. The top gate is tuned over a range from -15 to $+15$ V, and the corresponding charge noise measured in an identical fashion. The number of electrons in the dot is kept constant throughout the measurement.

In Fig 4.7a, it can be seen that C_S and C_D vary similarly with the top gate voltage, increasing rapidly at positive values of V_{TG} . This implies an expansion of the quantum dot towards the reservoirs (or, equivalently, an expansion of the reservoirs towards the quantum dot). Compared to the measurement of the movement of the dot with the back gate, there is a comparatively small change in the capacitance to the plunger gate, as indicated in Fig 4.7b. Fig 4.7c indicates that there is a strong increase in the size of the quantum dot with high V_{TG} . This strongly implies that the increased C_S and C_D do indeed correspond to a lateral extension of the dot, and cannot be purely explained by movement of the reservoirs. From these factors, we can conclude that the quantum dot remains relatively constant in both size and location from $V_{TG} = -10$ V up to $V_{TG} \approx 5$ V. Between $V_{TG} \approx 5$ V and $V_{TG} = 15$ V, the dot is significantly extended laterally along the axis of the wire, bringing

it closer to the source and drain.

The change in the charge noise experienced by the dot with V_{TG} is shown in Fig 4.7d. It can be seen that the charge noise remains approximately constant until a steady increase around $V_{TG} \approx 5$ V. It then sharply increases around $V_{TG} \approx 10$ V. Due to the strong correlation with C_S and C_D , we attribute this rapid increase in charge noise to the increased lateral extension of the dot. Attributing the source of the increased charge noise is not trivial, however, as increasing the size of the dot changes many factors.

Firstly, the increased coupling to the reservoirs implies an increased coupling to charge traps at the reservoir interface. However, the positive field applied by the top gate is also likely changing the reservoir interface, which in turn changes the population of traps that are active. The increased capacitive coupling to the reservoirs implies that the reservoir fluctuators will have an increased effect on the charge noise experienced by the dot, but we cannot conclude that this is the sole contributor.

A second factor to consider is the interfaces that the dot experiences as it is extended in the channel. The dominant interface that the dot is in contact with in the absence of external fields is the gate oxide, since it is formed by accumulation at the top gate. However, the gate is surrounded by silicon nitride spacers. These spacers are expected to contain a higher density of charge traps than the gate oxide, and indeed could be one of the main contributors to charge noise. The increase in charge noise as the dot extends laterally could therefore be attributed to an increased capacitive coupling to the traps in the spacers. A strong argument for this being the case is that the increase in charge noise is not observed when varying the back gate, despite a similar increase in C_S and C_D being observed. Indeed, as the reservoir capacitances are increased, the charge noise in that case continues to decrease. This may be an indicator that applying a positive top gate voltage shifts the interface between the dot and the reservoir towards the upper interface, increasing the coupling to traps at the upper interface. The location of the dot is likely unchanged, as indicated by the constant C_G .

Finally, it is currently unclear what effect the electric field fluctuations induced by charge traps have on the tunnel coupling. The tunnel barrier is a potential barrier created by the depletion region between the quantum dot and the reservoir. Charge traps in this region could cause fluctuations in the potential of this barrier, which would manifest as changes in the current through the quantum dot. As the potential barrier is decreased (equivalently, the tunnel coupling increased, as indicated by C_S and C_G), these fluctuations in the height of the tunnel barrier will have an increased effect on the magnitude of current through the dot. Thus, the increase in charge noise may be simply due to the increased contribution from fluctuations of the barrier potential. However, in the region of maximum sensitivity, this is indistinguishable from charge noise experienced by the on-site potential, so we cannot conclude if this is a significant contributor. Additionally, the inverse trend is observed with the back gate measurement. If this is, therefore, a factor in this case, it would imply that this is due to the interface being localized towards the upper interface, which is expected to have a higher density of charge traps. Alternatively - or additionally - it may only become a factor at higher tunnel coupling. In the back gate experiment, C_S and C_D are only increased to around 40 aF, whereas a capacitance of 80 aF is reached in the top gate experiment. As the tunnel coupling increases exponentially with a linear decrease

in the width of the barrier potential, this could explain why no competing mechanism was observed in the back gate experiment.

Indeed, the increase in charge noise with V_{TG} is likely to result from a combination of these factors. Further study would be required to isolate the relevant contributions from each factor and inform future fabrication techniques with a view to minimizing charge noise. If we consider a purely capacitive model, however, our conclusion is clear that increasing the coupling to the electron reservoirs induces a significant increase in charge noise. We see a variation from $6 \mu\text{eV}^2/\text{Hz}$ in the low noise regime ($V_{BG} = 5 \text{ V}$, $V_{TG} = 0 \text{ V}$) to $129 \mu\text{eV}^2/\text{Hz}$ in the high noise regime.

4.3.5 Identifying fluctuators

Thus far we have studied the effect of charge noise on the quantum dot, but the nature of the source of this charge noise remains unresolved. Temperature spectroscopy can allow comparison of the system with various models of charge noise, and investigation of the distribution of the fluctuators. In particular, it allows comparison of this system with the Dutta-Horn model of charge noise in semiconductors, as discussed previously.

Lorentzian spectra

The theory behind $1/f$ noise relies on the assumption that the two level systems have a continuous distribution of switching frequencies and activation energies, which may not be the case in a physical system. These assumptions are largely valid for defects in bulk material, but do not hold in general for the systems considered, that being interface defects acting as two-level fluctuators, with their charge state modulated via a tunnel coupling. Rather than defects in a bulk material, these are better described as a sparse bath of discrete defects. It follows that the number of fluctuators that couple to a quantum dot can be highly dependent on the size of the dot. For a sufficiently large dot, it can be anticipated that the observed noise should tend towards a $1/f$ spectrum as it couples to a continuous distribution of fluctuators. However, for a small dot, the number of fluctuators can be small, and no longer appropriately described by a continuous distribution. This can manifest as a deviation from the typical $1/f$ relationship of charge noise with frequency, either as a clear Lorentzian signature or a change in the linear gradient at 1 Hz described by γ . For a linear fit, this is equivalent to the value of β , the frequency exponent. These kind of deviations from the standard $1/f$ spectrum have been observed in all kinds of devices. A signature that a strong deviation from a $1/f$ spectrum can indicate is the presence of a single, very strong fluctuator close to the detector.

Simulations of the two extreme cases are shown in Fig 4.8. Fig 4.8a is a simulation of a single fluctuator with a characteristic frequency of 0.1 Hz. This single fluctuator has a lorentzian-type spectrum, described in Eqn. 4.5. This is typically indicated by the preferential fitting of a curve of the form $S_I(f) = \frac{A}{f^\beta} + \frac{B}{f^2/f_0^2 + 1}$ where the scaling factor B is comparable to or greater than A . In general, a curve that is indicative of a single fluctuator dominating within a distribution will manifest as $1/f$ noise with a Lorentzian distribution superimposed, with f_0 the characteristic switching frequency of the dominant fluctuator.

Fig 4.8b is a simulation of a group of fluctuators with a flat distribution of activation

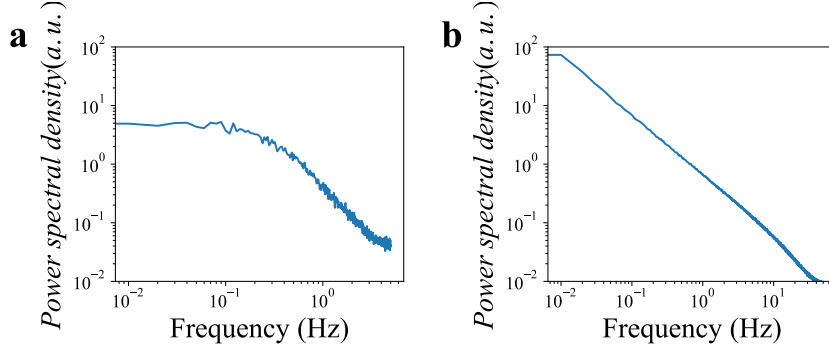


Figure 4.7: Bistable fluctuator simulations RTN simulations of two different populations of fluctuators. a) A single fluctuator with a characteristic frequency of 0.1 Hz. b) A population of fluctuators with energies distributed linearly between 1 eV and 1 μ eV.

energies. Clear $1/f$ noise is observed due to the summation of many individual lorentzian-noise fluctuators. However, it is common to see noise spectra that do not fall into either category. Spectra often follow a $1/f^\beta$ dependence, with β having a value typically between 1 and 2, with values between 1 and 1.4 most frequently reported in Silicon. This is an indication of a non-continuous activation energy distribution of fluctuators that is not dominated by one fluctuator in particular. Such a distribution $D(E)$ is described using the Dutta-Horn model:

$$S_\epsilon(\omega, T) = \frac{k_B T}{\omega} D(\tilde{E}) \quad (4.8)$$

Where $\tilde{E} = -k_B T \ln(\omega \tau_0)$. The gradient of the noise spectrum at 1 Hz is defined as $\gamma = -\partial \ln S_\epsilon / \partial \ln \omega$. If $\gamma \neq 1$, then $D(\tilde{E})$ must not be constant. A deviation from $\gamma = 1$ can therefore be used to indicate an uneven distribution of two-level systems. Charge noise in Si/SiGe quantum dots that deviated strongly from the $1/f$ spectrum was shown in [Con19] to qualitatively agree with the Dutta-Horn model of an uneven distribution of two-level systems.

Temperature dependence of charge noise

Temperature spectroscopy may be used to determine the origin of deviations from $1/f$ noise. The Dutta-Horn model, as described in eqn. 4.10, indicates that if $D(E)$ is not constant, $S_\epsilon(\omega, T)$ does not vary linearly with temperature. There is then a relationship between the noise power $S_\epsilon(\omega, T)$ and $\beta(\omega, T)$:

$$\beta(\omega, T) = 1 - \frac{1}{\ln(\omega \tau_0)} \left(\frac{\partial \ln S_\epsilon(\omega, T)}{\partial \ln T} - 1 \right) \quad (4.9)$$

Using these two relationships, the Dutta-Horn model can be applied to noise spectra obtained over a range of temperatures to determine the linearity of the distribution of activation energies $D(E)$. If S_ϵ varies linearly with temperature, then $D(E)$ can be approximated to a linear distribution. This can be interpreted by considering how the

temperature broadens the "window" of energy around the electrochemical potential that contains the active states of thermally activated fluctuators. States that lie beyond this window are either filled or emptied, and are not considered active fluctuators. If $D(E)$ is a linear distribution, as the temperature is increased, this window is similarly widened, and the number of active states will increase linearly, leading to higher charge noise (increasing S_ϵ). However, if $D(E)$ is not a linear distribution, then the increase in S_ϵ will not be linear with temperature.

A careful distinction should be made between the relationship between S_ϵ and T compared to the relationship between $S_\epsilon^{1/2}$ and T . Charge noise is often given in $\mu\text{eV}/\sqrt{\text{Hz}}$, as extracted from $S_\epsilon^{1/2}$, and as such the relationship between charge noise and temperature for a linear $D(E)$ should instead be a square root dependence. For clarity, S_ϵ is used here to discuss results from literature, adapted from the $S_\epsilon^{1/2}$ relationship if necessary.

In silicon quantum dots, [Pet18], the average charge noise has been seen to vary approximately quadratically with temperature, indicating that the system of fluctuators cannot be described simply by an even distribution of fluctuators in activation energy. The charge noise varies from $4 \mu\text{eV}^2/\text{Hz}$ to $144 \mu\text{eV}^2/\text{Hz}$ over 0.1-4 K.

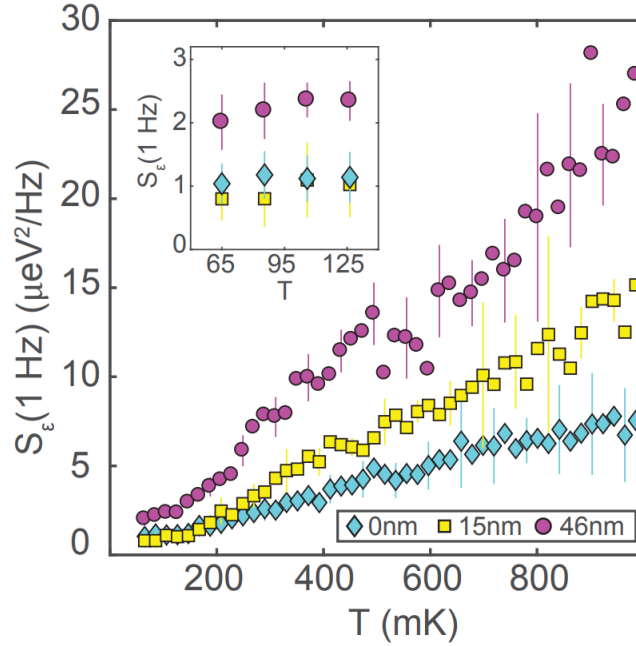


Figure 4.8: Temperature dependence of charge noise in Si/SiGe quantum dots

Temperature dependence of charge noise measured across three Si/SiGe samples with varying gate oxide thickness. Data averaged over multiple quantum dots per device. An approximately linear temperature dependence is observed for all devices. Adapted from [Con19].

A similar case is seen in Si/SiGe quantum dots, [Con19], where the charge noise follows an approximately linear temperature dependence on average, but individual dots demonstrate significant deviation. This is in qualitative agreement with the Dutta-Horn model and indicative of a non-uniform distribution of activation energies for individual quantum dots.

The magnitude of the increase with temperature was seen to be strongly dependent on the thickness of the gate oxide. The best case was seen for 0 nm gate oxide, where an increase from 1 $\mu\text{eV}^2/\text{Hz}$ to 6 $\mu\text{eV}^2/\text{Hz}$ was seen over a range from 65 mK to 1 K. A thick gate oxide of 46 nm was shown to significantly increase the charge noise, demonstrating an increase from 2 $\mu\text{eV}^2/\text{Hz}$ to 27 $\mu\text{eV}^2/\text{Hz}$ over the same range, and shown in Fig 4.9.

Deviations from such a linear dependence have been considered in [Ahn20] at the few-fluctuator level. It is shown that $1/f$ noise may be demonstrated over a broad frequency range for a single or few fluctuators when coupled to a microscopic thermal bath. The thermal bath model implies a range of phonon frequencies that may activate the two-level fluctuators in the system. This could be considered analogous to the activation energy of a single fluctuator, $E_{\alpha i}$, not being an invariable value, as assumed in the Dutta-Horn model. Instead, this energy can vary with temperature and cause deviation from the linear dependence typically seen. In this model including the microscopic sub-bath, the spectra are instead modelled using 4.12.

$$S(\omega) \approx \frac{4\tau\Delta^2}{1 + \omega^2\tau^2} + \frac{2E\Delta^2}{k_B\sigma_{sb}\omega\ln^2\frac{1}{\omega\tau}} e^{-\frac{(T-T_\omega)^2}{2\sigma_{sb}^2}} \quad (4.10)$$

The first term models the Lorentzian spectrum seen due to a single TLS, and the second term describes the influence of the sub-bath. Here the sub-bath is modelled as a portion of the thermal bath with an area A , giving a variance σ_{sb} :

$$\sigma_{sb}^2 = \frac{3\hbar^2 T}{\pi m A k_B} \quad (4.11)$$

This model was demonstrated to improve upon the raw Dutta-Horn model in describing data over a range from 50 mK to 1 K, where significant deviations from $1/f$ noise were demonstrated and non-linear temperature dependence was seen. In this case, only one or two fluctuators were needed to describe the temperature dependence seen.

We investigate the change in charge noise with temperature in the range 400 mK to 4 K. The temperature is varied by applying a voltage across a resistor which is thermally coupled to the device. It is measured via the internal thermometer of the refrigerator, which is thermally anchored at the top of the cold finger. A rise time in the thermometer is expected due to the need to thermalize the cold finger, so measurements are delayed for around 120 seconds after changing the temperature. This gives the fridge time to thermalize, and time for the electron temperature to reach equilibrium.

The quantum dot was biased into the low charge noise regime ($V_{BG} = 5$ V and $V_{TG} = 0$ V). The variation in S_ϵ with temperature can be seen in Fig 4.10 for two different numbers of electrons loaded into the dot. Such non-linear behaviour is indicative of a non-uniform distribution of fluctuators, as described by the Dutta-Horn model and outlined in Section 4.3.5. This can be seen simply by considering a single fluctuator with a Lorentzian-type noise spectrum. For frequencies below the fluctuator frequency f_i , the noise is essentially independent of frequency and appears similar to white noise in the power spectrum. It therefore decreases with temperature, as the activation frequency is increased $f_i = f_0 e^{-E_{\alpha i}/k_B T}$. However, at frequencies above the fluctuator frequency, the noise is

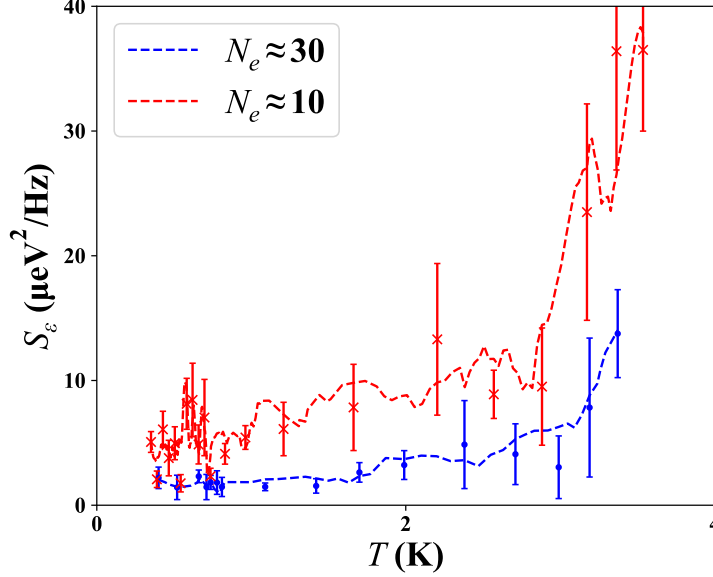


Figure 4.9: Temperature dependence of charge noise Measured for two different configurations of quantum dot. The quantum dot is biased in a low noise regime and the temperature varied from 400 mK to 4 K. Each point is representative of the average value of charge noise measured across a coulomb peak. Every 10th point is plotted with the error bar giving the standard deviation of charge noise across the peak, with the dashed line giving the moving average.

proportional to $1/f^2$ and exponentially increases with temperature.

Due to the small size of the quantum dot, it is expected that a few fluctuators dominate our noise spectrum. Taking a reasonable estimate of a few fluctuators per 100 nm² in the native oxide, the dot should experience significant coupling to only a few fluctuators. Two separate measurements were made at $N = 10$ and $N = 30$, shown in blue and red respectively in Fig 4.10. The temperature dependence in the two configurations is very similar, implying that the population of fluctuators experienced by the dot does not significantly change with its dimensions. However, the absolute value of the charge noise increases as the number of electrons in the dot decreases. This is believed to be a consequence of the decreased screening due to electron charges within the dot, meaning that a single charge fluctuation close to the quantum dot has a comparatively larger effect.

In order to support the argument that the charge noise is dominated by a few individual fluctuators, we turn to frequency analysis to identify the dominant fluctuators in the noise spectrum of the quantum dot.

Fluctuator energies

The fluctuation frequency of a thermally-activated charge trap is a temperature-dependent parameter, governed by a kinetic process. This process involves the tunnelling of an electron

into and out of the charge trap across a tunnel barrier. The frequency of fluctuator i can be described by the temperature-dependent equation given by 4.14.

$$f_i = f_0 e^{-\frac{E_{\alpha i}}{k_B T}} \quad (4.12)$$

Here, the parameter $E_{\alpha i}$, the characteristic activation energy of fluctuator i , is considered to be fixed. Additionally, the maximum attempt frequency, f_0 , is a characteristic property of the system, and can be considered an experimentally-derived fit parameter. It follows that the frequency of a single fluctuator is only dependent on the fluctuator energy and the temperature of the system. Similarly, the energy of a fluctuator may be described by 4.15.

$$E_{\alpha i} = -k_B T \ln \left(\frac{f_i}{f_0} \right) \quad (4.13)$$

Furthermore, by varying the temperature of the system, we can probe the energy of individual dominant fluctuators. When we change the temperature, we change the frequency of the fluctuators within the system, as indicated by Eqn 4.14. If we therefore fix our measurement frequency, we can vary the fluctuators that dominate at this frequency by varying the temperature. This allows exploration of a wider frequency range than would normally be possible with an experiment of reasonable duration at a fixed temperature. Additionally, it allows us to identify if fluctuators have a static activation energy, or if it changes with temperature as in the sub-bath model discussed in Section 4.3.5. We have shown that it is possible to identify a single dominant fluctuator by its characteristic Lorentzian spectrum. Through analysis of the spectra, it is possible to identify the fluctuators that dominate in a particular frequency range at different temperatures. This was done using two methods.

Spectral analysis

The first method uses analysis of the gradient of noise spectra at 1 Hz over a range of temperatures to reconstruct a spectrum in energy space that covers all values of f and T in a set window. Similarly to previous measurements, the noise spectrum for a quantum dot is extracted from current-time measurements with a duration and interval corresponding to the frequency space from 0.01 Hz to 50 Hz. The gradient of this spectrum at 1 Hz is extracted from many spectra across a coulomb peak and averaged. This gives the mean γ value for a particular configuration and temperature. The quantum dot configuration is kept constant through drift compensation using the plunger gate. The γ value is therefore extracted for a range of temperatures from 400 mK to 4 K in a consistent dot configuration. This is plotted in Figure 4.11a.

From the γ distribution, it is possible to re-create the distribution of S_ϵ in terms of $E_\alpha = -k_B T \ln \frac{f}{f_0}$. This can be seen in Fig 4.11b. This captures more precisely the noise spectrum over the full temperature range probed. For an ideal system where $D(E)$ does not vary with temperature, this approximate spectrum may be obtained directly through measurement at 4 K. However if this is not the case, which is apparent in our system due to the nonlinearity of S_ϵ with T , this method captures a variant $D(E)$, and allows us to predict the S_ϵ for any combination of f and T in this system. Additionally, through fitting

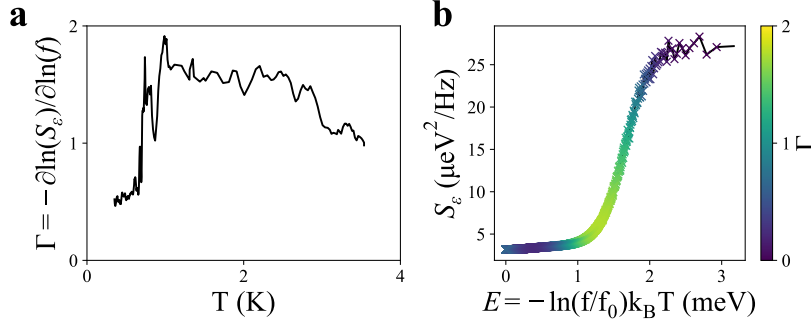


Figure 4.10: Variation of γ with temperature a) The average gradient of the noise spectrum at 1 Hz for temperature values between 400 mK and 4 K. b) A temperature-independent reconstruction of the averaged noise spectrum across the range of energies probed between 400 mK and 4 K. This is simulated by considering three crucial points in plot a). Firstly, the temperature at which the γ value changes from 0.5 towards 2 indicates a fluctuator with an activation energy that gives a characteristic fluctuation frequency of 1 Hz at this temperature. Secondly, the rate of decay towards 1 implies a lower frequency distribution of fluctuators that recover the $1/f$ curve but are low enough in frequency that they do not prevent the first fluctuator from dominating at lower temperatures. Thirdly, the initial value of γ , 0.5, implies a lack of a distribution of fluctuators in this energy range, which would tend the curve towards a γ value of 1. This gives a lower bound for a higher frequency cluster of fluctuators.

a double Lorentzian curve to this spectrum, we are able to obtain the activation energies of two primary fluctuators to a reasonable degree of precision. The fit gives these fluctuator energies to be 1.35 meV and 0.30 meV. The clear step in Fig 4.11a gives a good estimate for the lower frequency fluctuator at 1.35 meV. The other fluctuator has its characteristic frequency beyond the energy range probed and therefore we can only estimate its lower bound. In order to obtain a better estimate of the nature of these two fluctuators, we can turn to characteristic frequency analysis.

Characteristic frequency analysis

It is likely that, as the gate voltage is swept across a coulomb peak, a slightly different population of fluctuators may be activated. This manifests as a variation in the γ value in measurements across a coulomb peak for the same configuration. The variation can be as high as up to 0.4 across a peak. As discussed previously, the γ value can be an indication of the presence and central frequency of a dominant fluctuator. Since the fluctuator frequency f_i is a temperature dependent parameter, extracting the activation energy $E_{\alpha i}$ of a fluctuator directly requires varying the temperature. We also generally require the f_0 , often called the maximum attempt frequency. This parameter is usually derived from fitting the Dutta-Horn model. However, the accuracy of the value of f_0 obtained is unclear. For example, in [Con19], a maximum attempt frequency of only 5 Hz was obtained, with clear $1/f$ -like noise demonstrated to frequencies significantly greater than 5 Hz. In all calculations, we have assumed an f_0 of 200 Hz, as this is the maximum frequency to which $1/f$ noise was seen. However, it is possible to extract the energy of a

trap using an Arrhenius diagram [Bou17]. This method is independent of f_0 , allowing to extract the activation energy $E_{\alpha i}$ of individual fluctuators.

$$\log(\tau_i T^2) = E_{\alpha i} \frac{q}{k_B T} \quad (4.14)$$

To this end, we use Eqn 4.16 to plot the characteristic time constant ($\tau_i = 1/2\pi f$) of spectra that demonstrate single fluctuator behaviour. These spectra were selected from the data collected for the temperature spectroscopy experiment through analysis of the lorentzian fit covariance. Note that the characteristic lorentzian spectra are not identical across a coulomb peak. This results in a spread of results for each value of temperature, from which qualitative linear fits are made to extract fluctuator activation energies, as plotted in Fig 4.12.

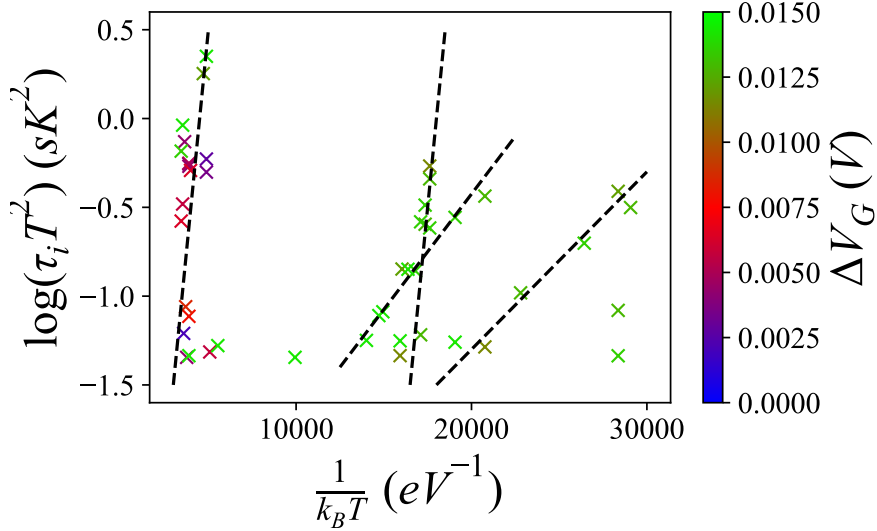


Figure 4.11: Arrhenius plot for individual fluctuator spectra Each point of data represents a single measured spectrum with a good lorentzian fit. The position of the spectrum across the coulomb peak is given in the colour scale, meaning similarly coloured points are likely to correspond to a fluctuator active at the same gate voltage. Qualitative fits are plotted as dashed black lines using the colour scale as a guide. A single fluctuator will exhibit a linear dependence on $1/k_B T$ with the gradient giving the activation energy E_{α} .

The gradient of these linear fits gives the $E_{\alpha i}$ for individual fluctuators. Four main qualitative fits give values for E_{α} that are in good agreement with two activation energies: 1.25 meV and 0.13 meV. These energies are in good agreement with those derived from spectral analysis, and can explain the variation in γ with temperature as seen in Fig 4.11a. We therefore conclude that there is a high probability that two fluctuators (or species of fluctuators) with activation energies close to 1.3 meV and 0.2 meV dominate the charge noise of the quantum dot in this temperature range.

4.4 Low frequency charge noise in the few-electron regime

Thus far, we have concerned ourselves with understanding the nature of the charge noise in the region of the quantum dot. However, this must be considered in context with the goal of the silicon quantum project. This is to enable creation of a reproducible, high-fidelity, controllable qubit at the nanoscale. In order to fulfil these requirements for a spin qubit, it is important to understand the main sources of decoherence, so that these can be minimized and the spin state preserved. To this end, the characterization of the charge noise in the many-electron regime is of limited value. Many sources in literature report low values of charge noise below $10 \mu\text{eV}^2/\text{Hz}$, but these are measured in quantum dots containing tens of electrons. Such values have little relevance or meaning at the single electron level, and robust characterization of charge noise at the few-electron level is lacking. In this chapter, we aim to fill this gap in the literature by characterizing the charge noise of a single electron using a simple and reliable method requiring only a nearby charge detector.

Currently, charge noise at the single-electron level is measured through Ramsay spectroscopy [Yon17]. However, this necessitates measurement and control of the single electron spin before measuring the charge noise. Here we present a simpler method which can be achieved using only a charge detector and weakly tunable tunnel coupling. We propose that this method could be used for rapid characterization of charge noise in single-electron quantum dots in different kinds of systems without significant requirements. We present measurements of the charge noise experienced by a quantum dot containing one, two and three electrons. The noise magnitude is extrapolated to the values obtained for tens of electrons. Finally, the spin-spin relaxation time in a charge noise limited regime under a typical magnetic field gradient is calculated and compared to the limitation of T_2^* in natural silicon.

4.4.1 Measurement method

A difficulty is encountered in attempting to directly measure charge noise at the single electron level in semiconductor quantum dots. The current passing through a single electron quantum dot is negligibly small, and well below the noise floor of our detectors. Additionally, as in this system, at the single-electron level, the electric potential of the channel can be dominated by geometric defects rather than electric field-induced potentials. Such defects, as well as poor definition of the quantum dot potential well, may contribute to resistive paths that divert current away from the quantum dot. This is the case in our system, where at the single-electron level, electrons are loaded into the dot via the adjacent dot, indicated by the high level of control over the tunnel rate with the occupation of the adjacent dot.

Therefore, we choose to use this dot, $T1$, as both charge sensor and reservoir. Fig 4.13a shows the transition of the first electron. The electron occupancies in each dot are labelled in the form (N_{T1}, N_{B1}) (the values for dot $T1$ are approximate and estimated from larger stability diagrams). No further transitions are seen at lower values of V_{B1} . This is a strong indicator that we do indeed see the transition of the first electron.

Stochastic charge transition events may be seen when the measurement rate is on the order of the tunnelling rate. This may be seen in Fig 4.13b. At the transition, electrons may tunnel into and out of the quantum dot. This is the degeneracy region, where the dot

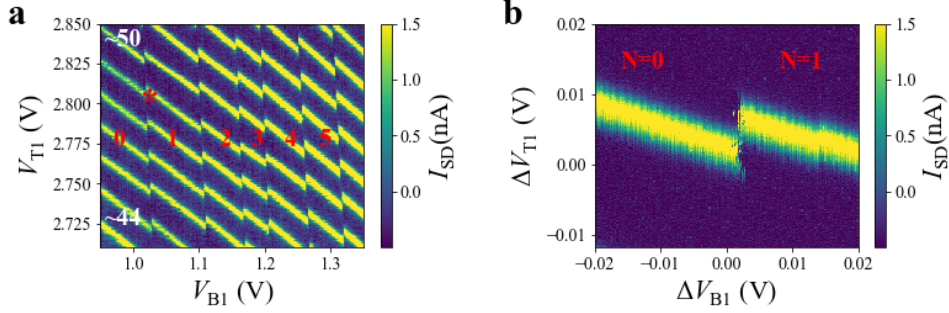


Figure 4.12: Transition of the first electron The colour scale indicates the current passing through the sensor dot $T1$. a) A wide-scale stability diagram demonstrating charge sensing of the first few electron transitions into dot $B1$. The bright coulomb peaks seen here correspond to energy levels in resonance in dot $T1$. The breaks in the coulomb peaks correspond to the capacitive shift caused by an electron loading into $B1$. The approximate occupancy of $B1$ is labelled in white, and the absolute occupancy of $B1$ labelled in red. b) A close-up stability diagram of the transition marked by $\star$ in a). Stochastic events can be seen by sharp switches between low and high current in the region of the transition. These correspond to electrons entering and leaving the dot as it lies in resonance with the reservoir.

$B1$ fluctuates between states $N_{B1} = 0$ and $N_{B1} = 1$. The frequency of this tunnelling is given by $\Gamma_{N=0 \leftrightarrow 1}$, and it is seen when the quantum dot is in the configuration indicated by the schematic in Fig 4.14a. If both dots are tuned to this region, we can obtain a current-time trace indicating the charge occupancy of $B1$ over time. This is shown in Fig 4.14b. The characteristic signature of such a trace is the two current levels, each of which indicates an occupancy state, with sharp transitions between the two. These two current levels can be assigned an occupancy state, which in this case is $N = 0$ for $I_{SD} > 5$ nA and $N = 1$ for $I_{SD} < 5$ nA. A threshold is defined at $I_{SD} = 5$ nA (dashed line in Fig 4.14b) to distinguish between the two states, and the current-time trace is digitized. Every transition between states, i.e. every time the trace crosses the threshold, is caused by an electron tunnelling onto or off the dot $B1$. Therefore by measuring such a current-time trace for a sufficient time, we may extract the tunnel rate as $\Gamma = N_{tr}/t_{meas}$ where N_{tr} is the number of tunnelling events and t_{meas} is the measurement duration.

The tunnel rate Γ will be the parameter used to measure the fluctuations induced by charge noise. In order to properly normalize this, we need to know the sensitivity of the detector, and the lever arm of the gate. The simpler of these to extract is the detector sensitivity. This can be obtained by reconstructing a tunnel rate coulomb peak, where Γ is measured as a function of V_{B1} . This peak is shown in Fig 4.14d. The sensitivity of the tunnel rate to fluctuations in the on-site potential of the quantum dot, $d\Gamma/dV_{B1}$, is extracted from this peak.

Extracting the lever arm of the gate is done via magnetospectroscopy since is not possible to use the typical coulomb diamonds method at the single electron level. Instead we investigate the change in addition energy with magnetic field. A more detailed explanation of this procedure is given in section 3.2.3. In brief, the addition energy, $E_{add} = E_C + \Delta E_Z$,

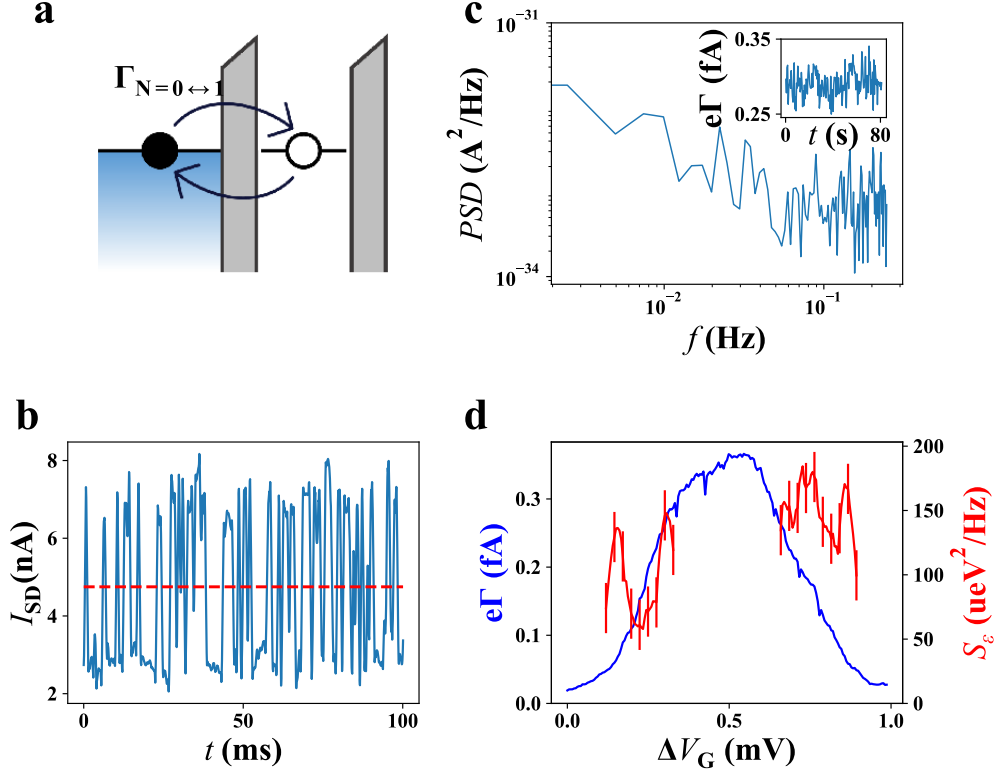


Figure 4.13: Measurement method for single electron charge noise (a) Schematic of the quantum dot configuration to be probed. The first available energy level in the quantum dot is brought into resonance with the reservoir potential. Electrons may tunnel through the barrier with a tunnel rate $\Gamma_{N_e=0 \leftrightarrow 1}$. The occupancy of the dot varies between $N_e = 0 \leftrightarrow N_e = 1$. (b) Current through the sensor dot measured over 100 ms (blue). Changes in the occupancy of the probed dot induces a capacitive shift in the sensor dot, switching the current between two states $N_e = 0$ and $N_e = 1$. A threshold may be defined (red) to distinguish between the two states. The number of switches between these two states gives the tunnel rate through the dot over 100 ms (c) Power spectral density extracted from the variation in tunnel rate over time for the probed dot. The tunnel rate is measured as indicated in (b) over ≈ 80 s (inset). (d) The power spectral density is measured whilst varying the plunger gate voltage on the probed dot. A coulomb peak (not visible in DC measurement) is reconstructed (blue) from the average tunnel rate value at each V_G . At each point, the $|de\Gamma/dV_G|$ is used to renormalize the measured PSD spectrum. The value at 1 Hz is extracted by extrapolation of a linear fit to 1 Hz and plotted as a function of V_G (red). Where $|de\Gamma/dV_G|$ is low, the PSD value is discarded to prevent over-renormalization.

changes with magnetic field, as $\Delta E_Z = \mu_B B g$, with $g \approx 2$ and μ_B the electron g-factor and the Bohr magneton respectively, both fixed parameters, is only dependent on B . Using this and the corresponding voltage shift, we can extract the lever arm at the first electron, which for our system is $\alpha_{N=1} = 0.27 \text{ eV/V}$.

$$S_\varepsilon(f, T) = S_{\Gamma_e}(f, T) \frac{\alpha_{N=1}}{\frac{d\Gamma}{dV_{B1}}} \quad (4.15)$$

$S_{\Gamma_e}(f, T)$ is extracted through measurement of variation in the tunnel rate with time. The tunnel rate is measured as described previously. The measurement duration is 1 s at 1 kHz acquisition rate. This is then measured over the course of several minutes to obtain a Γ_e time trace as seen in Fig 4.14c (inset). We then take the fourier transform of this time trace to obtain $S_{\Gamma_e}(f, T)$, Fig 4.14c. This is renormalized to obtain $S_\varepsilon(f, T)$ using Eqn 4.17. The duration of the tunnel rate measurement gives the upper bound for the frequency window which this spectrum can be measured for. In order to obtain the value for S_ε at 1 Hz, we extrapolate a linear fit to the spectrum. We repeat this process for many points across the tunnelling peak in order to extract an average charge noise, as demonstrated in Fig 4.14d. Whilst the experimental requirements are not trivial, this method is significantly simpler and faster than the current methods of measuring charge noise at the single electron level.

4.4.2 Charge noise at the single electron level

Examples of typical $1/f$ -like spectra measured using this method are shown in Fig 4.14c. For each spectrum, a linear fit is plotted and extrapolated to 1 Hz. It should be noted that many spectra measured in this way have a γ value below 1. This is indicative of higher frequency fluctuators (which operate outside the measurement bandwidth) dominating the charge noise in this regime. To obtain a precise charge noise value, up to 100 measurements are made across a tunnelling peak. The measured peak is indicated in Fig 4.14d. The shape of the peak is extracted from Γ_e -averaged time traces during the measurement. This gives a precise $d\Gamma_e/dV_{B1}$ value, which is the most important parameter to ensure accurate normalization. At each value of V_{B1} , the variation in the tunnel rate with time is measured, and the noise value S_ε calculated. As before, measurements where $d\Gamma_e/dV_{B1}$ is low are discarded, to avoid over-renormalization.

The calculated S_ε values are shown in Fig 4.14d. These give an average charge noise value of $(130 \pm 60) \mu\text{eV}^2/\text{Hz}$ at 1 Hz at 400 mK. This value means that the quantum dot experiences charge noise of a magnitude comparable to the electron temperature during 1 s.

Further measurements were conducted at 4 K and 200 mK using an identical method at the same charge transition. We measure $S_\varepsilon(1 \text{ Hz}, 4 \text{ K})$ to be $(1310 \pm 318) \mu\text{eV}^2/\text{Hz}$, which is approximately ten times greater than $S_\varepsilon(1 \text{ Hz}, 400 \text{ mK})$. This suggests an approximately linear increase in charge noise with temperature. It is also approaching the Zeeman energy for a device in a magnetic field of 1 T, which is a realistic value that may be used in future devices. Charge noise contributes to spin decoherence through several spin-to-charge mechanisms, including spin-orbit coupling and motion within a magnetic field gradient. Additionally, measurement fidelity can be degraded through direct disturbance

of the quantum dot potential, which induces current shifts and may contribute to false detection events. It also causes long-term changes in the configuration of the quantum dot, meaning that to maintain spin measurements over a long period of time, the system must be constantly re-tuned. However, the value for $S_\varepsilon(1 \text{ Hz}, 200 \text{ mK})$ is reduced to $(89 \pm 28) \mu\text{eV}^2/\text{Hz}$. As seen in the temperature-dependent measurement in section 4.3.5, the change in charge noise with temperature is significantly non-linear, and at low temperatures the change also appears to be less significant.

4.4.3 Evolution of charge noise with electron number

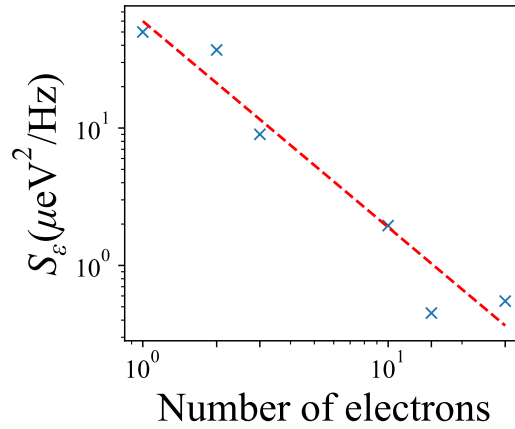


Figure 4.14: Charge noise with electron occupancy The reduction in charge noise as the number of electrons in the dot increases is plotted in log scale. The first three points, corresponding to $N_e = 1, 2, 3$, were measured using the tunnel rate method outlined in this section. The values for $N_e = 10, 15, 30$ were measured in the many-electron regime using the standard method as outlined in section 4.3.3. The red dashed line overlay is proportional to $1/N_e^{3/2}$.

Spin measurements are expected to most commonly be conducted at the first electron. However certain schema involve the use of quantum dots with two or more electrons in the same quantum dot or via exchange coupling between two adjacent dots, such as Pauli spin blockade readout [Fog18] or singlet-triplet exchange qubits [Deh14]. Given the significant reduction in charge noise at higher electron occupancy, these are expected to be more resistant to on-site charge noise, although as discussed previously, the exchange coupling is subject to additional effects from charge noise which could prove detrimental. Additionally, it has been shown that singlet-triplet qubits can have a higher spin-orbit coupling than expected in silicon [HC19], which implies a higher effect from charge noise. To quantify the effect of adding one or two additional electrons to the quantum dot, charge noise measurements were conducted at $N_{B1} = 1 \rightarrow 2$ and $N_{B1} = 2 \rightarrow 3$, as shown in Fig 4.15. These experiments were performed at 200 mK. The charge noise is reduced by an order of magnitude from $N = 1$ to $N = 3$. Values of $(37 \pm 15) \mu\text{eV}^2/\text{Hz}$ for $N = 2$ and

$(9 \pm 5) \mu\text{eV}^2/\text{Hz}$ for $N = 3$ were obtained from 100 measurements across the transition. We note that the error in these measurements is significantly higher than for the many-electron regime. This may be indicative of low frequency fluctuations of the tunnel rate Γ . However, it can also be seen in Fig 4.14d that the variations have an apparent periodicity, which could correspond to increased tunnelling when certain excited states become available with increasing V_G and offer additional tunnelling pathways. It may also be indicative that our method of renormalization of the tunnel rate is incomplete, and requires additional terms to be able to properly compare measurements across the coulomb peak. Also included in Fig 4.15 are extrapolations of data taken at 400 mK for $N = 10, 15, 30$ using the standard method of measurement without a charge detector dot. The charge noise appears to decrease following a trend proportional to $1/N_e^{3/2}$.

The mechanism responsible for such a reduction in charge noise is unclear. Intuition suggests that electric field screening could be responsible for the general reduction in charge noise with electron number. Electric field screening is an effective damping of electric fields in a conducting material by mobile charges. It follows that increasing the occupancy of the quantum dot increases the effective screening. Potential field simulations led by Michal and Niquet, however, suggest that the screening is a second-order effect. They find that expansion of the electron wavefunction towards a fluctuator has a greater effect than screening in the general case. This is a first order effect, induced by the increased spatial extension of the electron wavefunction. The only situation where the opposite is seen is when the fluctuator is within the quantum dot, in which case the first-order expansion reduces the charge noise experienced by the quantum dot. Simulations therefore suggest an opposite trend to that observed in experiment. However, the effect of interfaces and fluctuator positioning on the first-order expansion is not obvious, nor is the extent and directionality of the quantum dot expansion in the experimental device well-known. Further simulations will be required to explain the dependency seen in Fig 4.15. Additionally, comprehensive testing of similar devices would aid in understanding if the trends seen here are due to device-specific phenomena, or characteristic of nanowire devices in general.

Charge noise induced decoherence in quantum dots

Charge noise can cause decoherence of a spin qubit defined within the quantum dot via two main mechanisms: the spin-orbit interaction, and coupling to a magnetic field gradient. The spin-orbit interaction is a relativistic interaction between its spin and its orbit within the potential well, and acts to couple the electron spin to the electrical component of the electromagnetic field. While it is typically weak in silicon, the spin-orbit coupling can be significant at the interface, where the inversion symmetry is broken [Ber14]. It has also been shown to be strongly dependent on the magnetic field direction [Tan19], with a possibility of turning off the spin-orbit interaction altogether, but also with potentially strong spin-orbit coupling when the magnetic field is not precisely controlled. The Stark shift induced by electric fields can provide a further source of dephasing as it modulates the electron energy. Since it is a coupling between the energy and the electric field, electric field modulations in the form of charge noise will induce dephasing. Additionally, the Stark shift modulates the electron g-factor, which changes the qubit frequency. This causes dephasing during active operation.

Another process by which the spin can couple to the electric field is through experiencing a magnetic field gradient. This gradient is typically created by a micromagnet, and causes different spatial locations within the dot to experience a different magnetic field. Through electric field manipulation of the electron within the quantum dot, electric field fluctuations induce a magnetic field variation. This can be used for spin manipulation, by applying an oscillating electric field with a frequency equivalent to the Zeeman splitting E_Z such that the electron experiences an AC magnetic field at its resonant frequency. This method of spin control is commonly known as electron dipole spin resonance, or EDSR. However, uncontrolled electric field fluctuations can have a similar effect in the presence of a magnetic field gradient. Through this coupling to the magnetic field, electric field variations can induce stochastic spin decoherence and errors in EDSR-controlled manipulations.

For both cases, it can be shown that a modulation of the field can cause a loss of coherence in a qubit via stochastic degradation of the spin state (for a full derivation refer to [Pal14]). Here, decoherence manifests as a perturbation to an arbitrary field $\mathbf{F}$, which can be the sum of a controlled part, $\mathbf{F}_0$ and a fluctuating field $\mathbf{f}(t)$, which represents the noise. This distinction is purely one between "stochastic" and "non-stochastic"; the controlled field $\mathbf{F}_0$ may also be dynamic and may evolve with time in a controlled manner, whilst the fluctuating part is a random and uncontrolled perturbation. This perturbation prevents deterministic coherent evolution of the qubit, giving it a finite coherence time. As the magnitude of the field fluctuations increases, the coherence time will correspondingly decrease. This ultimately leads to a relationship between the pure dephasing time T_2^* and the noise spectrum $S_f(\omega)$ when the measurement time t is assumed to be long:

$$T_2^* = \frac{1}{\pi S_f(0)} \quad (4.16)$$

It can immediately be seen that this is a strange result, as for a $1/f$ noise spectrum the value of $S_f(0)$ approaches infinity, and therefore T_2^* decays to 0 for a long measurement time. To date, experiments have been conducted up to 1×10^{-6} Hz [Cal74], and the $1/f$ noise has been seen to remain consistent, implying that this relationship does not fully describe the effect of $1/f$ noise on the dephasing at low frequency. We can, however, interpret this as the low-frequency noise dominating the pure dephasing time for long-period measurements. T_2^* of 55 ns for a measurement time of ≈ 1 hour was measured in the presence of nuclear spin noise in natural Si [Pla12]. T_2^* of 18 μ s for a measurement time of 10 min was measured in the presence of charge noise of magnitude $0.22 \mu\text{eV}^2/\text{Hz}$ at 1 Hz in isotopically purified Si/SiGe [Str20], reducing to $T_2^* = 7 \mu$ s at long measurement time (6 h). This result from [Str20] is significant, as the monotonic reduction in T_2^* at long measurement time can be simply explained if it is due to noise proportional to $1/f$. This indicates that charge noise may be the limiting factor for coherence in spin qubits in silicon once the isotopic purification has eliminated the nuclear spin noise.

Another crucial effect that charge noise has which should be considered is the effect on the coupling between two adjacent qubits. As seen previously, the coupling between two quantum dots is dependent on both the tunnel coupling t_c , and the energy difference

between the two dots ΔE . t_c is strongly dependent on V_B , the barrier voltage. Both V_B and ΔE will be affected by charge noise, and therefore the effect on the exchange energy J can be more significant than for a single spin:

$$dJ = \left(\frac{2J}{t_c} \frac{dt_c}{dV_B} \right) dV_B - \left(\frac{J}{\Delta E} \right) d\Delta E \quad (4.17)$$

This dJ induces an error in any operation performed between the two qubits. For example, a SWAP gate corresponds to a π pulse, the duration of which is determined by the exchange energy J . A stochastic modulation of J via charge noise through the duration of the gate pulse will contribute an error to the angle of rotation, such that the pulse applied is instead $\pi + \delta$, with δ the gate rotation error. Then the resulting state will no longer be a clean exchange of single-qubit states, but will have a finite entangled component with a magnitude dependent on the magnitude of the error δ . Measurements made on the resulting state will therefore have a finite probability to project to the wrong state.

In addition, when the exchange coupling is on for long periods of time during the manipulation, the modulation of J will induce some dephasing to the spins. This contribution to the degradation of the two-qubit fidelity can be significant, especially when J is large ($dJ/dV = dJ/dV_B + dJ/d\Delta E \simeq 1$). In a physical system, it is more likely to see $dJ/dV \simeq 0.01$, in which case the T_2^* will be on the order of ten times the charge relaxation time T_1 . This suggests that charge noise may be one of the most significant sources of decoherence for exchange-coupled spin qubits in semiconductors, both during manipulation via gate errors, and during static operation due to the time-sensitive dephasing when the exchange coupling is active [Hu06].

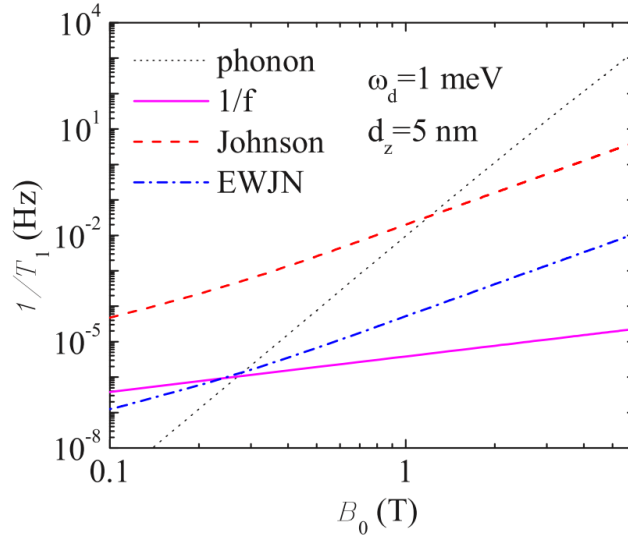


Figure 4.15: Relaxation sources Sources of spin relaxation in gate-defined silicon quantum dots at 150 mK. Adapted from [Hua14a].

Finally, charge noise can also be a source of spin relaxation [Hua14a] at long measurement time and low magnetic field. Simulations of relaxation sources in gate-defined silicon quantum dots are shown in Fig 4.2. It can be seen that below $B = 1$ T, Johnson (charge) noise starts to dominate the relaxation processes. In silicon, due to the large valley splitting, intravalley spin-orbit mixing dominates the spin relaxation [Yan13]. In the low-field regime the T_1 relaxation time is very long (tens of seconds). However, T_1 provides an upper bound for the spin coherence, and this further underpins the need to understand and minimize the effect of charge noise.

Charge noise was found to be the limiting factor in the dephasing time T_2^* for spin qubits in Si/SiGe in the presence of a magnetic field gradient [Yon17]. The charge noise here is observed to follow the $1/f$ law in excellent agreement, obtaining $\gamma = 1.01 \pm 0.05$. The noise was measured in two main frequency ranges, between $0.01 - 1$ Hz and around 10 kHz, with the extrapolation of each in close correlation to the other. An implication of this is that the $1/f$ charge noise dominates the noise in this device over a range from mHz to kHz, leading to a predicted value of T_2^* at $25 \mu\text{s}$.

Prediction of T_2^* in the presence of a magnetic field gradient

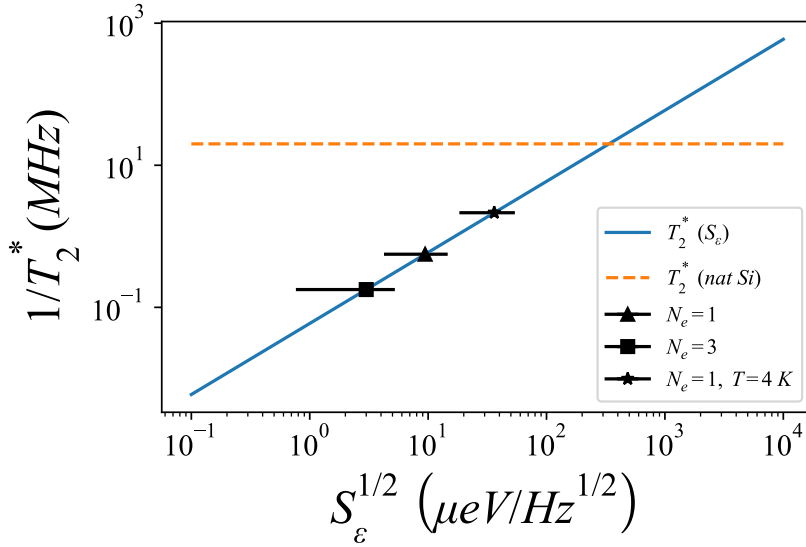


Figure 4.16: T_2^* simulation The increase in the dephasing frequency with increasing charge noise for a measurement duration of 1 s. The dashed line is the dephasing frequency due to the nuclear spin noise in natural silicon ($T_2^* = 0.05 \mu\text{s}$). The dephasing frequency limits given by the charge noise values measured for the first few electrons are indicated.

Here we simulate the effect of the charge noise experienced by our system on the T_2^* in the presence of a typical magnetic gradient required for EDSR. Fig 4.16 shows the simulated effect of charge noise on the T_2^* for $t_{meas} = 1$ s. The T_2^* limit due to the nuclear spin noise is $0.05 \mu\text{s}$ [Pla12].

The T_2^* is calculated as the dephasing due to the motion of the electron in a magnetic

field gradient along the spin quantization axis. We use a gradient of $dB_z/dz = 0.35 \text{ mT/nm}$, which is approaching the upper bound of typical magnetic field gradients in similar systems. The charge noise is modelled as a perturbation to the applied gate voltage, $dV_{rms} \propto S_\epsilon^{1/2}$. This can be considered to be the maximal effect the charge noise can have on the coherence, as it implies the displacement of the dot r_{rms} is purely along $\hat{z}$. It is assumed that for a small potential perturbation dV_{rms} the displacement r_{rms} is linear, and that for $dV_{rms} = 1 \text{ } \mu\text{eV}$, $r_{rms} = 1 \text{ pm}$, a value extracted from potential simulations of a similar device.

We calculate a T_2^* of $10.9 \text{ } \mu\text{s}$ and $34.5 \text{ } \mu\text{s}$ for $N_e = 1, 3$ respectively. This is more than two orders of magnitude longer than the T_2^* for natural silicon, and comparable to the $T_2^* \approx 20 \text{ } \mu\text{s}$ measured in highly purified Si/SiGe [Str20]. This is for the "worst case" and indicates a lower bound for T_2^* , as the charge noise is assumed to induce displacement in the dot along the quantization axis. It is a reasonable approximation, as the interface where the charge traps are expected to be prominent is located along the z direction. However, in a real system it is probable that most charge fluctuations will have a non-zero component orthogonal to the z axis relative to the quantum dot, and therefore the true magnitude of the z -axis fluctuations will be lower.

It demonstrates that we are not currently charge noise limited in terms of coherence time, and that it would be beneficial to move to isotopically purified silicon to extend the T_2^* as far as possible. The $1/f$ dependence of the T_2^* seen in [Str20] is a good indication that we would be charge noise limited after purification. The same holds true at higher operational temperatures. The measured value of $(1310 \pm 318) \text{ } \mu\text{eV}^2/\text{Hz}$ at 4 K yields a T_2^* of $2.9 \text{ } \mu\text{s}$, which is still 50-60 times longer than the nuclear spin noise limited T_2^* .

4.5 Conclusion

The charge noise experienced by a quantum dot in a CMOS silicon nanowire device has been characterized. We have shown that such devices experience the typical $1/f$ charge noise seen in similar semiconductor devices. Through comparison with the Dutta-Horn model of charge noise, we conclude that this $1/f$ noise spectrum is generated by a non-uniform distribution of two-level fluctuators in the region of the quantum dot.

We have investigated the influence of the dot position and shape in the channel on the charge noise. It was demonstrated that operating the quantum dot with a highly positive applied back gate reduces the charge noise by one order of magnitude. This was attributed to the movement of the dot in the channel, as indicated by the changes in relative capacitances, and the expected reduced population of charge traps at the rear interface.

It was demonstrated that extension of the wavefunction towards the electron reservoirs yields a rapid increase in charge noise. We attribute this to combined effect of the increased coupling of the quantum dot to the reservoirs, which contain a high density of implanted dopants that may act as charge traps, and the lateral extension of the wavefunction bringing it into the area of the interface with the silicon nitride spacers, which are expected to also contain a high density of charge traps. We investigated the population of fluctuators experienced by the quantum dot through temperature spectroscopy. It was found that our system is in weak agreement with the Dutta-Horn model of charge noise. The temperature

dependence of S_ε is non-linear, indicating a non-uniform distribution of fluctuators.

Additionally, the population of fluctuators experienced by the quantum dot in a low-noise configuration was analysed. We found that two fluctuators or groups of fluctuators with activation energies 1.3 meV and 0.2 meV dominate the noise spectrum in this regime.

We have demonstrated a novel technique for determining the charge noise of the first electron in a charge-sensed quantum dot. This technique was used to determine that the charge noise at the single electron level is approximately on the order of the temperature over a typical measurement time of 1 s. The same technique was used to investigate the charge noise of subsequent electrons, finding that the magnitude of the noise experienced by the quantum dot decreases with the number of electrons. The charge noise was also found to increase approximately linearly with temperature at the single-electron level, varying from $(130 \pm 60) \mu\text{eV}^2/\text{Hz}$ at 400 mK to $(1310 \pm 318) \mu\text{eV}^2/\text{Hz}$ at 4 K.

Finally, the impact of the charge noise at the single electron level on spin coherence was analysed. It was found that the T_2^* is limited to 10.9 μs for a $t_{meas} = 1$ s, $T = 200$ mK at the first electron due to charge noise. This is more than two orders of magnitude longer than the nuclear spin limited T_2^* , underpinning the benefits of isotopic purification on the spin coherence time.

CHAPTER 5

Characterisation of spin physics in a single CMOS quantum dot

5.1 Introduction

Research into large-scale quantum processors in silicon is motivated by the promising aspects of qubits based on electron spins. Electron spins are an obvious candidate for the implementation of a qubit. An electron in a static magnetic field forms a natural two-level system which is accessible and controllable via electric fields [Kop06]. Electrical control of electrons down to the single electron level is routinely possible in nanowire quantum dots [Cha20; CT20; Cor18], and high-fidelity readout of electron and hole spin states has been demonstrated in these devices both in DC measurement [Mau16] and with RF reflectometry [CT20; Urd19; Wes19]. In order to implement a large-scale silicon quantum processor, we must be able to simply and robustly characterise the spin physics of electrons in nanowire quantum dots.

Here we present a method of detection and addressable single-shot measurement of a single electron spin, which could be used to characterise large arrays of silicon CMOS nanowire quantum dots. We first present the readout method used and fidelity analysis, and demonstrate the measurement of the spin relaxation time T_1 with a spin state visibility greater than 90%. The relaxation time is analysed as a function of the magnetic field, and used to detect the spin-valley relaxation hotspot at $(297 \pm 5) \mu\text{eV}$. We demonstrate control over the valley splitting via electric field tuning from $297 \mu\text{eV}$ to $260 \mu\text{eV}$, detected via measurement of the relaxation hotspot. Finally, we measure the magnetic field anisotropy of the spin-valley mixing, demonstrating suppression of the relaxation mechanism in a field oriented along the main symmetry axis of the device.

5.2 Spin detection and readout

The first step towards a spin qubit in a silicon CMOS nanowire is enabling fast, high fidelity spin detection of a single electron. Charge detection of the first electron is relatively routine, both using an adjacent SET [Yan13], and through gate or source reflectometry [Urd19]. However, spin detection is more difficult, and requires a spin-to-charge conversion mechanism. The principle of such mechanisms are detailed in section 2.3.3. Here we employ an energy-selective readout method to convert spin information to easily-measurable charge information. We attain a maximum visibility of 92% in a measurement time of 1.2 ms.

5.2.1 Background

In section 2.3.3 we discuss the various spin-to-charge conversion and spin readout mechanisms which can be employed to measure the spin of a quantum dot. For a rapid and robust characterisation of spin physics in many quantum dots, energy-selective readout is a strong choice. It is the most commonly employed method of measurement for similar systems, being favoured over Pauli spin blockade (which requires an ancilla qubit) and tunnel rate selective readout (requiring a system with a difference in tunnel rate, such as a singlet-triplet qubit) due to its comparatively light hardware requirements and simplicity.

Energy-selective readout can be used to detect the spin of a single electron quantum dot capacitively coupled to a nearby SET, and used to probe the spin relaxation and spin-valley physics of a trapped electron. It was first demonstrated in GaAs/AlGaAs quantum dots in 2004 [Elz04]. It can also be possible to use such an SET as both detector and reservoir, as in [Mor10; Pla12]. Here, an electron was loaded onto a dopant via an SET which was both electrostatically and tunnel coupled to the dopant. It also demonstrates that the method is not inhibited by using the SET as the reservoir, despite it potentially having a non-reservoir-like density of states. One advantage demonstrated in [Pla12] is that the configuration used for energy-selective readout can also be used to deterministically initialize the dopant in the spin-down state. Rapid and deterministic initialisation, such as that using this method, is one of the crucial requirements for fast operation of qubits.

Readout of the second electron is also possible using this method. In [Yan13], energy-selective readout was used to detect the spin state of a quantum dot containing two electrons. In this case, the selectivity in energy is between the singlet and T_- triplet states. The method is the same, with a T_- state being detected via a current signal through the SET. It was also used to detect the spin of the third electron, with the method being identical to that of the first and the lowest two electrons uninvolved in the readout.

This readout method requires one qubit dot to be coupled to a reservoir (or a quasi-metallic SET which can act as a reservoir), and capacitively coupled to a charge sensor [Elz03], which can take the form of a QPC [Top00], SET [Fuj04; Ibb20], or an RF-coupled quantum dot or reservoir [CT20]. The main requirements for energy-selective readout are that the Zeeman splitting $E_Z = g\mu_B B$ must be greater than the reservoir's thermal broadening, and that the tunnel rate from the qubit dot to the reservoir must be less than the measurement bandwidth [Elz04]. The Zeeman splitting is dependent on the external magnetic field to 116 $\mu\text{eV/T}$ for electrons in silicon, where $g = 2$. Similarly, the thermal energy is dependent on the temperature at 86 $\mu\text{eV/K}$. At a magnetic field of 1 T, we require a temperature much lower than 1 K to have the requisite energy difference. To maximize the readout fidelity and avoid stochastic thermal tunnelling events, the energy difference should be as large as possible. Low temperature and high magnetic field therefore are necessary for high-fidelity energy-selective readout of single spins.

The principle of energy-selective readout is demonstrated in Fig 5.1. A three-step pulse sequence is used to read out the dot. First, the dot is emptied of all electrons by lifting the energy high above the reservoir potential μ_{res} . Second, an electron of unknown spin state is loaded into the dot by plunging it deep below the reservoir potential, so that an electron with an unknown spin state is loaded. Only one electron can be loaded due to coulomb

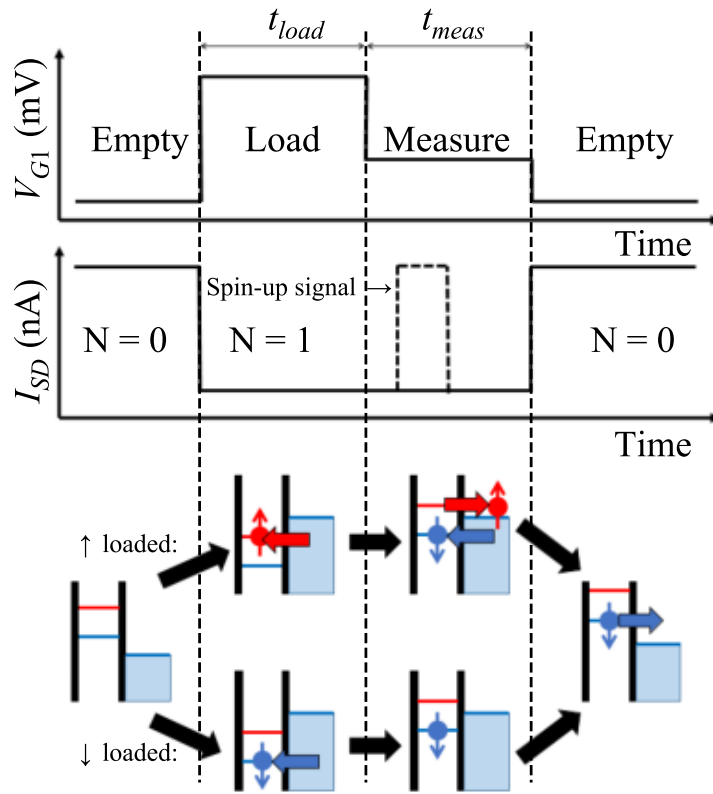


Figure 5.1: Energy selective readout of a single electron spin a) Pulse sequence of the voltage on the plunger gate controlling the qubit dot potential. b) Detector response schematic. c) Schematic of the quantum dot and reservoir during the pulse sequence. Depending on the spin state, the electron undergoes a different sequence. If the electron is spin-down (lower branch), it remains in the quantum dot, as it does not have the requisite energy to tunnel out. However, if the electron is spin-up (upper branch), it can tunnel out, leaving the quantum dot briefly empty. Another electron may then tunnel rapidly back in, restoring the original charge state of the dot. (Adapted from [Han07].)

repulsion, and it is loaded after a time based on the tunnel rate $t_{in} \approx \Gamma_{in}^{-1}$. The electron is held on the dot for a time t_{wait} . After t_{wait} , the spin state is measured by pulsing it to a position where the reservoir potential lies between the potential of the two spin states. If the electron was spin-down, it has potential $\mu_{\downarrow} < \mu_{res}$, and cannot tunnel out. However, if it was spin-up, it has $\mu_{\uparrow} > \mu_{res}$, and can tunnel out of the dot to the reservoir, leaving the quantum dot emptied of electrons. This change in the charge state of the quantum dot causes a characteristic jump in the charge detector current. A new spin-down electron can then tunnel back into the dot due to the lifting of the coulomb blockade, returning it to its original charge state.

The expected two-level readout current through the detector dot I_{SD} is depicted in Fig 5.1. The charge detector, which in this case is the adjacent quantum dot, is aligned such that an energy level lies within the bias window, allowing transport through the dot, when the qubit dot is empty of electrons. Therefore, when I_{SD} is high, the dot contains $N = 0$ electrons ($I_{N=0}$). When an electron is loaded onto the qubit dot, I_{SD} is reduced to $I_{N=1} \approx 0$ nA, as the capacitive shift causes the detector potential to shift out of resonance with the leads. It remains at $I_{N=1}$ for a time t_{wait} , as the electron is held on the quantum dot. Then, when the qubit dot is pulsed to the measurement position, the spin state of the electron dictates the current readout. If the electron was originally spin-down, the electron does not have sufficient energy to tunnel out of the dot, and remains loaded, leaving the current readout to remain $I_{N=1}$. However, if the electron was spin-up, then the electron is able to tunnel out, which it will do in a time t_{out} (which is typically short, less than 1 ms), resulting in an increase in current to $I_{N=0}$. Then, a new electron may tunnel back into the spin-down state of the empty qubit dot, resulting in the current returning to $I_{SD} = I_{N=1}$, indicating that the dot is filled. The dot can then be pulsed back above the reservoir potential μ_{res} to empty it in preparation for the next measurement, and the current returns to $I_{N=0}$.

It is important to note that the electron has a non-zero probability to relax from the excited state (spin-up) to the ground state (spin-down) whilst it is being held in the qubit dot (during t_{wait}). This relaxation is characterised by a time constant T_1 . The relaxation can be caused by phonon or electronic noise, coupled to the spin state via spin-orbit or spin-valley interaction, as discussed in section 5.3.2. The longer the electron is held, the higher the probability it will relax. Thus, in order to obtain the state of the initial electron, t_{wait} should be as short as possible. It is limited by the tunneling in time, t_{in} . However, t_{wait} can be varied in order to probe the relaxation time of the spin state by measuring the relative state probabilities, as discussed in section 5.2.4.

The method has some drawbacks. The energy requirement can be significant. The required temperatures are only achievable in a dilution refrigerator, presenting an immediate experimental complication. Similarly, the necessary fields are high, upwards of 1 T, requiring a powerful electromagnet. Additionally, the sensitive positioning of the dot potential with respect to the reservoir means significant charge noise induced switches can displace the dot from the measurement position [Jun04]. The nature of the reservoir, if it is atypical, can also be detrimental. The precise positioning of the quantum dot potential can be disrupted by a non-uniform density of states in the reservoir, preventing energy-selective readout in certain configurations. This can be especially prevalent when a quasi-metallic

quantum dot is used as a reservoir, compared to a 2DEG or metallic contact, for example. Finally, the presence of additional energy states can disrupt the energy configuration.

5.2.2 Locating spin

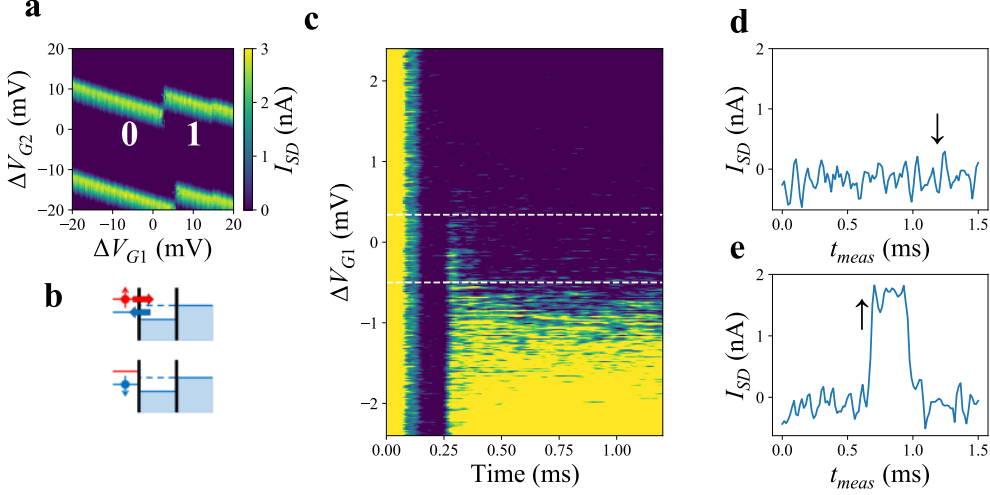


Figure 5.2: Elzerman spin readout Detail of the energy-selective readout mechanism used to measure the spin. a) Stability diagram of the transition of the first electron. The coulomb peak corresponding to a transition in the detector dot ($N_{det} \approx 40$). The prominent break in the line indicates a potential shift due to an electron entering the adjacent qubit dot. The occupancy of the qubit dot (N_e) is indicated. b) Elzerman spin readout mechanism. An electron of unknown spin is loaded from the reservoir into the quantum dot. The quantum dot is then brought to the measurement position, where the reservoir energy (dashed line) is between the energy of the spin-up state (red) and the spin-down state (blue). If the electron is in the spin-up state (red), it will rapidly tunnel out of the dot. A new spin-down electron can then tunnel into the empty dot. This process is very fast and typically takes place within a few milliseconds of pulsing to the measurement position. If the original electron is in the spin-down state (blue), it will remain in the dot, as it does not have the requisite energy to tunnel out. c) The current through the detector dot I_{SD} as a function of the measurement time, for several V_{meas} values across the transition. Around 100 traces are acquired and averaged for each V_{meas} . High current (yellow) indicates that the qubit dot is empty, while low current (blue) indicates that it is full. The white dashed lines indicate the measurement window, in which a single out-in event is seen, as evidenced by the "tail" of transient current seen most strongly at $t = 0.25$ ms. The gate voltage scale is relative to the chosen measurement position at $\Delta V_{G1} = 0$. d) The current trace obtained when the loaded electron is spin-down. No tunnelling event is seen, and the electron remains loaded onto the dot. e) The trace obtained for a spin-up electron. A single, short tunnelling event is seen, indicating an electron rapidly leaving the dot, then a new one tunnelling back in. By setting a current threshold I_{th} , we can distinguish between a spin-up and spin-down measurement. If the current trace passes the threshold within t_{meas} , we assign it spin-up. If it remains below the threshold, it is spin-down.

To read the charge on the qubit dot, we require a nearby charge sensor which is capacitively coupled to the quantum dot. In our system, we use a large quantum dot formed under the opposite gate, G_2 , as a charge sensor to probe our qubit dot, which is

controlled by G_1 . It is in the many-electron regime ($N_{\text{sensor}} \approx 45$, measured by counting transitions in the sensor dot), and has a charging energy of $E_C \approx 1$ meV. The stability diagram of the transition of the first electron is shown in Fig 5.2a. The break in the horizontal resonance peak corresponds to detection of a charge shift in the nearby qubit dot, which is strongly coupled to V_{G1} . No further breaks are seen at lower V_{G1} , so we can be sure that this is the first electron in the qubit dot.

The qubit dot is tunnel coupled to an electron reservoir, indicated by the addition of electrons. The tunnel rate into and out of the dot is measurable via a two-level trace at the degeneracy point (where stochastic in-out events are seen on the stability diagram). Such a two-level trace can be digitized, and the number of times the electron tunnels in a given time calculated to find the tunnel rate, Γ_{1e} . The tunnel rate was measured here to be in the regime of $\Gamma_{1e} \approx 100$ kHz. This was chosen to be faster than the minimum integration time for a single measurement point $t_{\text{int}} = 30$ μ s, but slow enough that a single in-out tunnelling event within a single integration point would register as a signal beyond our noise level.

The tunnel rate Γ_{1e} is highly tunable in this system by varying the number of electrons (and therefore size) of the sensor dot. For this reason, it is believed that the sensor dot is acting as the electron reservoir for the qubit dot. Electrons are therefore tunnelling from the sensor dot into the qubit dot. This can complicate the picture, as the energy states in the sensor dot now become relevant when considering the energy diagram in 5.2b. In particular, if an electron has to tunnel from the reservoir to the sensor dot, then from the sensor dot to the qubit dot, then it is possible to imagine that spin blockade could occur, whereby the available energy states in the sensor dot only allow one species of electron to be transmitted. If this were the case, then spin measurement would be prevented at certain N_{sensor} . For this dot, however, we were able to perform spin measurement for all transitions up to $N_{\text{sensor}} \pm 2$, with only the tunnel rate changing significantly. It is likely, therefore, that the sensor dot is in a configuration where many degenerate states are available, such that spin blockade does not occur.

To obtain spin readout, we have to employ a spin-to-charge conversion technique. Here, we use the method outlined in [Elz04], often referred to as "Elzerman readout" or "energy-selective readout", with the basic theory addressed in 2.3.3. A schematic of the readout pulse sequence is shown in Fig 5.2b. Before the sequence, the potential of the qubit dot is raised far above the Fermi energy (into the "empty" region), and held there for a few fractions of a millisecond to ensure that the dot is fully emptied of the last electron. The dot is then plunged briefly below the Fermi energy into the "load" region to load an electron with an unknown spin onto the quantum dot. This simulates any processes that take place before the measurement which may rotate the electron spin from the ground state into an unknown state. There is no significant difference in the tunnel rate for different spin states, so the population of spins should be around 50/50. We wait a length of time t_{wait} in the load region, before pulsing to the measurement region. The measurement region is chosen by sweeping the measurement point across the transition to determine where the chemical potential of the dot V_{meas} is well-positioned such that the Fermi energy lies between the two spin states. The pulse sequence is plotted for varying V_{meas} in Fig 5.2c, with $t_{\text{wait}} \approx 0.5$ ms.

At $t_{meas} = 0$ ms, the dot is empty. It is then initialized by pulsing to the region indicated by $N_e = 1$ in the stability diagram (a). By $t_{meas} = 0.5$ ms, the electron is loaded, as indicated by the low detector current. The qubit dot is then pulsed to a point across the transition indicated by ΔV_{G1} . When ΔV_{G1} is low (pulsing far into the "empty" region), the electron rapidly tunnels out of the dot no matter what the original spin state is, and the dot remains emptied, as indicated by the high detector current. Similarly, when ΔV_{G1} is high (pulsed far into the "load" region), both spin states lie deep below the Fermi energy, and no matter the spin state the electron remains captured, indicated by low detector current. There are two regions where this is not the case: firstly, the "thermal" region between $\Delta V_{G1} = -1 \rightarrow -0.3$ mV, where electrons can tunnel into and out of the dot due to thermally-activated tunnel events. Secondly, around $\Delta V_{G1} = 0$ mV, there is a region where an initial increase in current is seen (indicating that the dot has been emptied), then it is re-filled and remains full. This signature cannot be due to stochastic thermally-activated tunnelling such as that seen between $\Delta V_{G1} = -1 \rightarrow -0.3$ mV, as the dot remains in the $N_e = 1$ state and does not fluctuate. It is indicative of the single tunnelling event characteristic of a spin-up electron tunnelling out, and a spin-down electron tunnelling back in. This voltage range is considered the "Elzerman window", as it is the region where energy-selective tunnelling is possible. As such, V_{meas} is chosen to lie within this measurement region.

The temperature was lowered to 120 mK after spin was not detected at 400 mK. Measurements are made at high magnetic field, between 1.5 T and 3 T. The current-time trace at V_{meas} is plotted in Fig 5.2d & e. Two representative time traces are seen for two different cases. The first trace, in which no change in the current through the sensor is seen, indicates that the electron remains on the qubit dot for the duration of the measurement time. This is indicative of the loaded electron being spin-down, as this state remains below the Fermi energy and cannot normally tunnel out of the dot. The second trace has a characteristic current jump, representing an electron tunnelling out of the dot. The increase to the higher current level indicates that the qubit dot has been emptied and is in the charge state. Because there is a state below the Fermi energy, a new electron rapidly tunnels back into the qubit dot, returning it to the previous current level, indicating it is filled. This is the signature of a spin-up electron tunnelling out of the dot, and then a new spin-down electron tunnelling in. We define a current threshold I_{th} between the two current levels. If the current crosses the threshold within the integration time t_{meas} , we define the trace to be a measurement of a spin-up electron. If the current does not cross the threshold, we attribute it to the measurement of a spin-down electron. This method enables readout of the spin state of a single electron with a single shot measurement. It is a destructive method of measurement, in that the electron is lost if it is spin-up. However, it allows us to probe the spin dynamics of the qubit dot.

5.2.3 Measurement fidelity

In order to maximise the visibility of the spin state, we turn to analysis of the readout traces. Once the detection window has been found, we can analyze the fidelity to optimize our quantum dot configuration.

The fidelity of this measurement is characterised by three figures: $F_{\uparrow}$, $F_{\downarrow}$, and V . $F_{\uparrow}$

is the fidelity of a spin-up measurement. This is given by $1 - P(\downarrow | \uparrow)$, where $P(\downarrow | \uparrow)$ is the probability to measure a spin-down electron when its true state is actually spin-up. Similarly, $F_\downarrow$ is the fidelity of a spin-down measurement, given by $1 - P(\uparrow | \downarrow)$. V is a parameter called the "visibility", and is defined as $V = F_\uparrow F_\downarrow$ (in some cases, it is also defined $(F_\uparrow + F_\downarrow)/2$, or $1 - (P(\uparrow | \downarrow) + P(\downarrow | \uparrow))/2$, which are similar ways to quantify the average readout fidelity of a random spin state). To obtain the maximum measurement visibility, we aim to maximize $F_\uparrow$ and $F_\downarrow$ (equivalently, minimize $P(\downarrow | \uparrow)$ and $P(\uparrow | \downarrow)$). $P(\downarrow | \uparrow)$ is the probability to measure a spin-down electron when it was originally spin-up, and therefore corresponds to events in which the spin has relaxed. Thus, $P(\downarrow | \uparrow)$ can be minimized through faster measurement (shorter t_{wait}). $P(\uparrow | \downarrow)$ is the probability to see a spin-up measurement when the electron is spin-down. This corresponds to thermally-activated tunnelling or excitation of the spin state, and can be minimized by lowering the electron temperature.

Here, we assess the fidelity of our readout method with regards to distinguishing the two states from a current-time trace, in a typical configuration used for T_1 measurement. Thermal excitation can give false positive measurements, and these are determined through measurement of the spin-up population at long wait time. Additionally, the initialization fidelity is measured from the maximum spin-up population at short wait time. We also measure the readout fidelity at very short t_{wait} as a function of the loading position across the transition.

State visibility

Around 10,000 measurements of the spin state are taken at short t_{wait} . The spin has minimal time to relax, so we expect to obtain approximately 50% population of either spin state. Additionally, the tunnel-in time is longer than the pulse time, so the transition is not adiabatic with respect to the tunnel-in time. To analyse the fidelity of the spin, we create a cumulative histogram of the maximum current level in each trace. To allow analysis of the tunnelling out and in time for spin-down electrons, the integration time for this measurement is 1.2 ms.

A representative histogram is plotted in Fig 5.3a. The two main peaks correspond to the distributions of the two current levels I_0 and I_{max} . Each peak takes a gaussian shape and is described by equation 5.1.

$$fit = Ae^{-B(I-I_\downarrow)^2} + Ce^{-\frac{t_m}{T_1} - B(I-I_\uparrow)^2} + f_{decay}(I) \quad (5.1)$$

The pure gaussian fit to the histogram consists of two gaussian curves which describe the noise-broadened current distribution about $I_\downarrow$ and $I_\uparrow$. These take the form $Ae^{-B(I-I_\downarrow)^2}$ and $Ce^{-\frac{t_m}{T_1} - B(I-I_\uparrow)^2}$ respectively, with A, B, C as fit parameters. t_m is the integration time of a single measurement point ($t_m \approx 30 \mu s$ in this case). T_1 is the spin-lattice relaxation time, obtained from measurements detailed later in this chapter. $I_\uparrow$ and $I_\downarrow$ are the current values about which the two gaussians are centered for the spin-up and spin-down populations respectively. The integral of the overlap between these two gaussians gives $1 - V$, the error in the readout visibility. This error is the chance to obtain a false measurement, detecting spin-up when the electron is spin-down, or vice versa.

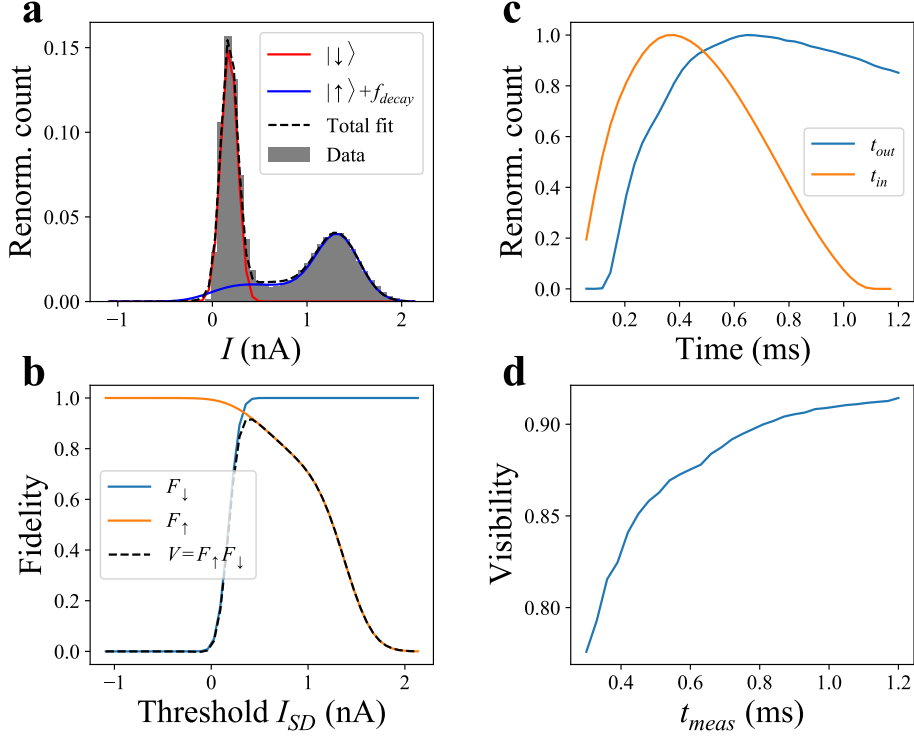


Figure 5.3: Fidelity analysis Analysis of the single state fidelity and state visibility. a) Binned maximum current data from around 10,000 traces. The blue and red traces are gaussians describing the noise-broadened spin-down and spin-up distributions respectively. Black dashed line is the total distribution with an additional $f_{decay}(I)$ to account for the loss of information during one measurement point. b) Fidelity plots for single state fidelity and state visibility. The blue and orange curves are the single state fidelities for spin-down and spin-up measurements respectively as a function of the placement of the threshold current, I_{th} . The black dashed curve is the state visibility, defined as $V = F_{\uparrow}F_{\downarrow}$. c) Histogram of tunnel-out times t_{out} (blue) and tunnel-in times t_{in} (orange). d) State visibility as a function of the measurement time t_{meas} .

In the case of two pure gaussians fit to our data, we obtain a state visibility of upwards of 99%. This represents the charge fidelity - that is, the visibility of the two charge states. However, it does not capture the spin physics, and does not fully fit the obtained histogram.

Whilst the current trace is typically a two-level trace, there is a finite probability to obtain a maximum current value between $I_{\downarrow}$ and $I_{\uparrow}$ which cannot be simply described by the noise broadening of the gaussian current distributions. This occurs when the spin-up electron tunnels out and a new electron tunnels back in during a single measurement point, leading to a lower integrated current being recorded. The tunneling times T_{out} and T_{in} are plotted in Fig 5.3c. The peak T_{out} and T_{in} are 610 μ s and 350 μ s respectively. However, it

is notable that the T_{in} curve does not start at 0 count, supporting that there is a finite (significant) probability for the electron to tunnel back in faster than the measurement interval.

$$f_{decay}(I) = \int_{I_{\downarrow}}^{I_{\uparrow}} A \frac{t_m P_{\uparrow}}{T_1(I_{\uparrow} - I_{\downarrow})} e^{-\frac{t_m}{T_1(I_{\uparrow} - I_{\downarrow})}(I_d - I_{\downarrow})^2} e^{-B(I - I_d)^2} dI_d \quad (5.2)$$

To account for this, we introduce an additional term $f_{decay}(I)$, Eqn 5.2. A, B, C are the gaussian fit parameters. $P_{\uparrow}$ is the population fraction which is measured to be spin-up (approximately 0.5). The full fit is plotted in Fig 5.3a. We can then analyse the measurement fidelity. We define the state fidelity to be $F_{\downarrow}(I_{th}) = 1 - \int_0^{I_{th}} f_{\downarrow}(I) dI$ for the spin-down, and $F_{\uparrow}(I_{th}) = 1 - \int_{I_{th}}^{I_{max}} f_{\uparrow}(I) dI$ for the spin-up electrons, where $f_{\downarrow}(I)$ is $Ae^{-B(I - I_{\downarrow})^2}$, and $f_{\uparrow}(I)$ is $Ce^{-B(I - I_{\uparrow})^2} + f_{decay}(I)$. The fidelity is plotted in Fig 5.3b as a function of I_{th} , the threshold current used to distinguish between spin-up and spin-down states. We define the visibility to be $V = F_{\uparrow}F_{\downarrow}$. It is plotted in Fig 5.3b as the dashed line. The optimum I_{th} is taken at the maximum of the visibility plot to obtain the best fidelity for both states. We obtain measurement fidelities of $F_{\downarrow} = 99\%$ and $F_{\uparrow} = 93\%$, giving a maximum visibility of $V = 92\%$. This method of obtaining the optimum I_{th} is used when processing data to ensure maximum visibility in each measurement. The state visibility is plotted as a function of the measurement time t_{meas} in Fig 5.3d. The visibility is initially low due to the noise broadening of the gaussian functions for each state, but this is reduced and the visibility improved with a longer integration time. At very long t_{meas} , the visibility will be reduced again through spin-lattice relaxation and the increased probability to observe stochastic, uncorrelated charge transitions.

In [Mor10], a similar readout fidelity of more than 90% was achieved. They obtained spin-down fidelity 99% and spin-up fidelity 93%, giving a state visibility of 92%, which is closely comparable with the fidelities obtained here. Whilst a higher fidelity would be preferable for an ideal qubit in large-scale error correction protocols due to the additional physical qubit overhead required for low fidelities, this has been proved to be more than sufficient for characterisation of spin physics.

Initialization error

When the quantum dot potential is plunged below the Fermi potential, the tunnelling probability for spin-up and spin-down electrons is expected to be 50% for each, as both energy levels lie deep below the Fermi energy. However, a non-continuous density of states in the reservoir (which is, in this case, a quasi-metallic SET), can alter this picture. The probability to load an electron of state spin-up can be measured by pulsing to the load position V_{load} for a very short t_{wait} . This allows an electron enough time to load into the dot, but does not allow time for the electron to relax. The spin-up population of repeated measurements at very short t_{wait} gives the initialization error. For the configuration discussed in the state fidelity analysis, the initialization error is $13 \pm 2\%$, from a maximum measured spin-up population of $37 \pm 2\%$.

Thermally activated tunnelling

Thermally-activated charge transitions can be detected as false spin-up signatures. These are not distinguishable in the short fidelity analysis, but by analysing the spin-up population at long t_{wait} the false readings can be identified. From T_1 measurements (detailed in the next section) we can be confident that at long t_{wait} , $P_{\uparrow}$ is less than 1%, so any current values beyond the spin-down distribution must be due to stochastic tunnelling events. Depending on the position within the measurement window (closer to the transition increases the likelihood of seeing stochastic charge transitions), we see a minimum spin-up population of $15 \pm 3\%$ at $t_{meas} \approx 1.2$ ms. This suggests that $15 \pm 3\%$ of detected spin-up electrons are due to thermal transitions, limiting our spin-up fidelity and the state visibility. It is important to note that this is highly dependent on both the measurement position and the measurement time.

Readout fidelity

Table 5.1: Readout and initialization errors

Charge visibility	99%
Spin visibility	92%
Thermally activated errors	15%
Total readout fidelity	77%
Initialization error	13%

A summary of the relevant fidelities is given in Table 5.1. The charge visibility reflects the maximum possible visibility in this configuration, and is determined by the capacitive coupling of the charge detector to the quantum dot relative to the temperature and the noise induced in the signal by the amplification chain. The spin visibility and errors due to thermal activation of tunnelling events limit the total readout fidelity to around 77%. An additional initialization error of 13% is present, likely due to the non-ideal reservoir used to load the electrons. Fast spin relaxation could also play a part in the initialization error, but the measured T_1 indicates that the spin lifetime should be long compared to the time required for the electron to tunnel out (t_{out} : Fig 5.3c), so the impact of this should be minor. This fidelity is configuration-specific, and is only a typical value, varying depending on the loading and measurement voltages, as well as the magnetic field, the occupancy of the charge sensor dot, and the integration time.

Loading spectrum

The energy spectrum of the quantum dot can be probed by analyzing the probability to load a spin-up electron as a function of the detuning across the transition. The probability to load a spin-up electron can be measured through repeated single-shot measurements to determine the population at very short t_{wait} , such that the spin does not have time to

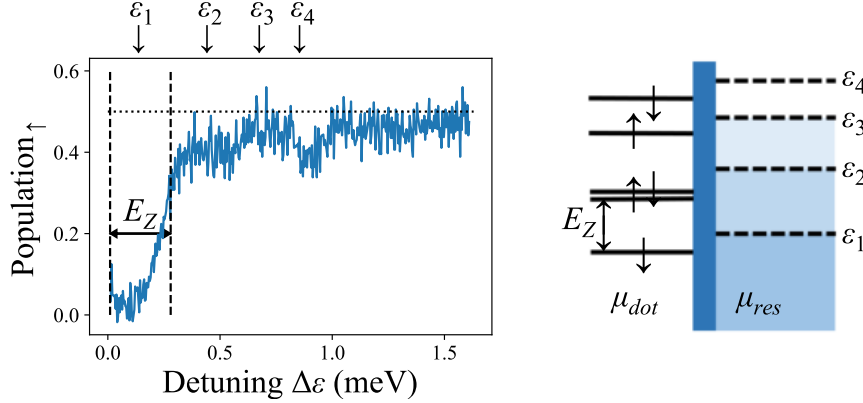


Figure 5.4: Loading Spectrum a) The spin-up population at short t_{wait} is measured as a function of the qubit gate voltage at V_{load} , converted to energy via the lever arm $\alpha = 0.27$ eV/V. The magnetic field for this measurement is 2.3 T, with the corresponding Zeeman energy $E_Z = 0.27$ meV labelled. The transition is located at $\Delta\epsilon = 0$ meV, with the measurement window indicated between the dashed lines. Values of qubit potential below the transition $\Delta\epsilon = 0$ meV are not included, as the potential is above the Fermi sea at this point and almost zero electrons are able to tunnel onto the dot. $P_{\uparrow} = 50\%$ is indicated by the black dotted line. The population has been renormalized based on the initialization and readout errors. b) A schematic of an energy level structure in the quantum dot that may yield the population map in a). The dot contains ground and excited states which have spin states separated by E_Z . In this case, the spin-up of the ground state and the spin-down of the first excited state are likely close in energy. The reservoir as positioned with respect to the quantum dot is labelled with $\epsilon_{1,2,3,4}$ corresponding to the detuning positions marked in a). At ϵ_1 , there is no spin-up state available in the quantum dot, so the population of spin-up electrons is nominally zero. At ϵ_2 , the ground spin-up state and the excited spin-down states (which are close in energy) are both below the reservoir potential. The probability to load a spin-up electron is therefore 33%, so the spin-up population increases. At ϵ_3 , the excited spin-up state is below the reservoir potential, and the probability to load a spin-up electron is maximum - nominally 50%. At ϵ_4 , a further decrease in the spin-up population below the maximum may indicate the presence of a further spin-down state entering the window, again reducing the probability to load spin-up.

relax. The spin-up population at $t_{wait} = 100$ μ s, for a measurement time $t_{meas} = 1.2$ ms, is plotted in Fig 5.4a for varying loading position V_{load} . It has been renormalized based on the readout and initialization error. As $\Delta\epsilon$ increases, the potential of the qubit dot is plunged deeper below the Fermi sea. Below $\Delta\epsilon = 0$ meV, the qubit is above the Fermi sea, and no electrons can tunnel onto the dot.

At $0 - 0.25$ meV (ϵ_1), the dot enters the measurement region, where the population of spin-up approaches zero. This is due to the only available energy level below the Fermi energy being a spin-down state, so spin-up tunnelling is prevented.

We then seen an increase to a population of around 40% from $0.25 - 0.6$ meV (ϵ_2), indicated in Fig 5.4 in red. In this region, the Zeeman energy $E_Z = 0.27$ meV ($B_{ext} = 2.3$ T has been exceeded, meaning that both spin-up and spin-down states are available. However, we are close to the excited state energy, and within a few μ eV, the spin-down state of the

next excited state becomes available; see Fig 5.4b. This means we expect to see around 33% spin-up, assuming equal probability to tunnel into any of the three available states. The spin-up population is higher than expected, possibly indicating preferential loading of spin-up due to the non-ideal reservoir.

After another 0.27 meV , we overcome the spin splitting again, and four states are available: the spin-up and spin-down states of both of the lowest valleys. We therefore see an increase towards the maximum spin-up population.

Further investigation would be interesting regarding the slight secondary minimum to around 40%, seen at $\Delta\varepsilon \approx 2.2\text{ meV}$. Two possible explanations for this could arise in our system. Firstly, it may indicate the presence of the next orbital or valley state entering the picture (bringing the expected spin-up probability to 40%). Secondly, it could result from the non-continuous density of states in the quantum dot acting as the reservoir. This would be indicative of spin-dependent tunnelling at certain qubit potentials, resulting from an inhomogeneous density of states.

Conclusion

Due to the rapid electron tunnelling back in to the dot, the fidelity for this measurement method is degraded in comparison to other similar methods (although is comparable with other similar measurements using the same method [Mor10]). Additionally, the tunnelling out of the spin-up state is fairly slow, which increases the required t_{meas} and reduces the readout fidelity due to thermal transitions. We find fairly high thermally-activated ($15 \pm 4\%$) and initialization ($13 \pm 2\%$) errors in our region of measurement. These may be a consequence of stochastic tunnelling, as well as a non-reservoir-like density of states in the quantum dot which acts as the reservoir. It is likely that the energetic picture of our quantum dot system is different to the ideal case. In order to rectify this, it may be possible to operate the sensor dot at a larger scale, i.e. with more electrons. This also significantly changes the dimensions of the sensor dot, and modifies the tunnel coupling to the qubit dot by orders of magnitude. It would therefore require improved control over the tunnel coupling, therefore, to be able to tune the qubit dot into a regime where spin is measurable. This may be possible through manipulation of the top gate and back gate. There is a wide parameter space to explore, and fully tuning a device to the ideal regime could be a long process. For characterisation of the device, however, the fidelity obtained here is sufficient; it is significantly higher than 50%, meaning we are not prevented from probing spin physics. The parameter which limits any spin readout and manipulation timescale is the spin-lattice relaxation time T_1 . In the next section we probe the T_1 of the first electron spins in the qubit dot.

5.2.4 T_1 measurement

The relaxation time of the spin T_1 can be measured through repeated spin measurements of a quantum dot after allowing some fraction of the spins to relax to their ground state. We initialize the dot with an electron of unknown spin state by pulsing deep into the load region (V_{load}) where the maximum spin-up population is obtained (see Fig 5.4). After a short time t_{wait} , we then pulse into the measurement position V_{meas} and detect the spin state. This may be repeated many times to determine the population of spin-up electrons

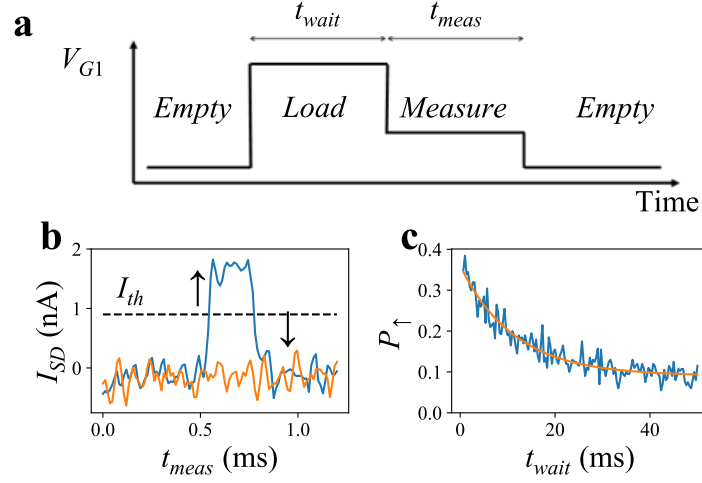


Figure 5.5: T_1 measurement sequence a) A pulse sequence schematic demonstrating the T_1 measurement method. The dot is initially emptied of all electrons by raising its potential above the Fermi energy for a long time (≈ 5 ms). It is then pulsed below the Fermi energy to load an electron. It remains in this state for a duration t_{wait} . The dot is then moved into the measurement region, where a spin-up electron can escape but a spin-down electron cannot. The maximum current recorded during t_{meas} is used to determine the parity of the spin state. The dot is then re-emptied of the remaining electron to begin the next measurement. b) Sample two-level traces from a typical T_1 measurement. The maximum current is compared to I_{th} , which is dynamically computed for each set of data. This is repeated approximately 1,000 times for each t_{wait} to extract $P_{\uparrow}$, the spin-up population. c) $P_{\uparrow}$ as a function of t_{wait} . The decay function is fitted with $P_{\uparrow} = P_{\uparrow}(0)e^{-t_{wait}/T_1}$.

measured in the quantum dot. We measure the spin-up population $P_{\uparrow}$ as a function of the waiting time t_{wait} . The spin-up state is the excited state of the system, and therefore $P_{\uparrow}$ will decay with increased wait time as the probability for a spin-up electron to relax to the ground state is increased.

A typical T_1 measurement is plotted in Fig 5.5. At $t_{wait} = 0$ ms, $P_{\uparrow}(0)$ approaches 50%. It is in reality reduced to 35 – 40% by initialization errors, consistent with the maximum population seen in Fig 5.4. $P_{\uparrow}$ decays to a minimum of $\approx 11\%$ over a few tens of ms. This minimum value of $P_{\uparrow}$ is again consistent with Fig 5.4 and results from false detection events caused by stochastic charge transitions. The decay of $P_{\uparrow}$ with t_{wait} can be described by an exponential, $P_{\uparrow}(t_{wait}) = P_{\uparrow}(0)\exp(-t_{wait}/T_1)$. The maximum T_1 measured for this device was (80 ± 5) ms at 1.5 T. Longer T_1 measurement was not possible due to thermal broadening of the transition preventing reliable spin measurement below 1.5 T.

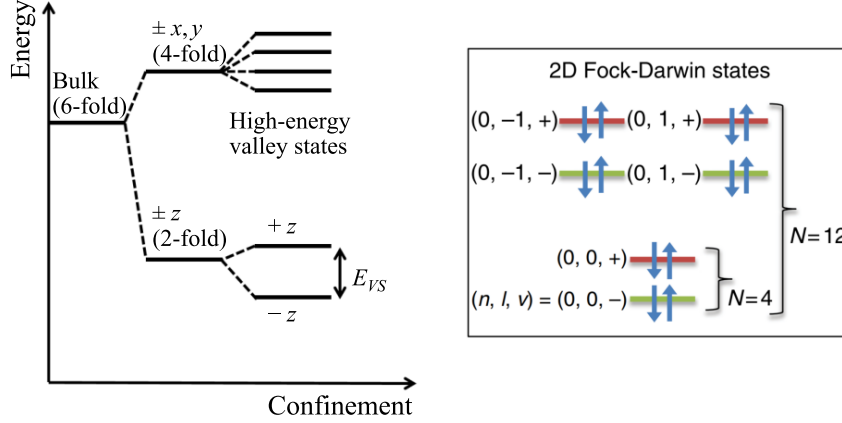


Figure 5.6: Valley states in silicon a) The lifting of valley state degeneracy through confinement. At low confinement (bulk silicon), the six valley states are degenerate in energy. However, close to an interface, the degeneracy is partly lifted, giving two low-energy states ($\pm z$) and four high-energy states ($\pm x, y$). Further in-plane confinement, such as that experienced by a 2D quantum dot, further lifts the degeneracy and separates the six valley states. The states of interest are $+z$ and $-z$, the two lowest energy states, which are separated by an energy E_{VS} . b) Spin and orbital state filling of the lowest two valleys. The shell and orbital quantum numbers n and l are indicated, with the valley states $\pm z$ labelled $+$ and $-$. In typical operation, we generally only consider the lowest four states, which are in the same orbital but separated by valley and spin energy. Adapted from [Yan13].

5.3 Spin-valley interactions

5.3.1 Valley states in silicon

The presence of energy-degenerate valley states in bulk silicon can inhibit the creation of a well-isolated two-level system [Pen20]. In bulk silicon, there are six degenerate valley states which prevent formation of a two-level spin system. These arise due to the periodicity of the silicon lattice structure. The free-space energy of an electron in k space can be written $E(\mathbf{k}) = \frac{\hbar^2 k^2}{2m}$. In the presence of degenerate valley states, this can be separated as follows:

$$E(\mathbf{k}) = E_0 + \frac{\hbar^2}{2} \left[\frac{k_x^2}{m_x} + \frac{k_y^2}{m_y} + \frac{k_z^2}{m_z} \right] \quad (5.3)$$

Where m_x , m_y and m_z are not necessarily equal. These are the effective electron mass in the associated direction, and depend on the curvature of the conduction band energy. In bulk silicon, the effective masses are the same due to the crystal isotropy, leading to the six-fold valley degeneracy ($\pm x$, $\pm y$ and $\pm z$ valleys) as depicted in Fig 5.6. However, interfaces introduce discontinuities to the conduction band, which can alter the effective electron masses and lift the valley degeneracy. Due to fabrication processes using stacking materials, such interfaces tend to be along the z axis, the growth axis of the material.

In the two main types of z -confined quantum dot, Si/SiO₂ and Si/SiGe, have a typical energy difference between the remaining ground states ($\pm z$) and the lifted x and y states of 20 meV and 200 meV respectively [Pen20]. This energy scale is much larger than that typically probed in quantum dots, especially at the single-electron level, being higher than the orbital energy E_{orb} and the charging energy. However, the energy splitting of the two remaining $\pm z$ valleys can be much less than the charging energy, and approach the Zeeman energy. This is denoted E_{VS} and it becomes highly relevant for spin physics in confined quantum dots. In general, where $E_{VS} \ll E_{orb}$, the first excited state is the next valley state. In other materials or types of quantum dot, $E_{orb} \ll E_{VS}$, and the first excited state is instead an orbital state.

We will denote the lowest energy valleys v_- and v_+ separated by the valley splitting energy E_{VS} . This splitting E_{VS} can be very variable between devices. It is determined by a combination of factors including interface disorder and atomic-scale defects, such as local strain or potential minima. Typical values of E_{VS} range between tens of μeV to a few meV. In Si/SiO₂, it is typically around 1 meV, and about 0.1 meV for Si/SiGe [Pen20]. As will be demonstrated later in this chapter, for our particular quantum dots we find that the E_{VS} is on the order of the Zeeman energy. The variability of this E_{VS} could impede the development of spin-based qubits in silicon, so it is important to understand its effects. It can then be possible to control E_{VS} , and even exploit it for valley-based spin qubits [Cul12].

Control of E_{VS} using an electrostatic gate potential was demonstrated in gate-defined silicon quantum dots [Yan13]. Electric field control over a wide range of 0.3 to 0.8 meV valley splitting was demonstrated using a local plunger gate, biased to potentials between 1.4 and 2.1 V. This range of demonstrated E_{VS} is impressive, but required use of two separate techniques for different magnetic field ranges. Nevertheless, the agreement between the two is excellent, with the plunger gate providing strong control over E_{VS} . Note that the plunger gate is also used to control the on-site potential of the quantum dot. This implies that the E_{VS} can change significantly depending on the potential required to reach the first electron. This could therefore be a good measure of subtle variability between devices.

5.3.2 Relaxation mechanisms

Electron spins are two-level systems that preferentially couple to the magnetic field. They have a magnetic moment, $\mu_B = 5.8 \mu\text{eV/T}$, which is relatively weak [Han07] meaning magnetic field fluctuations have only a limited effect on the spin. Electric field fluctuations have to be coupled to the spin through a spin-charge interaction. The main mediator in silicon quantum dots is the spin-orbit interaction. The spin orbit-interaction is the interaction between the spin state of an electron and its motion within the orbital potential of its host atom (or, equivalently, quantum dot). In bulk silicon, the spin-orbit interaction is relatively weak. However, it can be significantly enhanced via interaction with the valley states. This gives rise to a second interaction channel, the spin-valley interaction, which can also couple the electric field and the spin state.

Through these two channels, electric field fluctuations can cause the spin to relax to the ground state. The rate at which the spin relaxes, T_1^{-1} , sets an upper bound on the spin coherence time. Silicon has been shown to have long T_1 times, with T_1 values approaching

an hour reported in phosphorous donors [Tyr12], and a few seconds in electrons in quantum dots [Sim11].

$$T_1^{-1} = \Gamma_{J,SV} + \Gamma_{ph,SV} + \Gamma_{J,SO} + \Gamma_{ph,SO} + \Gamma_{const} \quad (5.4)$$

This relaxation is mediated in a strongly confined (> 1 meV) silicon quantum dot by four main relaxation mechanisms (see Eqn 5.4) [Zha20]. Each of these corresponds to a source of noise felt by the electron spin through two different coupling mechanisms. The main sources of noise are Johnson noise (J) and phonon noise (ph). $1/f$ charge noise also plays a role, but it has a relatively minor effect on the relaxation. These noise sources are coupled to the spin system via the two main channels which couple the spin with the electric field [Hua14b]: the spin-orbit coupling (SO), which arises from the motion of the electron within its orbital, and the spin-valley coupling (SV), arising from the mixing of spin and valley states.

The source of the noise is the most dominant contributor to its dependence on the magnetic field as the sources, phonon and Johnson noise, have different spectral densities. The phonon spectral density has an approximately f^5 dependence, whilst Johnson noise is approximately constant with frequency.

Spin relaxation due to phonon noise typically dominates at high magnetic field due to the higher density of phonons at higher frequency (where E_Z is large). The dependence of the relaxation due to phonon noise through the spin-valley and spin-orbit interactions ($\Gamma_{ph,SV}$ and $\Gamma_{ph,SO}$) on the applied magnetic field can be determined through correlation functions as detailed in [Hua14a]. It is found that $\Gamma_{ph,SV}$ has a B^5 dependence, whilst $\Gamma_{ph,SO}$ has a B^7 dependence. The total spin relaxation due to phonon noise is given by $\Gamma_{ph} = \Gamma_{ph,SV} + \Gamma_{ph,SO}$. We can therefore expect that at high magnetic field the spin relaxation is dominated by spin-orbit mediated phonon noise. However, at very high field, a "bottleneck" effect begins to cause the dependence to deviate from the dominant B^7 curve. This bottleneck results from an effective reduction in the coupling interaction for high frequency phonons, and comes into play as the field approaches 8 – 10 T.

Johnson noise is the other main source of spin relaxation. Contrary to phonon noise, Johnson noise has a spectral density which (in an ideal resistor) is flat across all frequencies. It is found that the Johnson noise through the spin-valley interaction, $\Gamma_{J,SV}$, is approximately linear with magnetic field when the zeeman energy is lower than the valley splitting: $E_Z \ll E_{VS}$ [Hua14a]. However, the spin-orbit-coupled Johnson noise $\Gamma_{J,SO}$ demonstrates a B^3 dependence. At very low field, the spin-valley mediated relaxation dominates the Johnson noise contribution. However once the condition $E_Z \gg E_{VS}$ is met, i.e. at high magnetic fields beyond the valley energy, the dependence of $\Gamma_{J,SV}$ becomes proportional to B^{-1} ; that is, the relaxation rate via the spin-valley coupling decreases at higher field. In this regime, the spin-orbit coupling will dominate the Johnson contribution.

5.3.3 Valley relaxation hotspot

In order to explore the different contributions to the relaxation mechanism, we measure the T_1 as a function of the static magnetic field B_Z . The results are plotted in Fig 5.7a. The fit is calculated from a combination of contributions to the relaxation mechanisms, described

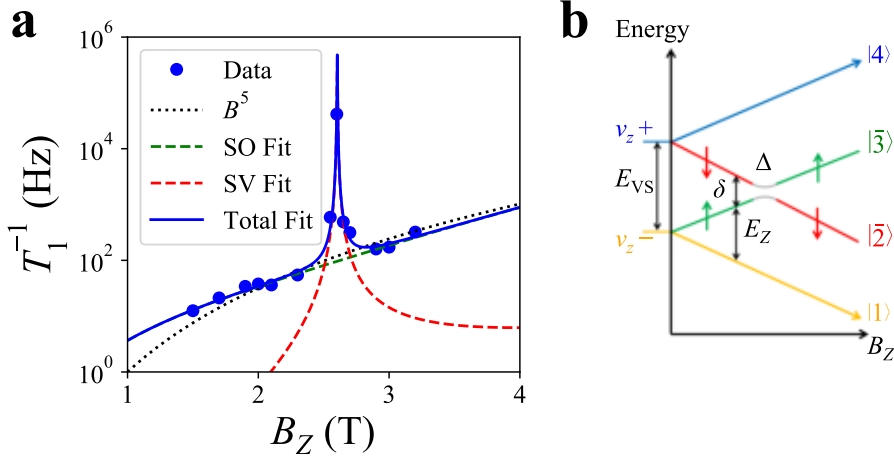


Figure 5.7: Spin-valley hotspot a) Variation of T_1 with magnetic field B_Z in the range 1.5 T to 3.2 T. The black dotted line is proportional to B^5 . Green and red dashed lines demonstrate the spin-orbit and spin-valley contributions respectively. The spin-orbit contribution $\Gamma_{SO} = \Gamma_{ph,SO} + \Gamma_{J,SO}$ follows an approximately B^5 trend. The spin-valley contribution $\Gamma_{SV} = \Gamma_{ph,SV} + \Gamma_{J,SV}$ peaks at $B_Z = E_{VS}/(g\mu_B)$, creating a point known as a spin-valley relaxation hotspot. The solid blue line gives the sum of the spin-orbit and spin-valley contributions. b) Energy diagram of the spin-valley hotspot. As B_Z approaches $E_{VS}/(g\mu_B)$, the states $|2\rangle$ and $|3\rangle$ hybridize, forming new eigenstates which are linear combinations of $|2\rangle$ and $|3\rangle$. These new eigenstates mix the spin and valley degrees of freedom and provide a rapid spin relaxation route when $E_{VS} = E_Z$. As the hotspot is approached, the energy difference between $|\uparrow, z_+\rangle$ and $|\downarrow, z_+\rangle$, $\delta = E_{VS} - E_Z$, approaches zero. Δ is the splitting at the anticrossing, also known as the valley coupling.

in Eqn 5.4. Noise is modelled as a combination of the phonon lattice deformation and Johnson (electrical) noise contributions. Each of these noise contributions normally cannot induce spin-flips. However they are coupled to the spin state via two main interaction mechanisms: spin-orbit (SO) and spin-valley (SV) interactions, as discussed in 5.3.2. The internal angular momentum of an electron is intrinsically linked to its orbital energy. Fluctuations in the electric field cause this orbital energy to shift, inducing a rotation to the electron spin.

The spin-valley interaction is depicted in Fig 5.7b. In bulk silicon, there are six degenerate valley states which an electron can occupy [Pen20]. Under strong confinement, such as that provided by an interface, the degeneracy of these valley states can be lifted. Two of these valley states, v_- and v_+ , remain relatively close in energy. v_+ and v_- correspond to the two z-axis valley states, with v_- the twofold-degenerate ground state of the system at $B_Z = 0$ T. They are separated by the valley splitting energy E_{VS} . At finite magnetic field, the spin degeneracy is lifted, and the states separate into a mixture of both spin and valley states: $|1\rangle$, $|2\rangle$, $|3\rangle$, $|4\rangle$. These are defined as $|1\rangle = |v_-, \downarrow\rangle$, $|2\rangle = |v_-, \uparrow\rangle$, $|3\rangle = |v_+, \downarrow\rangle$, and $|4\rangle = |v_+, \uparrow\rangle$. At $B_Z \ll E_{VS}/(g\mu_B)$, these are the eigenstates of the system. However,

as the Zeeman energy becomes comparable to the Valley splitting, the states $|2\rangle$ and $|3\rangle$ hybridize and form an anticrossing point. At this point it makes more sense to describe the hybrid system in terms of redefined eigenstates $|\bar{2}\rangle$ and $|\bar{3}\rangle$:

$$|\bar{2}\rangle = \sqrt{(1-a)/2}|2\rangle - \sqrt{(1+a)/2}|3\rangle \quad (5.5)$$

$$|\bar{3}\rangle = \sqrt{(1+a)/2}|2\rangle - \sqrt{(1-a)/2}|3\rangle \quad (5.6)$$

Where $a = -\delta/\sqrt{\delta^2 + \Delta^2}$ with $\delta = E_{VS} - E_Z$ the detuning from the mixing point and Δ the splitting at the anticrossing. Importantly, both of these states have a rapid relaxation mechanism to the ground state $|1\rangle$. Both have a contribution from the state $|3\rangle$, which has a spin component which is parallel to $|1\rangle$. This allows direct phonon-induced relaxation, which is normally forbidden [Yan13]. This significantly increases the spin-flip rate and reduces T_1 by several orders of magnitude where this interaction is maximal. Far from the anticrossing (i.e. where $\delta \gg \Delta$) the original eigenstates $|2\rangle$ and $|3\rangle$ are recovered, albeit with their energetic ordering swapped at high field, and the spin-valley mixing effect is reduced.

The effect of the noise through the two interaction mechanisms is plotted in Fig 5.7a. We label these as spin-orbit (SO) and spin-valley (SV) interactions for clarity. However, it is important to note that spin-valley interactions are also a spin-orbit effect with inter-valley state mixing [Zha20]. We label spin-valley as a separate interaction to distinguish it from the intra-valley spin-orbit interaction with higher orbital states, which we simply call the spin-orbit interaction. The spin-orbit contribution follows a commonly-reported approximately B^5 dependence, generally seen for dominant spin-orbit interactions in the absence of spin-valley coupling. This is due to the two noise sources, phonons and Johnson noise, having a B^7 [Tah14] and B^3 [Hua14b] dependence respectively. The B^7 dependence of the spin-orbit coupled phonon noise is characteristic of silicon quantum dots since it arises from interaction with the next orbital level [Yan13].

Conversely, the spin-valley coupling contribution is generally weaker than the spin-orbit coupling, but has a significant peak at $B_Z = 2.6$ T. The spin-valley interaction is strongly dependent on the value of E_{VS} . Its magnitude is proportional to $1 - \left(1 + \frac{\Delta^2}{(E_{VS} - E_Z)^2}\right)^{-1/2}$, which is maximal when $E_{VS} = E_Z$ [Hua14b]. Δ is the magnitude of the spin-valley coupling at the anticrossing, and 2Δ describes the width of the peak. The fitting for our data gives $\Delta \approx 200$ neV. Where $E_Z \ll E_{VS}$, the spin-valley contribution approaches a small constant Δ^2/E_{VS}^2 . When $E_Z \gg E_{VS}$ it becomes proportional to Δ^2/E_Z^2 and eventually decreases towards 0, offset by the B^7 and B^3 dependence of the phonon and Johnson noise contributions to produce the relatively low, constant contribution seen in Fig 5.7a. From this fit, we extract $E_{VS} = (297 \pm 5)$ μ eV, with the relaxation "hotspot" occurring at $B_Z \approx 2.6$ T.

A valley splitting of (681 ± 23) μ eV and (571 ± 27) μ eV was reported for two devices with a similar structure, fabricated using similar CMOS-compatible 300 mm wafer technology [CT20]. The spin-valley hotspot was not directly measured through magnetic field relaxometry, but was extracted from excited state spectroscopy by varying the loading position.

A device identical to that presented here in terms of structure, dimensions and fabrication process was also measured (data unpublished, personal communication with Cardoso-Paz), with $E_{VS} = 191 \mu\text{eV}$ found via magnetic field relaxometry. The E_{VS} measured here is therefore in line with other CMOS-fabrication quantum dots in silicon. The variation observed from device to device is significant, however. A difference of a few $100 \mu\text{eV}$ in nominally identical devices corresponds to a difference of more than 1 T in terms of the static magnetic field. For schema that seek to exploit the spin-valley interaction, precise tuning of the E_{VS} could be critical, and high variance in E_{VS} across multiple qubits could prove problematic for scalable manipulation. Fortunately, E_{VS} is not a fixed parameter, but is strongly dependent on the electric fields present in the device [Bou18; Ibb18]. In the next section we demonstrate tuning of the valley splitting by applying a universal electric field to the device.

5.3.4 Valley energy tuning

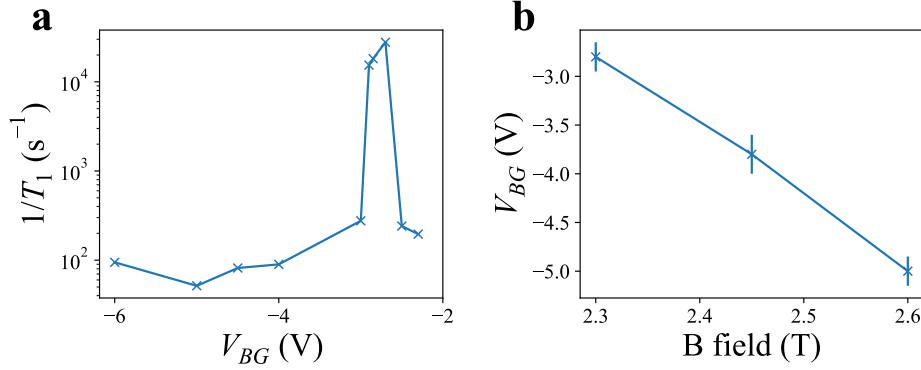


Figure 5.8: Valley energy tuning a) For a fixed B_Z , we find the hotspot by tuning the back gate voltage. The relaxation rate T_1^{-1} is plotted as a function of V_{BG} . It is approximately constant in the range -6 to -3 V. However, a prominent increase in the relaxation rate is found around $V_{BG} = 2.8$ V. This is due to the tuning of the E_{VS} with the back gate, bringing the valley splitting into resonance with the spin splitting. At this field and back gate voltage, the condition $E_{VS} = E_Z$ is fulfilled. b) The spin-valley relaxation hotspot plotted as a function of the back gate voltage V_{BG} and the magnetic field B_Z . The valley energy $E_{VS} = g\mu_B B_Z$ is tuned by over $40 \mu\text{eV}$ from 297 to $260 \mu\text{eV}$.

The valley splitting is highly dependent on the confinement of the quantum dot. A linear dependence of E_{VS} with the electric field was predicted for quantum dots in 2DEG type devices using effective mass theory [Fri07; Sar11], and demonstrated in MOS type devices over a wide range of 0.3 to 0.8 meV [Yan13]. Simulations have shown [Bou18] that such a linear dependence is predicted in CMOS nanowire devices through manipulation of the electric fields with the universal back gate.

We measure E_{VS} by finding the magnetic field value at which the T_1 becomes minimal B_{min} , giving $E_{VS} = g\mu_B B_{min}$. At a fixed magnetic field, we can measure T_1 for different V_{BG} to find the hotspot. Such a measurement is shown in Fig 5.8a for a magnetic field of 2.3 T . The hotspot is detected at $V_{BG} \approx (-2.8 \pm 0.2) \text{ V}$. This corresponds to

$E_{VS} = (260 \pm 10) \mu\text{eV}$. Thus we demonstrate that we are able to vary the value of E_{VS} by up to $40 \mu\text{eV}$ by applying a back gate potential in the range -5 V to -3 V as shown in Fig 5.8. We obtain a linear coefficient of 0.14 T/V , consistent with the trend predicted in [Bou18]. We were not able to obtain data outside of this range due to the limited range in which spin signature was visible. This is a drawback of the use of external fields to manipulate the valley splitting; applied fields, even non-local fields, necessarily interact with the dot potential in the channel. However, it is a key result that we are able to vary the valley energy up to $40 \mu\text{eV}$ even within a limited range and maintain spin detection.

Some leading schema for spin control in quantum dots make use of the spin-valley coupling to perform EDSR. Such a protocol has been demonstrated in silicon nanowire quantum dots [Cor18] using spin blockade. This allows spin control using electric fields, which is a more scalable and addressable method of control.

5.3.5 Relaxation field anisotropy

The geometry of a device and on the applied fields can have a very significant effect on the valley splitting, and on the spin-valley mixing strength Δ . It has been shown that spin-valley mixing can be highly anisotropic with the applied magnetic field in gate-defined quantum dots in silicon [Zha20]. In particular, the directionality of the spin-orbit coupling (determined by symmetries in the device) allows the spin-valley mixing to occur. In the presence of multiple orthogonal planes of symmetry, the spin-valley mixing can vanish. It has been shown that, for a device with a single major symmetry axis, the spin-orbit coupling can be defined by an effective magnetic field, $\mathbf{B}_{SO}$. When the external applied magnetic field is parallel to $\mathbf{B}_{SO}$, the spin-orbit coupling which allows spin-valley mixing can be highly suppressed. This can limit or even prevent relaxation and manipulation through valley states. In [Zha20], the point of maximum suppression of spin-valley mixing was found to have a 180-degree periodicity, with the suppression dominant along the x axis of the device. This introduces a limitation on the directionality of the external field if such mixing is desired, either for fast relaxation for initialization, or for spin-valley EDSR manipulation. However, it also provides a "sweet spot" that the field can be placed at to suppress valley mixing, in the case that it is detrimental to qubit operation.

This anisotropy was also investigated for CMOS nanowire devices [CT20]. Despite the highly symmetric device design, a single-symmetry periodicity was not observed. This was believed to be due to distortion of the quantum dot by local electric field disorder. Interface disorder can break symmetry, and also introduce new quasi-symmetric planes which complicate the geometric picture. If this is a prominent problem in nanowire devices, it could prove detrimental to spin-valley qubits and limit their reproducibility - one of their main advantages. As such, it is an important spin property to characterise.

Our system is almost identical to that outlined in [CT20], with the quantum dot strongly confined along the z axis through the positive field applied by the plunger gate and secondary fields generated by the back gate and top gate. This lifts the x and y valley degeneracy, leaving the $\pm z$ valleys as the lowest energy states in the quantum dot. Each of these valleys are twofold degenerate with spin states. These spin states are separated by the Zeeman energy E_Z , as before, along the direction of the external magnetic field B_{ext} . The spin valley states can provide a channel for spin relaxation when the spin and valley

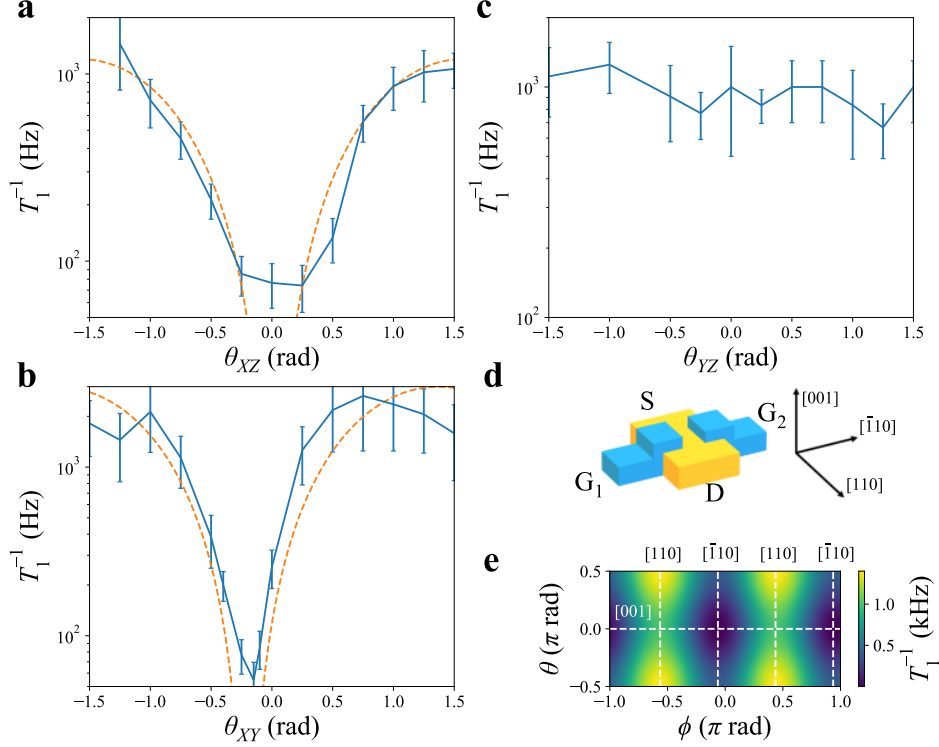


Figure 5.9: Magnetic field anisotropy Spin relaxation rate measurements to determine the spin-valley anisotropy in a quantum dot. a-c) Spin relaxation rate as a function of the angle of the magnetic field, rotated around the three main crystalline axes of the nanowire device. d) Schematic representation of the nanowire device with the crystalline axes labelled. e) A spherical heatmap of the relaxation rate, generated from $\sin^2$ fits to the relaxation curves measured in a) and b). Dashed lines corresponding to cuts of the rotation of B_{ext} about one of the crystalline axes are labelled.

states are mixed.

The mechanism that couples the spin states at the valley hotspot is an inter-valley spin mixing. This mixing is caused by spin-orbit interaction, and can be considered to be an effective magnetic field B_{SO} [Zha20]. When B_{SO} is orthogonal to B_{ext} , the applied magnetic field, this spin-valley interaction can cause the spin to flip through rotation in the plane of the B_{ext} axis, providing a relaxation channel. When the external magnetic field B_{ext} is applied parallel to B_{SO} , this relaxation channel can be suppressed, reducing the relaxation rate and producing a π rad-periodic "sweet spot" where the spin-valley relaxation components are minimized. This is generally referred to as the spin-valley relaxation anisotropy, and has been observed in Si quantum dots [Zha20], GaAs quantum dots [Hof17], and with a periodicity indicative of an unintuitive symmetry in similar CMOS nanowires

to that measured here [CT20].

The directionality of B_{SO} is determined by the symmetry of the quantum dot. Spin-orbit coupling is due to the interaction between the momentum of an electron and the electric fields in the device. Along the x axis of the device, the fields are approximately symmetric. In this direction, the momentum p_x is symmetric under reflection. Due to the device geometry, the directionality of the spin-orbit interaction is effectively fixed. We therefore probe the spin-valley relaxation anisotropy by varying the directionality of the external applied magnetic field $B_{ext}(\phi, \theta)$.

The magnetic field is rotated in three planes, labelled XY , XZ , and YZ . X , Y and Z here correspond to the three cardinal axes of the magnetic coil. The XZ and YZ planes which are swept here are rotated with respect to the coil axes by 0.2rad in order to align the X axis with the channel of the device, which was not precisely aligned with the coil axes. We can therefore associate each plane with the crystalline axes, labelling them with the axial vector aligned to the plane vector. The XY , XZ and YZ rotations therefore correspond to rotations around the crystalline axes $[001]$, $[\bar{1}10]$ and $[110]$ respectively. The axes are labelled with respect to the nanowire device in Fig 5.9d.

We measure the spin relaxation rate T_1^{-1} as a function of the angle of rotation of the magnetic field in each plane. The measured T_1^{-1} for a π rotation in each plane is plotted in Fig 5.9a-c. A prominent minimum in the relaxation rate is observed at -0.2rad in the XY plane rotation as shown in Fig 5.9b. It is fitted with $\sin^2(\theta_{XY})$, the dashed curve in Fig 5.9b. The relaxation rate minimum -0.2rad is chosen as the starting angle for the XZ rotation plane, with the YZ plane rotated by $\pi/2$ about the Z axis from XZ . The rotation along the XZ rotation plane is plotted in Fig 5.9a. A similar minimum is found at $\theta_{XZ} = 0\text{rad}$. The T_1^{-1} also follows an approximately $\sin^2(\theta_{XZ})$ behaviour. The T_1^{-1} behaviour along the YZ plane stands out as having no variation on the scale of that seen in the XY and XZ planes.

To understand this $\sin^2$ dependence, we recall the interaction of the spin and valley states. The states which are coupled here are $|2\rangle$ and $|3\rangle$ as depicted in Fig 5.7, which have spin and valley components $|2\rangle = |v_-, \uparrow\rangle$ and $|3\rangle = |v_+, \downarrow\rangle$. The energy of the two branches which interact at the anticrossing is given by [Bou18]:

$$E_{\pm} = \frac{1}{2}(E_{v_-} + E_{v_+}) \pm \frac{1}{2}\sqrt{(E_{VS} - g\mu_B B_{ext})^2 + 4|\Delta|^2} \quad (5.7)$$

The important term which defines the magnitude of the mixing between the valley states at the hotspot (when $E_{VS} \approx g\mu_B B_{ext}$) is $|\Delta|^2$. This mixing is necessary to have the rapid relaxation characteristic of the valley hotspot. When $\Delta \rightarrow 0$, the hotspot relaxation is suppressed, and the T_1 can be long. Here the spin-valley states are coupled via the spin-orbit hamiltonian given in Eqn 5.8.

$$H_{SO} = \sum_k \mathbf{p}_k \boldsymbol{\sigma}_k \quad (5.8)$$

$\boldsymbol{\sigma}_k$ are the standard Pauli matrices, and $\mathbf{p}_k$ are a set of momentum operators determining the spin-orbit interaction. Note that we are considering only the two states $|2\rangle$ and $|3\rangle$, as the states $|1\rangle$ and $|4\rangle$ are very well separated in energy. The coupling term is given by

[Bou18]:

$$\Delta = \langle 2|H_{SO}|3 \rangle \quad (5.9)$$

This defines the interaction between the strength of the spin-valley coupling and the geometry of the spin-orbit interaction. H_{SO} is necessarily a geometric hamiltonian, as it contains momentum components p_x , p_y , etc. which have a symmetry (or asymmetry) defined by the geometry of the channel. Taking the 2D case where $H_{SO} = a\sigma_x p_y + b\sigma_y p_x$, we can consider the presence of a plane of symmetry R , which is perpendicular to the x axis. It therefore only acts on the p_x component of the momentum operators. If we consider the interaction of p_x with the valley states:

$$\langle v_- | p_x | v_+ \rangle = \langle R v_- | R^\dagger p_x R | R^\dagger v_+ \rangle \quad (5.10)$$

Since v_- and v_+ are geometrically defined along the z axis, they are invariant under R ($\langle R v_- | = \langle v_- |$, $| R^\dagger v_+ \rangle = | v_+ \rangle$). The reflection operation on the momentum operator gives, by symmetry, $R^\dagger p_x R = -p_x$. Assuming there is only one plane of symmetry, this is the only axis along which this is true (there is asymmetry in p_y and p_z). Therefore the interaction of the x component of the momentum, $\langle v_- | p_x | v_+ \rangle$, must be zero, from $\langle v_- | p_x | v_+ \rangle = -\langle v_- | p_x | v_+ \rangle = 0$. So when such a mirror plane acts on H_{SO} , we obtain for Δ :

$$\Delta = \langle v_- \uparrow | H_{SO} | v_+ \downarrow \rangle = \langle v_- | p_y | v_+ \rangle \langle \uparrow | \sigma_x | \downarrow \rangle \quad (5.11)$$

When $\mathbf{B}_{ext}$ is along x , $\uparrow$ and $\downarrow$ are eigenstates of σ_x so that $\langle \uparrow | \sigma_x | \downarrow \rangle = 0$, and therefore $\Delta = 0$, giving no spin-valley mixing.

This is the basis of the symmetry argument for the spin-valley mixing "sweet spot". To consider how the magnitude of the mixing varies as a function of the magnetic field, it is simpler to think in terms of field vectors. From our definition of H_{SO} , we can define a dominant spin-orbit effective magnetic field $\mathbf{B}_{SO}$, which is orthogonal to the plane of symmetry. We use the relationship $\langle v_- | p | v_+ \rangle = (im^* E_{VS}/\hbar) \langle v_- | r | v_+ \rangle$ [Zha20] to obtain the interaction between the spin-orbit hamiltonian and the valley states:

$$\langle v_- | H_{SO} | v_+ \rangle = \langle v_- | \alpha_m p_y \sigma_x + \alpha_p p_x \sigma_y | v_+ \rangle = \frac{im^* E_{VS}}{\hbar} (\alpha' \sigma_x + \beta' \sigma_y) \quad (5.12)$$

Here m^* is the effective mass of the electron, and $\alpha' = \alpha_m r_y^{+-}$ and $\beta' = \alpha_p r_x^{-+}$ contain the inter-valley dipole matrix elements between the valley eigenstates on the y and x axes, and $\alpha_m = \beta - \alpha$ and $\alpha_p = \beta + \alpha$, constants giving the interaction strength from Dresselhaus (β) and Rashba (α) contributions. The precise definition of α' and β' is not important for the definition of $\mathbf{B}_{SO}$; a detailed explanation is given in [Zha20]. Then we can define an effective spin-orbit field $\mathbf{B}_{SO}$ such that $\langle v_- | H_{SO} | v_+ \rangle = \mathbf{B}_{SO} \cdot \boldsymbol{\sigma}$, where $\boldsymbol{\sigma}$ is the generalized pauli spin matrix:

$$\mathbf{B}_{SO} = \frac{im^* E_{VS}}{\hbar \gamma} (\alpha' \hat{\mathbf{x}} + \beta' \hat{\mathbf{y}}) \quad (5.13)$$

The symmetry argument can similarly be extended to more planes of symmetry, and it is found that in the presence of more than one perpendicular symmetry plane there can be no spin-orbit mixing with the valley states [Bou18], as there cannot be a $\mathbf{B}_{SO}$ which is orthogonal to both or all symmetry planes.

To examine how the spin-orbit mixing changes as a function of $\mathbf{B}_{ext}$, we examine how Δ changes as a function of H_{SO} . The valley-orbit interaction can be characterized by the interaction between the valley states and the spin-orbit coupling [Zha20]:

$$H_{SV} = \langle v_- | H_{SO} | v_+ \rangle \quad (5.14)$$

The magnitude of H_{SV} determines the coupling between the spin states, and therefore the rate of relaxation. From the definition of $\mathbf{B}_{SO}$:

$$\Delta = \langle \uparrow | H_{SV} | \downarrow \rangle = \langle \uparrow | \mathbf{B}_{SO} \cdot \boldsymbol{\sigma} | \downarrow \rangle \quad (5.15)$$

The states $|\uparrow\rangle$ and $|\downarrow\rangle$ are the eigenstates of $\mathbf{B}_{ext} \cdot \boldsymbol{\sigma}$. From vector transformation, we obtain the condition for $\Delta = 0$ to be $\mathbf{B}_{ext} \times \mathbf{B}_{SO} = 0$ [Zha20]. It follows that the spin-valley mixing Δ is maximum when the external field $\mathbf{B}_{ext} \perp \mathbf{B}_{SO}$, and zero when $\mathbf{B}_{ext} \parallel \mathbf{B}_{SO}$.

To determine how the magnitude of Δ varies as a function of the magnetic field angle, we can consider rotating the spin-orbit field $\mathbf{B}_{SO}$ with respect to the quantization axis. $|\uparrow\rangle$ and $|\downarrow\rangle$ are quantized along $\mathbf{B}_{ext}$. Then for $\mathbf{B}_{SO}$ at an angle θ to the quantization axis:

$$|\langle \uparrow | \mathbf{B}_{SO} \cdot \boldsymbol{\sigma} | \downarrow \rangle| \propto |\sin\theta| \quad (5.16)$$

Note that the relaxation rate T_1^{-1} is dependent on the square of the spin-valley mixing, $|\Delta|^2$. Thus we obtain:

$$|\Delta|^2 \propto \sin^2\theta. \quad (5.17)$$

Eqn 5.17 is the basis for the $\sin^2$ dependence of the relaxation rate T_1^{-1} [CT20; Cor18]. This argument can be extended to three dimensions, as the angular dependence is the same under transformation of the axes. The measurements probe the orientation of $\mathbf{B}_{SO}$ by varying the quantization axis in three dimensions. From the measurements in Fig 5.9a-c, it can be seen that the minimum mixing occurs at $\theta_{XY} = -0.2$ rad, $\theta_{XZ} = 0$ rad. We can conclude therefore that $\mathbf{B}_{SO}$ is along this axis, which is approximately parallel to the [110] crystalline axis. This coincides with the axis of maximal symmetry in the nanowire device, along which we expect to find $\mathbf{B}_{SO}$.

We fit our data to $\sin^2$ along the XY and XZ axes, seen in Fig 5.9a and b. From these fits we can construct a 2D map of the T_1^{-1} , shown in Fig 5.9e. The white dashed lines correspond to rotations of B_{ext} about the labelled axis, with θ and ϕ the corresponding angles in spherical coordinates. The minimum of the T_1^{-1} is due to the suppression of the spin-valley mixing ($\Delta \rightarrow 0$). Notably, the spin relaxation rate is not reduced to 0. There are two main factors which can determine the minimum of the anisotropy curve.

First is that interface disorder can distort the electron wavefunction such that there is no longer a single symmetry plane. Indeed, the relaxation sweet spot is predicted to occur even with a small degree of symmetry breaking [CT20]. In this case, the $\mathbf{B}_{SO}$ can

be tilted away from the channel axis. A small degree of symmetry breaking could cause the minimum valley mixing to be finite, such that a degree of the spin-valley relaxation contribution remains.

Secondly, it is important to note that the "pure" spin-orbit interaction is still present, and this can also have a degree of anisotropy with the magnetic field direction through a similar argument to that made here for the mixing mechanism. It is highly likely that the "pure" spin-orbit T_1 contribution will also vary with the angle of $\mathbf{B}_{ext}$ (that which does not involve the valleys). This can also lead to an variation in the effective "pure" spin-orbit contribution along the symmetry axis. However, given the magnitude of the "pure" spin-orbit contribution, this alone cannot explain the T_1^{-1} behaviour seen here, and so we conclude that the anisotropic behaviour is dominated by the effects of the spin-valley mixing.

In conclusion, we have measured the anisotropy of the spin-valley mixing. The relaxation rate can be tuned over more than an order of magnitude through suppressing the spin-valley mixing mechanism. The axis of minimal mixing was found to align with the axis of maximum symmetry in the device. It is in agreement with anisotropy measurements in other types of quantum dots [Hof17; Zha20] and tight-binding simulations of a device of similar design and fabrication [Bou18].

However, such an orientation dependence was not seen in another geometrically similar device [CT20]. The dependence observed there did not correspond to the $\sin^2\theta$ dependence expected with a single symmetry axis. It is believed that this is due to the presence of multiple non-orthogonal axes of symmetry, such that the mixing sweet spot is less strongly geometrically defined. However, given the similar geometric symmetry in the device structure, it is curious that these symmetry axes do not arise in our data.

One explanation for this could be that our dot is comparatively less disordered. Interface disorder can perturb the shape of the dot wavefunction and significantly change the potential environment the electron experiences. The two devices were fabricated using the same method, but fabrication issues meant that some devices were missing an important element of the gate stack. This has been seen on devices demonstrating unusual characteristics. Thanks to parametric testing, we are sure that the device measured in this thesis lies in the group which contains this element, offering a potential explanation for the discrepancy across similar devices. It would imply that in the "normal" case, the disorder experienced by the quantum dot in this type of device is relatively low, with a strong mirror plane observed; a promising conclusion for future development.

5.3.6 T_1 charge noise

When operated close to the valley splitting hotspot, the quantum dot is in a sensitive configuration. Fluctuations in the potential of the quantum dot over a long duration can change the effective field it experiences, leading to variation in E_{VS} . This can be probed by measuring the variation in the spin relaxation rate over time. The variation in the relaxation time T_1 over more than three hours is plotted in Fig 5.10a. The low-frequency fluctuations, already apparent in the time trace, are characteristic of $1/f$ noise. Each point represents an average of 200 traces over the course of 4 min. The magnetic field is aligned such that the spin relaxation is dominated by the spin-valley interaction; that is, on the

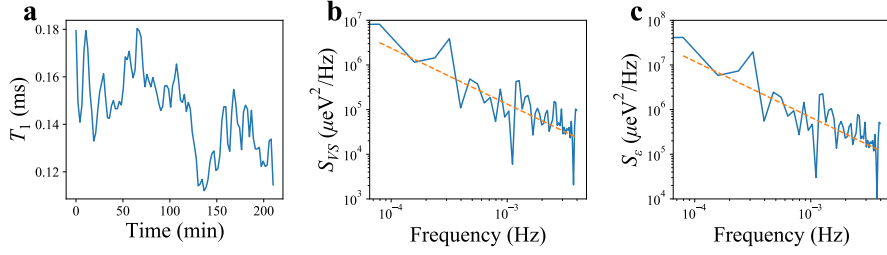


Figure 5.10: Charge noise measurement via spin relaxation The charge noise experienced by the first electron is investigated by measuring the change in E_{VS} with time via the spin relaxation rate. a) A time trace of the T_1 over several hours. Low-frequency oscillations are visible, induced via low-frequency noise. Each measurement point represents a T_1 measurement of duration 4 min. b) Renormalized power spectral density of the T_1 low-frequency noise. The time trace in a) is converted to the frequency domain via fourier transform and renormalized using the method detailed in the text. The orange dashed line is a fit proportional to $1/f^{1.25}$.

edge of the relaxation hotspot peak. To obtain the variation in valley splitting energy, we take the fourier transform of the time trace to obtain S_{T_1} and use the gradient at the side of the valley hotspot peak:

$$S_{VS} = \frac{S_{T_1}}{|dT_1/dB_Z|^2} \quad (5.18)$$

For this to hold, we assume the gradient dT_1/dB_Z calculated from the hotspot fit in Fig 5.7 is approximately constant for small fluctuations. The variation in the valley splitting with frequency is plotted in Fig 5.10b. A clear $1/f^\beta$ trend is seen, with a frequency exponent of $\beta = 1.25$. Extrapolating to 1 Hz yields a variation in the valley splitting energy of around $23 \mu\text{eV}^2/\text{Hz}$, which corresponds to a variation of a few μeV over a measurement time of one second. The decoherence of the qubit due to this fluctuating effective magnetic field can be estimated by considering the frequency of Larmor precession in this field. In an external field of magnitude B which is not aligned with the spin, an electron will precess at a frequency given by:

$$f = \gamma B \quad (5.19)$$

Where γ is the gyromagnetic ratio for an electron 28.0 GHz/T . Fluctuations in the T_1 of $5 \mu\text{eV}$ correspond to an effective field of 43 mT . The frequency of the fluctuations in the precession of the spin is therefore 1.2 GHz . This is much faster than the natural T_2^* due to the hyperfine field in silicon, which is about 20 MHz , and is on the order of the decoherence rate of a charge qubit (1 GHz , [Kim15]). It is highly dependent on the configuration of the quantum dot due to the $|dT_1/dB_Z|$ term, and suggests that the quantum dot undergoes fast decoherence as it approaches the mixing point. This is an important result when considering the use of spin-valley mixing for spin manipulation.

Spin-valley mixing can be used to implement EDSR spin manipulation [Cor18]. EDSR uses fluctuating electrical fields to manipulate an electron spin. A mechanism is required to

couple the spin and electric fields. Typically this is done via the spin-orbit interaction. In silicon, the spin-orbit coupling for electrons is weak. This is sometimes overcome by using a micromagnet to generate a magnetic field gradient, allowing a pseudo-SOC interaction as the electron is moved within the gradient. However, the spin-valley interaction allows even a weak spin-orbit coupling to induce fast spin rotation. It was used in [Cor18] to drive selective spin rotation in a blockaded quantum dot, allowing addressable manipulation. The spin states can be coupled via the valley states when the zeeman energy E_Z is close to the valley splitting E_{VS} . At this point, the spin-up state of the lower valley is mixed with the spin-down state of the upper valley, strongly enhancing the Rabi frequency near the anti-crossing. The Rabi frequency here can reach tens to hundreds of MHz [Bou18].

However, the magnitude of the T_1 fluctuations indicates that close to the mixing point, the spin coherence time could be very short, around 1 ns. This could prove detrimental for EDSR, as the manipulation time would have to be much shorter than the coherence time. If this method of spin manipulation is to be used, it may be necessary to find a sweet spot where the spin-valley mixing is sufficiently strong, but the coherence time remains long enough for spin manipulation.

To investigate the origin of this noise, we can compare it to the first electron charge noise by converting it to a potential fluctuation induced by a global gate. We can use previous analyses to compare the valley splitting variation to the potential fluctuations induced by charge noise, with some assumptions. First is that the shift in the valley splitting is caused by deformation of the potential of the quantum dot in a manner similar to charge noise. Second is that the potential shift affects the valley splitting E_{VS} in a similar way to the electric field generated by the voltage applied to the back gate V_{BG} . This allows us to use both dE_{VS}/dV_{BG} from the hotspot tuning measurement and the lever arm of the back gate to relate low-frequency fluctuations in the relaxation time T_1 to the potential shift induced by charge noise. To relate the variation in the valley splitting to the potential fluctuations at the dot, we create the following renormalization equation using these assumptions:

$$S_\varepsilon = \frac{\alpha_{BG}^2 S_{VS}}{|dE_{VS}/dV_{BG}|^2} \quad (5.20)$$

The power spectral density of the T_1 is calculated via fourier transform. In order to renormalize the PSD such that the potential fluctuation of the dot is being considered, we use Eqn 5.20. The back gate lever arm is measured to be $\alpha_{BG} = 0.0378$ meV/V. The change in energy with the back gate is $dE_{VS}/dV_{BG} = 16.8$ μ eV/V.

The renormalized power spectral density is plotted in Fig 5.10c. We extrapolate the $1/f$ behaviour to high frequency and find a value of 121 μ eV²/Hz. This is in good agreement with the 90 μ eV²/Hz measured using tunnel rate fluctuations.

Note that this may be an over-estimation of the magnitude of the potential shift, as fluctuations to the spin relaxation time due to variations in the phonon or Johnson noise may also contribute. As discussed in section 5.3.2, phonon and Johnson noise have a much greater direct effect on the spin relaxation than the $1/f$ charge noise. However, these noise sources are expected to dominate at higher frequencies, and as such should be averaged out in this measurement, so the dominant contribution to the low-frequency fluctuation

should be largely from variations in the electric field controlling E_{VS} .

The magnitude of the potential shift at the qubit dot may also be underestimated when only considering the T_1 variation, as only charge fluctuations which affect E_{VS} will cause T_1 fluctuations. It is likely that some charge fluctuations induce potential shifts which do not affect E_{VS} but still contribute to charge noise (for example, potential shifts along the main symmetry axis).

Because of these factors - and the assumptions made when converting S_{VS} to S_ϵ - the precise value of S_ϵ should be considered critically. The good agreement with the tunnel rate fluctuation measurement suggests it is a reasonable estimate, and that the T_1 fluctuations close to the spin-valley mixing hotspot are likely dominated by $1/f$ charge noise. The magnitude of the fluctuations in the valley splitting energy E_{VS} due to charge noise are large, and induce an effective field of 43 mT, corresponding to a coherence time of 1.2 ns. This suggests that attempts to further reduce the charge noise experienced by the quantum dot would be beneficial for spin qubits operated in a spin-valley mode.

5.4 Conclusion

The behaviour of a single electron spin in a silicon CMOS nanowire quantum dot has been measured and characterised. Spin readout using the Elzerman energy-selective readout method has been demonstrated with a single-state fidelity of up to 99% for spin-down and 93% for spin-up electrons, and a state visibility of 92%. Thermally-induced tunnelling errors reduce the readout fidelity to $< 80\%$. Even with these limitations, the readout method remains robust enough to characterise the spin physics of the quantum dot.

T_1 values between 80 ms and 100 μ s have been measured in various device configurations, to a minimum magnetic field of 1.5 T. The evolution of T_1 with the external applied magnetic field has been measured and fit to a spin-orbit spin-valley relaxation model, with a valley energy of 297 μ eV and valley coupling of 200 neV extracted. It was shown that the T_1 follows an approximately B^5 dependence with magnetic field, consistent with phonon and Johnson noise interacting through spin-orbit coupling far from the spin-valley relaxation hotspot. It was demonstrated that the valley energy E_{VS} can be tuned over 40 μ eV through manipulation of the electric field using a universal back gate in the range -6 V to -3 V.

The spin-valley anisotropy was measured, with a relaxation sweet spot located along the $[110]$ crystalline axis, parallel to the channel of the device. At the sweet spot, the spin relaxation rate is suppressed by more than an order of magnitude due to a strong symmetry plane breaking the spin-orbit interaction. This anisotropy is consistent with results in similar qubit implementations, indicating that local disorder is low in CMOS nanowire devices.

Finally, the effect of the charge noise on the spin relaxation time was investigated, finding a charge noise value of 121 μ eV²/Hz when converted to a potential fluctuation. This is in good agreement with the single-electron charge noise measured in section 4.4, suggesting that low frequency T_1 fluctuations are charge noise-induced. They corresponds to fluctuations of the valley splitting energy on the order of 5 μ eV over the course of 1 s. The effective field induced by these fluctuations is 43 mT, leading to a T_2^* of 1.2 GHz. This is

very fast, and suggests that a balance must be found between the increased Rabi frequency at the mixing point and keeping the coherence time long enough for manipulation.

Conclusion

The goal of this thesis was to characterize the charge noise and spin physics of quantum dots in a silicon CMOS-fabricated nanodevice. Despite the rapid advancements towards CMOS qubits in recent years, the quantum dots formed in these devices are not well-understood. In-depth characterization of these quantum dots is necessary to enable development of reproducible devices with low variability in terms of local disorder, gate controls, and dot geometries. Here, we aim to understand the properties of CMOS quantum dots that could inhibit qubit control.

The $1/f$ charge noise experienced by a quantum dot in a CMOS silicon nanowire device was characterized, and we conclude that this $1/f$ noise spectrum is generated by a non-uniform distribution of two-level fluctuators in the region of the quantum dot. The influence of the dot position and shape in the channel on the charge noise was investigated, with a noise reduction of over an order of magnitude possible by tuning the position of the quantum dot. It was found via temperature spectroscopy that two fluctuators or groups of fluctuators with activation energies 1.3 meV and 0.2 meV dominate the noise spectrum in the low-noise regime. Additionally, we demonstrated a simple technique for determining the charge noise of the first electron in a charge-sensed quantum dot, finding that the charge noise at the single electron level is approximately on the order of the temperature over a typical measurement time of 1 s. It was also used to investigate the charge noise of subsequent electrons, finding that the magnitude of the noise experienced by the quantum dot decreases with the number of electrons. Finally, the impact of the charge noise at the single electron level on spin coherence was analysed. It was found that the T_2^* is limited to 10.9 μ s for a $t_{meas} = 1$ s, $T = 200$ mK at the first electron due to charge noise. This is more than two orders of magnitude longer than the nuclear spin limited T_2^* , underpinning the benefits of isotopic purification on the spin coherence time.

The behaviour of a single electron spin in a silicon CMOS nanowire quantum dot was also characterized, using the Elzerman energy-selective readout method to detect single electron spins with a visibility of 92%. T_1 values between 80 ms and 100 μ s were measured in various device configurations, to a minimum magnetic field of 1.5 T. The evolution of T_1 with the external applied magnetic field has been measured and fit to a spin-orbit spin-valley relaxation model, with a valley energy of 297 μ eV and valley coupling of 200 neV extracted. It was shown that the T_1 follows an approximately B^5 dependence with magnetic field, consistent with phonon and Johnson noise interacting through spin-orbit coupling far from the spin-valley relaxation hotspot. The anisotropy in the spin-valley mixing was measured,

with a relaxation sweet spot located along the $[110]$ crystalline axis, parallel to the channel of the device. At the sweet spot, the spin relaxation rate is suppressed by more than an order of magnitude due to a strong symmetry plane breaking the spin-orbit interaction. This anisotropy is consistent with results in similar qubit implementations, indicating that local disorder is low in CMOS nanowire devices. Finally, the effect of the charge noise on the spin relaxation time was investigated, finding a charge noise value of $121 \text{ } \mu\text{eV}^2/\text{Hz}$ when converted to a potential fluctuation. This corresponds to fluctuations of the valley splitting energy on the order of $5 \text{ } \mu\text{eV}$ over the course of 1 s . The effective field induced by these fluctuations is 43 mT , leading to a T_2^* of 1.2 GHz . This is very fast, and suggests that a balance must be found between the increased Rabi frequency at the mixing point and keeping the coherence time long enough for manipulation.

These results represent an important step towards large scale characterization of CMOS qubits. The characterizations presented here are relatively simple to implement, with low hardware overhead and fully DC control, and can be used for routine analysis of CMOS quantum dots. We find that the charge noise compares favourably with similar types of quantum dot implementations in the many-electron regime. The charge noise at the first electron is likely to be a more useful measure for CMOS qubits, and it was found that fluctuations over one second are consistently on the order of the temperature, which is promising for future development. Additionally, characterization of the spin-valley mixing anisotropy suggests low disorder in the environment of the quantum dot. The measured T_1 is similar to that seen in other silicon devices, and is expected to approach seconds at 1 T and below.

An important result of the charge noise measurement is the indication that the coherence time is not charge noise limited in the general case. The isotopic purification of silicon is a complex and expensive process, so it is vital to be sure that spin qubits can benefit from it and are not limited by other factors. We did however find that the coherence time is strongly reduced close to the valley mixing hotspot due to fluctuations in the valley splitting energy which are charge noise-like. This has significant implications for spin-valley mediated EDSR manipulation, whereby the system is pulsed into the mixing regime where the decoherence is enhanced for the duration of the manipulation. If the manipulation is slower than the decoherence, coherent qubit manipulation becomes difficult.

Further characterization is needed to have a more complete understanding of these devices. By characterizing a device containing more quantum dots, it would be possible to measure the effect of charge noise on the tunnel barriers between dots, and distinguish this from the effect of fluctuations at quantum dot sites. With more quantum dots available, it could be possible to isolate a dot from the reservoirs and measure the charge noise in isolation via charge sensing, to determine if the reservoirs have a significant effect. The effect of charge noise on spin manipulation should be studied, especially in the spin-valley mixing regime, where fluctuations in the spin splitting can be significant. The $1/f$ nature of these fluctuations suggests that charge noise is a likely candidate for their origin, but further investigation is necessary to confirm this. Whilst EDSR has been demonstrated using the spin-valley mixing [Cor18], the spin coherence time has not been measured and enhanced decoherence could be detrimental to electrical control of coupled spins. The characterization of coupled spins, and the effect of charge noise on two-qubit gates (via

fluctuations in the coupling, for example), will be necessary further down the line, when such operations have been demonstrated. Finally, efforts are underway in collaboration with our industrial partner to study correlations between room-temperature and low-temperature behaviour. Mass automated characterization at 2 K is now possible with a 300 mm wafer cryo-prober, and many of the techniques presented here are possible at this temperature and within the limitations of the cryo-prober, opening the door to obtaining statistics on mass batch measurement of many devices from cycle to cycle.

CMOS qubits have come a long way since their first development. Though they lag behind superconducting and semiconductor heterostructure quantum dots, their rapid advancement in recent years is exciting for the field. As inter-device variability is reduced and the design perfected, it is likely that further milestones will be reached in the coming months and years. The push for progress has enormous weight behind it, with a powerhouse of the semiconductor industry devoting significant resources towards a CMOS qubit, and the longstanding expertise of the solid state research community to guide development. Whether or not CMOS qubits are the eventual qubit of choice for a large scale quantum computer, the progress that has been made in this field will contribute significantly towards that end, as well as having implications for the wider semiconductor industry.

Bibliography

- [Ahn20] Seongjin Ahn, S Das Sarma, and JP Kestner: ‘Microscopic bath effects on noise spectra in semiconductor quantum dot qubits’, *arXiv preprint arXiv:2007.03689* (2020) (see p. 72).
- [Aru19] Frank Arute, Kunal Arya, Ryan Babbush, Dave Bacon, Joseph C Bardin, Rami Barends, Rupak Biswas, Sergio Boixo, Fernando GSL Brandao, David A Buell, et al.: ‘Quantum supremacy using a programmable superconducting processor’, *Nature* **574** (2019), 505–510 (see pp. 7, 8).
- [Bar14] Rami Barends, Julian Kelly, Anthony Megrant, Andrzej Veitia, Daniel Sank, Evan Jeffrey, Ted C White, Josh Mutus, Austin G Fowler, Brooks Campbell, et al.: ‘Superconducting quantum circuits at the surface code threshold for fault tolerance’, *Nature* **508** (2014), 500–503 (see p. 7).
- [Ben08] Jan Benhelm, Gerhard Kirchmair, Christian F Roos, and Rainer Blatt: ‘Towards fault-tolerant quantum computing with trapped ions’, *Nature Physics* **4** (2008), 463–466 (see p. 9).
- [Ben84] Charles H Bennett and Gilles Brassard: ‘Quantum cryptography: Public key distribution and coin tossing’, *Proceedings of IEEE International Conference on Computers, Systems and Signal Processing* **175** (1984), 8 (see p. 2).
- [Ber14] Adam Bermeister, Daniel Keith, and Dimitrie Culcer: ‘Charge noise, spin-orbit coupling, and dephasing of single-spin qubits’, *Applied Physics Letters* **105** (2014), 192102 (see p. 58).
- [Ber15] Benoit Bertrand: ‘Long-range transfer of spin information using individual electrons’, Thèse de doctorat dirigée par Meunier, Tristan Physique de la matière condensée et du rayonnement Grenoble Alpes 2015, PhD thesis, 2015 (see p. 33).
- [Blo46] Felix Bloch: ‘Nuclear induction’, *Physical review* **70** (1946), 460 (see p. 3).
- [Bou17] Dimitri Boudier, Bogdan Cretu, Eddy Simoen, Anabela Veloso, and Nadine Collaert: ‘On quantum effects and low frequency noise spectroscopy in Si Gate-All-Around Nanowire MOSFETs at cryogenic temperatures’, *2017 Joint International EUROSIOI Workshop and International Conference on Ultimate Integration on Silicon (EUROSIOI-ULIS)*, IEEE, 2017, 5–8 (see p. 76).

- [Bou18] Léo Bourdet and Yann-Michel Niquet: ‘All-electrical manipulation of silicon spin qubits with tunable spin-valley mixing’, *Physical Review B* **97** (2018), 155433 (see pp. 106, 107, 109–112, 114).
- [Bro11] Kenton R Brown, Andrew C Wilson, Yves Colombe, C Ospelkaus, Adam M Meier, Emanuel Knill, Dietrich Leibfried, and David J Wineland: ‘Single-qubit-gate error below 10^{-4} in a trapped ion’, *Physical Review A* **84** (2011), 030303 (see p. 9).
- [Bru09] Daniel Brunner, Brian D Gerardot, Paul A Dalgarno, Gunter Wüst, Khaled Karrai, Nick G Stoltz, Pierre M Petroff, and Richard J Warburton: ‘A coherent single-hole spin in a semiconductor’, *science* **325** (2009), 70–72 (see p. 56).
- [Bru19] Colin D Bruzewicz, John Chiaverini, Robert McConnell, and Jeremy M Sage: ‘Trapped-ion quantum computing: Progress and challenges’, *Applied Physics Reviews* **6** (2019), 021314 (see p. 9).
- [Cal74] MA Caloyannides: ‘Microcycle spectral estimates of $1/f$ noise in semiconductors’, *Journal of Applied Physics* **45** (1974), 307–316 (see p. 59).
- [Cam17] Earl T Campbell, Barbara M Terhal, and Christophe Vuillot: ‘Roads towards fault-tolerant universal quantum computation’, *Nature* **549** (2017), 172–179 (see p. 7).
- [Cha20] Emmanuel Chanrion, David J Niegemann, Benoit Bertrand, Cameron Spence, Baptiste Jadot, Jing Li, Pierre-André Mortemousque, Louis Hutin, Romain Maurand, Xavier Jehl, et al.: ‘Charge detection in an array of CMOS quantum dots’, *Physical Review Applied* **14** (2020), 024066 (see p. 87).
- [Chi13] Lilian Childress and Ronald Hanson: ‘Diamond NV centers for quantum computing and quantum networks’, *MRS bulletin* **38** (2013), 134–138 (see p. 9).
- [Cir95] Juan I Cirac and Peter Zoller: ‘Quantum computations with cold trapped ions’, *Physical review letters* **74** (1995), 4091 (see p. 9).
- [CT20] Virginia N Ciriano-Tejel, Michael A Fogarty, Simon Schaal, Louis Hutin, Benoit Bertrand, M Fernando Gonzalez-Zalba, Jing Li, Y-M Niquet, Maud Vinet, and John JL Morton: ‘Spin readout of a CMOS quantum dot by gate reflectometry and spin-dependent tunnelling’, *arXiv preprint arXiv:2005.07764* (2020) (see pp. 87, 88, 105, 107, 109, 111, 112).
- [Con19] Elliot J Connors, JJ Nelson, Haifeng Qiao, Lisa F Edge, and John M Nichol: ‘Low-frequency charge noise in Si/SiGe quantum dots’, *Physical Review B* **100** (2019), 165305 (see pp. 56, 62, 71, 72, 76).
- [Cor18] Andrea Corna, Léo Bourdet, Romain Maurand, Alessandro Crippa, Dharmraj Kotekar-Patil, Heorhii Bohuslavskyi, Romain Laviéville, Louis Hutin, Sylvain Barraud, Xavier Jehl, et al.: ‘Electrically driven electron spin resonance mediated by spin–valley–orbit coupling in a silicon quantum dot’, *npj quantum information* **4** (2018), 1–7 (see pp. 13, 14, 87, 107, 111, 113, 114, 118).

- [Cre13] Marco Crescentini, Marco Bennati, Marco Carminati, and Marco Tartagni: ‘Noise limits of CMOS current interfaces for biosensors: A review’, *IEEE transactions on biomedical circuits and systems* **8** (2013), 278–292 (see p. 56).
- [Cro97] SM Cronenwett, SR Patel, CM Marcus, K Campman, and AC Gossard: ‘Mesoscopic fluctuations of elastic cotunneling in Coulomb blockaded quantum dots’, *Physical review letters* **79** (1997), 2312 (see p. 10).
- [Cul09] Dimitrie Culcer, Xuedong Hu, and S Das Sarma: ‘Dephasing of Si spin qubits due to charge noise’, *Applied Physics Letters* **95** (2009), 073102 (see p. 56).
- [Cul12] Dimitrie Culcer, AL Saraiva, Belita Koiller, Xuedong Hu, and S Das Sarma: ‘Valley-based noise-resistant quantum computation using Si quantum dots’, *Physical review letters* **108** (2012), 126804 (see p. 102).
- [Deh14] Juan P Dehollain, Juha T Muhonen, Kuan Y Tan, Andre Saraiva, David N Jamieson, Andrew S Dzurak, and Andrea Morello: ‘Single-shot readout and relaxation of singlet and triplet states in exchange-coupled p 31 electron spins in silicon’, *Physical review letters* **112** (2014), 236801 (see p. 81).
- [Den16] Vasil S Denchev, Sergio Boixo, Sergei V Isakov, Nan Ding, Ryan Babbush, Vadim Smelyanskiy, John Martinis, and Hartmut Neven: ‘What is the computational value of finite-range tunneling?’, *Physical Review X* **6** (2016), 031015 (see p. 10).
- [DiV00] David P DiVincenzo: ‘The physical implementation of quantum computation’, *Fortschritte der Physik: Progress of Physics* **48** (2000), 771–783 (see p. 4).
- [Edw91] Arthur H Edwards: ‘Interaction of H and H 2 with the silicon dangling orbital at the < 111> Si/SiO 2 interface’, *Physical Review B* **44** (1991), 1832 (see p. 57).
- [Eke91] Artur K Ekert: ‘Quantum cryptography based on Bell’s theorem’, *Physical review letters* **67** (1991), 661 (see p. 2).
- [Elz03] JM Elzerman, R Hanson, JS Greidanus, LH Willems Van Beveren, S De Franceschi, LMK Vandersypen, S Tarucha, and LP Kouwenhoven: ‘Few-electron quantum dot circuit with integrated charge read out’, *Physical Review B* **67** (2003), 161308 (see p. 88).
- [Elz04] JM Elzerman, R Hanson, LH Willems Van Beveren, B Witkamp, LMK Vandersypen, and Leo P Kouwenhoven: ‘Single-shot read-out of an individual electron spin in a quantum dot’, *nature* **430** (2004), 431–435 (see pp. 88, 92).
- [Fik95] W Fikry, G Ghibaudo, H Haddara, S Cristoloveanu, and M Dutoit: ‘Method for extracting deep submicrometre MOSFET parameters’, *Electronics Letters* **31** (1995), 762–764 (see p. 44).
- [Fla18] Fulvio Flamini, Nicolo Spagnolo, and Fabio Sciarrino: ‘Photonic quantum information processing: a review’, *Reports on Progress in Physics* **82** (2018), 016001 (see p. 8).

- [Fog18] MA Fogarty, KW Chan, B Hensen, W Huang, T Tanttu, CH Yang, A Laucht, M Veldhorst, FE Hudson, KM Itoh, et al.: ‘Integrated silicon qubit platform with single-spin addressability, exchange control and single-shot singlet-triplet readout’, *Nature communications* **9** (2018), 1–8 (see p. 81).
- [Fri07] Mark Friesen, Sucismita Chutia, Charles Tahan, and SN Coppersmith: ‘Valley splitting theory of Si Ge/ Si/ Si Ge quantum wells’, *Physical Review B* **75** (2007), 115318 (see p. 106).
- [Fuj04] T Fujisawa, T Hayashi, Y Hirayama, HD Cheong, and YH Jeong: ‘Electron counting of single-electron tunneling current’, *Applied physics letters* **84** (2004), 2343–2345 (see p. 88).
- [Fuj98] Toshimasa Fujisawa, Tjerk H Oosterkamp, Wilfred G Van der Wiel, Benno W Broer, Ramón Aguado, Seigo Tarucha, and Leo P Kouwenhoven: ‘Spontaneous emission spectrum in double quantum dot devices’, *Science* **282** (1998), 932–935 (see p. 144).
- [Gam96] D Gammon, ES Snow, BV Shanabrook, DS Katzer, and D Park: ‘Homogeneous linewidths in the optical spectrum of a single gallium arsenide quantum dot’, *Science* **273** (1996), 87–90 (see p. 10).
- [Ger98] Neil Gershenfeld and Isaac L Chuang: ‘Quantum computing with molecules’, *Scientific American* **278** (1998), 66–71 (see p. 9).
- [Ghi88] Gérard Ghibaudo: ‘New method for the extraction of MOSFET parameters’, *Electronics Letters* **24** (1988), 543–545 (see p. 44).
- [Gil95] David L Gilden, Thomas Thornton, and Mark W Mallon: ‘1 f noise in human cognition’, *Science* **267** (1995), 1837–1839 (see p. 55).
- [Han07] Ronald Hanson, Leo P Kouwenhoven, Jason R Petta, Seigo Tarucha, and Lieven MK Vandersypen: ‘Spins in few-electron quantum dots’, *Reviews of modern physics* **79** (2007), 1217 (see pp. 30, 89, 102).
- [HC19] Patrick Harvey-Collard, N Tobias Jacobson, Chloé Bureau-Oxton, Ryan M Jock, Vanita Srinivasa, Andrew M Mounce, Daniel R Ward, John M Anderson, Ronald P Manginell, Joel R Wendt, et al.: ‘Spin-orbit interactions for singlet-triplet qubits in silicon’, *Physical review letters* **122** (2019), 217702 (see p. 82).
- [Hau14] M Hauck, F Seilmeier, SE Beavan, A Badolato, PM Petroff, and A Högele: ‘Locating environmental charge impurities with confluent laser spectroscopy of multiple quantum dots’, *Physical Review B* **90** (2014), 235306 (see p. 56).
- [Hei03] Thomas Heinzel and Igor Zozoulenko: *Mesoscopic electronics in solid state nanostructures*, vol. 3, Wiley Online Library, 2003 (see p. 23).
- [Hof17] Andrea Hofmann, Ville F Maisi, Tobias Krähenmann, Christian Reichl, Werner Wegscheider, Klaus Ensslin, and Thomas Ihn: ‘Anisotropy and suppression of spin-orbit interaction in a gaas double quantum dot’, *Physical review letters* **119** (2017), 176807 (see pp. 108, 112).

- [Hoo94] Friits N Hooge: ‘1 f noise sources’, *IEEE Transactions on electron devices* **41** (1994), 1926–1935 (see p. 56).
- [Hu06] Xuedong Hu and S Das Sarma: ‘Charge-fluctuation-induced dephasing of exchange-coupled spin qubits’, *Physical review letters* **96** (2006), 100501 (see p. 60).
- [Hua20] He-Liang Huang, Dachao Wu, Daojin Fan, and Xiaobo Zhu: ‘Superconducting quantum computing: a review’, *Science China Information Sciences* **63** (2020), 1–32 (see p. 7).
- [Hua14a] Peihao Huang and Xuedong Hu: ‘Electron spin relaxation due to charge noise’, *Physical Review B* **89** (2014), 195302 (see pp. 60, 103).
- [Hua14b] Peihao Huang and Xuedong Hu: ‘Spin relaxation in a Si quantum dot due to spin-valley mixing’, *Physical Review B* **90** (2014), 235315 (see pp. 103, 105).
- [Hut16] L Hutin, R Maurand, D Kotekar-Patil, A Corna, H Bohuslavskyi, X Jehl, S Barraud, S De Franceschi, M Sanquer, and M Vinet: ‘Si CMOS platform for quantum information processing’, *2016 IEEE Symposium on VLSI Technology*, IEEE, 2016, 1–2 (see p. 10).
- [Ibb18] David J Ibberson, Léo Bourdet, José C Abadillo-Uriel, Imtiaz Ahmed, Sylvain Barraud, María J Calderón, Y-M Niquet, and M Fernando Gonzalez-Zalba: ‘Electric-field tuning of the valley splitting in silicon corner dots’, *Applied Physics Letters* **113** (2018), 053104 (see p. 106).
- [Ibb20] David J Ibberson, Theodor Lundberg, James A Haigh, Louis Hutin, Benoit Bertrand, Sylvain Barraud, Chang-Min Lee, Nadia A Stelmashenko, Jason WA Robinson, Maud Vinet, et al.: ‘Large dispersive interaction between a CMOS double quantum dot and microwave photons’, *arXiv preprint arXiv:2004.00334* (2020) (see p. 88).
- [Ito14] Kohei M Itoh and Hideyuki Watanabe: ‘Isotope engineering of silicon and diamond for quantum computing and sensing applications’, *MRS communications* **4** (2014), 143–157 (see p. 83).
- [Jad21] Baptiste Jadot, Pierre-André Mortemousque, Emmanuel Chanrion, Vivien Thiney, Arne Ludwig, Andreas D Wieck, Matias Urdampilleta, Christopher Bäuerle, and Tristan Meunier: ‘Distant spin entanglement via fast and coherent electron shuttling’, *Nature Nanotechnology* (2021), 1–6 (see p. 10).
- [Jak98] C Jakobson, I Bloom, and Y Nemirowsky: ‘1 f noise in CMOS transistors for analog applications from subthreshold to saturation’, *Solid-State Electronics* **42** (1998), 1807–1817 (see p. 56).
- [Jes17] Jan Jeske, Desmond WM Lau, Xavier Vidal, Liam P McGuinness, Philipp Reineck, Brett C Johnson, Marcus W Doherty, Jeffrey C McCallum, Shinobu Onoda, Fedor Jelezko, et al.: ‘Stimulated emission from nitrogen-vacancy centres in diamond’, *Nature communications* **8** (2017), 1–8 (see p. 10).

- [Jon01] Jonathan A Jones: ‘Quantum computing and nuclear magnetic resonance’, *PhysChemComm* **4** (2001), 49–56 (see p. 9).
- [Jun06] Hak Kee Jung and Sima Dimitrijevic: ‘Analysis of subthreshold carrier transport for ultimate DGMOSFET’, *IEEE transactions on electron devices* **53** (2006), 685–691 (see p. 44).
- [Jun04] SW Jung, T Fujisawa, Y Hirayama, and YH Jeong: ‘Background charge fluctuation in a GaAs quantum dot device’, *Applied Physics Letters* **85** (2004), 768–770 (see pp. 33, 56, 90).
- [Kan98] Bruce E Kane: ‘A silicon-based nuclear spin quantum computer’, *nature* **393** (1998), 133–137 (see p. 10).
- [Kim15] Dohun Kim, DR Ward, CB Simmons, John King Gamble, Robin Blume-Kohout, Erik Nielsen, DE Savage, MG Lagally, Mark Friesen, SN Coppersmith, et al.: ‘Microwave-driven coherent operation of a semiconductor quantum dot charge qubit’, *Nature nanotechnology* **10** (2015), 243–247 (see p. 113).
- [Kim19] J-S Kim, Thomas M Hazard, Andrew A Houck, and Stephen A Lyon: ‘A low-disorder metal-oxide-silicon double quantum dot’, *Applied Physics Letters* **114** (2019), 043501 (see p. 62).
- [Kon14] Mikhail Konnik and James Welsh: ‘High-level numerical simulations of noise in CCD and CMOS photosensors: review and tutorial’, *arXiv preprint arXiv:1412.4031* (2014) (see p. 56).
- [Kop06] Frank HL Koppens, Christo Buizert, Klaas-Jan Tielrooij, Ivo T Vink, Katja C Nowack, Tristan Meunier, LP Kouwenhoven, and LMK Vandersypen: ‘Driven coherent oscillations of a single electron spin in a quantum dot’, *Nature* **442** (2006), 766–771 (see p. 87).
- [Kou01] Leo P Kouwenhoven, DG Austing, and Seigo Tarucha: ‘Few-electron quantum dots’, *Reports on Progress in Physics* **64** (2001), 701 (see pp. 19, 22).
- [Kum90] Arvind Kumar, Steven E Laux, and Frank Stern: ‘Electron states in a GaAs quantum dot in a magnetic field’, *Physical Review B* **42** (1990), 5166 (see p. 10).
- [Lan14] Trevor Lanting, Anthony J Przybysz, A Yu Smirnov, Federico M Spedalieri, Mohammad H Amin, Andrew J Berkley, Richard Harris, Fabio Altomare, Sergio Boixo, Paul Bunyk, et al.: ‘Entanglement in a quantum annealing processor’, *Physical Review X* **4** (2014), 021041 (see p. 10).
- [Lee06] Jeong-Hyun Lee, Sang-Yun Kim, Il-Hyun Cho, Sung-Bo Hwang, and Jong-Ho Lee: ‘1 f noise characteristics of sub-100 nm MOS transistors’, *JSTS: Journal of Semiconductor Technology and Science* **6** (2006), 38–42 (see p. 56).
- [Len01] Patrick M Lenahan, TD Mishima, J Jumper, TN Fogarty, and RT Wilkins: ‘Direct experimental evidence for atomic scale structural changes involved in the interface-trap transformation process’, *IEEE Transactions on Nuclear Science* **48** (2001), 2131–2135 (see p. 57).

- [Li94] X Li and LKJ Vandamme: ‘Normalised 1 f noise: a more sensitive diagnostic tool for hot-carrier degradation in submicron MOSFET’s’, *conference; Proc. 5th European Symposium on Reliability of Electron Devices, Failure Physics and Analysis, Glasgow, UK, October 1994*, 1994, 131–134 (see p. 56).
- [Liu09] T-JK Liu and L Chang: ‘Transistor scaling to the limit’, *Into the Nano Era*, 2009, 191–223 (see p. 1).
- [Liv96] C Livermore, CH Crouch, RM Westervelt, KL Campman, and AC Gossard: ‘The Coulomb blockade in coupled quantum dots’, *Science* **274** (1996), 1332–1335 (see p. 10).
- [Los98] Daniel Loss and David P DiVincenzo: ‘Quantum computation with quantum dots’, *Physical Review A* **57** (1998), 120 (see p. 10).
- [Lun02] Kent H Lundberg: ‘Noise sources in bulk CMOS’, *Unpublished paper* **3** (2002), 28 (see p. 56).
- [Mac11] Chris A Mack: ‘Fifty years of Moore’s law’, *IEEE Transactions on semiconductor manufacturing* **24** (2011), 202–207 (see p. 1).
- [Mat86] WH Matthaeus and ML Goldstein: ‘Low-frequency 1 f noise in the interplanetary magnetic field’, *Physical review letters* **57** (1986), 495 (see p. 55).
- [Mau16] R Maurand, X Jehl, D Kotekar-Patil, A Corna, H Bohuslavskiy, R Laviéville, L Hutin, S Barraud, M Vinet, M Sanquer, et al.: ‘A CMOS silicon spin qubit’, *Nature communications* **7** (2016), 1–6 (see pp. 13, 14, 87).
- [McW57] AL McWhorter and RH Kingston: ‘Semiconductor surface physics’, *University of Pennsylvania Press, Philadelphia, PA* (1957), 207–228 (see pp. 56, 61).
- [Mon95] Chris Monroe, David M Meekhof, Barry E King, Wayne M Itano, and David J Wineland: ‘Demonstration of a fundamental quantum logic gate’, *Physical review letters* **75** (1995), 4714 (see p. 9).
- [Mon13] Christopher Monroe and Jungsang Kim: ‘Scaling the ion trap quantum processor’, *Science* **339** (2013), 1164–1169 (see p. 8).
- [Mon16] Thomas Monz, Daniel Nigg, Esteban A Martinez, Matthias F Brandl, Philipp Schindler, Richard Rines, Shannon X Wang, Isaac L Chuang, and Rainer Blatt: ‘Realization of a scalable Shor algorithm’, *Science* **351** (2016), 1068–1070 (see p. 9).
- [Moo21] Gary J Mooney, Gregory AL White, Charles D Hill, and Lloyd CL Hollenberg: ‘Whole-device entanglement in a 65-qubit superconducting quantum computer’, *arXiv preprint arXiv:2102.11521* (2021) (see p. 7).
- [Mor10] Andrea Morello, Jarryd J Pla, Floris A Zwanenburg, Kok W Chan, Kuan Y Tan, Hans Huebl, Mikko Möttönen, Christopher D Nugroho, Changyi Yang, Jessica A Van Donkelaar, et al.: ‘Single-shot readout of an electron spin in silicon’, *Nature* **467** (2010), 687–691 (see pp. 11, 88, 96, 99).

- [Mor18] Pierre-Andre Mortemousque, Emmanuel Chanrion, Baptiste Jadot, Hanno Flentje, Arne Ludwig, Andreas D Wieck, Matias Urdampilleta, Christopher Bauerle, and Tristan Meunier: ‘Coherent control of individual electron spins in a two dimensional array of quantum dots’, *arXiv preprint arXiv:1808.06180* (2018) (see p. 10).
- [Mye08] AH Myerson, DJ Szwed, SC Webster, DTC Allcock, MJ Curtis, G Imreh, JA Sherman, DN Stacey, AM Steane, and DM Lucas: ‘High-fidelity readout of trapped-ion qubits’, *Physical Review Letters* **100** (2008), 200502 (see p. 9).
- [Nak99] Yasunobu Nakamura, Yu A Pashkin, and Jaw Shen Tsai: ‘Coherent control of macroscopic quantum states in a single-Cooper-pair box’, *nature* **398** (1999), 786–788 (see p. 7).
- [Nat14] S Natarajan, M Agostinelli, S Akbar, M Bost, A Bowonder, V Chikarmane, S Chouksey, A Dasgupta, K Fischer, Q Fu, et al.: ‘A 14nm logic technology featuring 2 nd-generation finfet, air-gapped interconnects, self-aligned double patterning and a 0.0588 μm 2 sram cell size’, *2014 IEEE International Electron Devices Meeting, IEEE*, 2014, 3–7 (see p. 1).
- [Nie01] Michael A Nielsen and Isaac L Chuang: ‘Quantum computation and quantum information’, *Phys. Today* **54** (2001), 60 (see p. 6).
- [OC02] Adelmo Ortiz-Conde, FJ Garcia Sánchez, Juin J Liou, Antonio Cerdeira, Magali Estrada, and Y Yue: ‘A review of recent MOSFET threshold voltage extraction methods’, *Microelectronics reliability* **42** (2002), 583–596 (see p. 43).
- [Pal14] E Paladino, YM Galperin, G Falci, and BL Altshuler: ‘1/f noise: Implications for solid-state quantum information’, *Reviews of Modern Physics* **86** (2014), 361 (see pp. 58, 61).
- [Pen20] Nicholas Penthorn: ‘Investigating Valley States and their Interactions in Silicon/Silicon-Germanium Quantum Dots’, PhD thesis, UCLA, 2020 (see pp. 101, 102, 104).
- [Pet99] MG Peters, JI Dijkhuis, and LW Molenkamp: ‘Random telegraph signals and 1/f noise in a silicon quantum dot’, *Journal of applied physics* **86** (1999), 1523–1526 (see p. 56).
- [Pet18] L Petit, JM Boter, HGJ Eenink, G Droulers, MLV Tagliaferri, R Li, DP Franke, KJ Singh, JS Clarke, RN Schouten, et al.: ‘Spin lifetime and charge noise in hot silicon quantum dot qubits’, *Physical review letters* **121** (2018), 076801 (see pp. 62, 72).
- [Pet05] Jason R Petta, Alexander Comstock Johnson, Jacob M Taylor, Edward A Laird, Amir Yacoby, Mikhail D Lukin, Charles M Marcus, Micah P Hanson, and Arthur C Gossard: ‘Coherent manipulation of coupled electron spins in semiconductor quantum dots’, *Science* **309** (2005), 2180–2184 (see p. 56).
- [Pla12] Jarryd J Pla, Kuan Y Tan, Juan P Dehollain, Wee H Lim, John JL Morton, David N Jamieson, Andrew S Dzurak, and Andrea Morello: ‘A single-atom electron spin qubit in silicon’, *Nature* **489** (2012), 541–545 (see pp. 10, 59, 83, 84, 88).

- [Rio99] Michael Riordan, Lillian Hoddeson, and Conyers Herring: ‘The invention of the transistor’, *More Things in Heaven and Earth* (1999), 563–578 (see p. 1).
- [Sar11] AL Saraiva, MJ Calderón, Rodrigo B Capaz, Xuedong Hu, S Das Sarma, and Belita Koiller: ‘Intervalley coupling for interface-bound electrons in silicon: an effective mass study’, *Physical Review B* **84** (2011), 155320 (see p. 106).
- [Sim11] CB Simmons, JR Prance, BJ Van Bael, Teck Seng Koh, Zhan Shi, DE Savage, MG Lagally, Robert Joynt, Mark Friesen, SN Coppersmith, et al.: ‘Tunable spin loading and T 1 of a silicon spin qubit measured by single-shot readout’, *Physical review letters* **106** (2011), 156804 (see p. 103).
- [Slu19] Sergei Slussarenko and Geoff J Pryde: ‘Photonic quantum information processing: A concise review’, *Applied Physics Reviews* **6** (2019), 041303 (see p. 7).
- [Ste99] Andrew M Steane: ‘Efficient fault-tolerant quantum computing’, *Nature* **399** (1999), 124–126 (see p. 4).
- [Str20] Tom Struck, Arne Hollmann, Floyd Schauer, Olexiy Fedorets, Andreas Schmidbauer, Kentarou Sawano, Helge Riemann, Nikolay V Abrosimov, Łukasz Cywiński, Dominique Bougeard, et al.: ‘Low-frequency spin qubit energy splitting noise in highly purified 28 Si/SiGe’, *npj Quantum Information* **6** (2020), 1–7 (see pp. 59, 84).
- [Swi98] M Switkes, AG Huibers, CM Marcus, K Campman, and AC Gossard: ‘High bias transport and magnetometer design in open quantum dots’, *Applied physics letters* **72** (1998), 471–473 (see p. 10).
- [Tah14] Charles Tahan and Robert Joynt: ‘Relaxation of excited spin, orbital, and valley qubit states in ideal silicon quantum dots’, *Physical Review B* **89** (2014), 075302 (see p. 105).
- [Tan19] Tuomo Tanttu, Bas Hensen, Kok Wai Chan, Chih Hwan Yang, Wister Wei Huang, Michael Fogarty, Fay Hudson, Kohei Itoh, Dimitrie Culcer, Arne Laucht, et al.: ‘Controlling spin-orbit interactions in silicon quantum dots using magnetic field direction’, *Physical Review X* **9** (2019), 021028 (see p. 58).
- [Tar96] Seigo Tarucha, DG Austing, T Honda, RJ Van der Hage, and Leo P Kouwenhoven: ‘Shell filling and spin effects in a few electron quantum dot’, *Physical Review Letters* **77** (1996), 3613 (see p. 22).
- [Top00] MA Topinka, Brian J LeRoy, SEJ Shaw, EJ Heller, RM Westervelt, KD Maranowski, and AC Gossard: ‘Imaging coherent electron flow from a quantum point contact’, *Science* **289** (2000), 2323–2326 (see p. 88).
- [Tow] *Towards Data Science*, <https://towardsdatascience.com/can-we-teach-a-computer-quantum-mechanics-part-i-c3e724e31e1a>, Accessed: 2021-04-22 (see p. 3).

- [Tsu99] Morikazu Tsuno, Masato Suga, Masayasu Tanaka, Kentaro Shibahara, Mitiko Miura-Mattausch, and Masataka Hirose: ‘Physically-based threshold voltage determination for MOSFET’s of all gate lengths’, *IEEE transactions on electron devices* **46** (1999), 1429–1434 (see p. 43).
- [Tuo17] Jani Tuorila, Matti Partanen, Tapio Ala-Nissila, and Mikko Möttönen: ‘Efficient protocol for qubit initialization with a tunable environment’, *npj Quantum Information* **3** (2017), 1–12 (see p. 4).
- [Tyr12] Alexei M Tyryshkin, Shinichi Tojo, John JL Morton, Helge Riemann, Nikolai V Abrosimov, Peter Becker, Hans-Joachim Pohl, Thomas Schenkel, Michael LW Thewalt, Kohei M Itoh, et al.: ‘Electron spin coherence exceeding seconds in high-purity silicon’, *Nature materials* **11** (2012), 143–147 (see pp. 10, 103).
- [Urd19] Matias Urdampilleta, David J Niegemann, Emmanuel Chanrion, Baptiste Jadot, Cameron Spence, Pierre-André Mortemousque, Christopher Bäuerle, Louis Hutin, Benoit Bertrand, Sylvain Barraud, et al.: ‘Gate-based high fidelity spin readout in a CMOS device’, *Nature nanotechnology* **14** (2019), 737–741 (see pp. 15, 87).
- [Van94] Lode KJ Vandamme, Xiaosong Li, and Dominique Rigaud: ‘1/f noise in MOS devices, mobility or number fluctuations?’, *IEEE Transactions on Electron Devices* **41** (1994), 1936–1945 (see p. 56).
- [Van01] Lieven MK Vandersypen, Matthias Steffen, Gregory Breyta, Costantino S Yannoni, Mark H Sherwood, and Isaac L Chuang: ‘Experimental realization of Shor’s quantum factoring algorithm using nuclear magnetic resonance’, *Nature* **414** (2001), 883–887 (see p. 9).
- [Vel17] M Veldhorst, HGJ Eenink, Chih-Hwan Yang, and Andrew S Dzurak: ‘Silicon CMOS architecture for a spin-based quantum computer’, *Nature communications* **8** (2017), 1–8 (see p. 10).
- [Vel15] Menno Veldhorst, CH Yang, JCC Hwang, W Huang, JP Dehollain, JT Muhonen, S Simmons, A Laucht, FE Hudson, Kohei M Itoh, et al.: ‘A two-qubit logic gate in silicon’, *Nature* **526** (2015), 410–414 (see pp. 12, 13).
- [Voi13] Sorin Voinigescu: *High-frequency integrated circuits*, Cambridge University Press, 2013 (see p. 1).
- [Wes19] Anderson West, Bas Hensen, Alexis Jouan, Tuomo Tanttu, Chih-Hwan Yang, Alessandro Rossi, M Fernando Gonzalez-Zalba, Fay Hudson, Andrea Morello, David J Reilly, et al.: ‘Gate-based single-shot readout of spins in silicon’, *Nature nanotechnology* **14** (2019), 437–441 (see p. 87).
- [Wes89] Bruce J West and Michael F Shlesinger: ‘On the ubiquity of 1 f noise’, *International Journal of Modern Physics B* **3** (1989), 795–819 (see p. 55).
- [Wie02] Wilfred G Van der Wiel, Silvano De Franceschi, Jeroen M Elzerman, Toshimasa Fujisawa, Seigo Tarucha, and Leo P Kouwenhoven: ‘Electron transport through double quantum dots’, *Reviews of Modern Physics* **75** (2002), 1 (see p. 137).

- [Yan13] CH Yang, A Rossi, R Ruskov, NS Lai, FA Mohiyaddin, S Lee, C Tahan, Gerhard Klimeck, A Morello, and AS Dzurak: ‘Spin-valley lifetimes in a silicon quantum dot with tunable valley splitting’, *Nature communications* **4** (2013), 1–8 (see pp. 60, 87, 88, 101, 102, 105, 106).
- [Yon17] J Yoneda, K Takeda, T Otsuka, T Nakajima, MR Delbecq, G Allison, T Honda, T Kodera, S Oda, Y Hoshi, et al.: ‘A > 99.9%-fidelity quantum-dot spin qubit with coherence limited by charge noise’, *arXiv preprint arXiv:1708.01454* (2017) (see pp. 10–12, 55, 60, 78).
- [Zha20] Xin Zhang, Rui-Zi Hu, Hai-Ou Li, Fang-Ming Jing, Yuan Zhou, Rong-Long Ma, Ming Ni, Gang Luo, Gang Cao, Gui-Lei Wang, et al.: ‘Giant anisotropy of spin relaxation and spin-valley mixing in a silicon quantum dot’, *Physical Review Letters* **124** (2020), 257701 (see pp. 103, 105, 107, 108, 110–112).
- [Zho99] X Zhou, KY Lim, and D Lim: ‘A simple and unambiguous definition of threshold voltage and its implications in deep-submicron MOS device modeling’, *IEEE Transactions on Electron Devices* **46** (1999), 807–809 (see p. 44).

List of Figures

1.1	Bloch Sphere	3
1.2	Quantum algorithm	5
1.3	Qubit realizations	8
1.4	High-fidelity readout	11
1.5	Single- and two-qubit gates	12
1.6	CMOS qubit manipulation	14
1.7	Single-shot readout	15
2.1	Quantum Dot model	18
2.2	Coulomb blockade model	21
2.3	Coulomb peaks	22
2.4	Coulomb diamonds	23
2.5	Charge detector	26
2.6	Double dot spin states	30
2.7	Energy selective readout	32
2.8	Tunnel rate readout	33
2.9	Spin blockade readout	34
3.1	Experiment schematic	39
3.2	Device design	41
3.3	Room temperature characterization	43
3.4	Low-temperature stability diagram	46
3.5	Coulomb diamonds	47
3.6	Addition energy	50
3.7	Lever arm calculation	51
4.1	Dangling bonds	57
4.2	Charge noise induced relaxation	60
4.3	1/f charge noise	62
4.4	Charge Noise measurement method	63
4.5	Charge noise across a coulomb peak	64
4.6	Vertical manipulation	66
4.7	Lateral manipulation	68

4.8	Bistable fluctuator simulations	70
4.9	Si/SiGe charge noise temperature dependence	72
4.10	Temperature dependence of charge noise	74
4.11	Variation of γ with temperature	75
4.12	Arrhenius plot for individual fluctuator spectra	77
4.13	First electron transition	78
4.14	Single electron charge noise measurement method	80
4.15	Charge noise with electron occupancy	82
4.16	T_2^* simulation	83
5.1	Energy selective readout	89
5.2	Readout	91
5.3	Measurement fidelity	95
5.4	Loading Spectrum	98
5.5	T_1 measurement sequence	100
5.6	Valley states	101
5.7	Spin-valley hotspot	104
5.8	Tuning	106
5.9	Magnetic field anisotropy	108
5.10	Charge noise measurement via spin relaxation	113
A.1	Double dot	138
A.2	Double dot potential schematic	141
A.3	Double dot stability diagram	142
A.4	Triple points	143

List of Tables

5.1 Readout and initialization errors	97
---	----

APPENDIX A

Double Quantum Dots

A.1 Double dot energy

While a single quantum dot is sufficient to define a single qubit, to fully exploit the potential of quantum computers, we require the ability to entangle two qubits. While there are proposals for methods to transfer quantum information over long distances, the simplest and most direct method of quantum information transfer between adjacent qubits is to directly transmit an electron from one quantum dot to another. The double quantum dot system is an extension of the single dot system that enables more dynamic measurement and manipulation of electrons. Similarly to how a single dot can be considered an artificial atom, double dots (or more complex systems of multiple dots) are analogous to artificial molecules.

The double dot system is constructed using a second quantum dot that is tunnel coupled to the first, as demonstrated in Fig. A.1. Each dot is also coupled to a single reservoir, to allow movement of electrons in and out of each dot independently. In addition, the on-site potential energy of each dot is manipulated through a plunger gate, which for a single dot had an applied voltage V_G . To distinguish between the potential applied to each dot in this case, we use V_{G1} and V_{G2} instead. Additionally, the capacitive coupling of the dot to the plunger gate is C_{G1} and C_{G2} , and the relevant alpha factor α_{G1} and α_{G2} . The dots have N_1 and N_2 electrons respectively. Similarly, each dot has its own self-capacitance, which is the sum of all capacitances the dot experiences, which we denote $C_{1\Sigma}$ and $C_{2\Sigma}$. For each dot, this will be:

$$C_{i\Sigma} = C_{\text{reservoir}} + C_{Gi} + C_{12} \quad (\text{A.1})$$

Where *reservoir* is the Source for dot 1 and Drain for dot 2, and C_{12} is the inter-dot capacitance. The charge on each dot is, as before, the sum of $C_i V_i$ for each relevant capacitance and voltage acting on the dot, therefore:

$$Q_{1(2)} = C_{S(D)}(V_{1(2)} - V_{S(D)}) + C_{G1(2)}(V_{1(2)} - V_{G1(2)}) + C_{12}(V_{1(2)} - V_{2(1)}) \quad (\text{A.2})$$

This can be more succinctly expressed in matrix form, as $\vec{Q} = \mathbf{C}\vec{V}$ [Wie02], where:

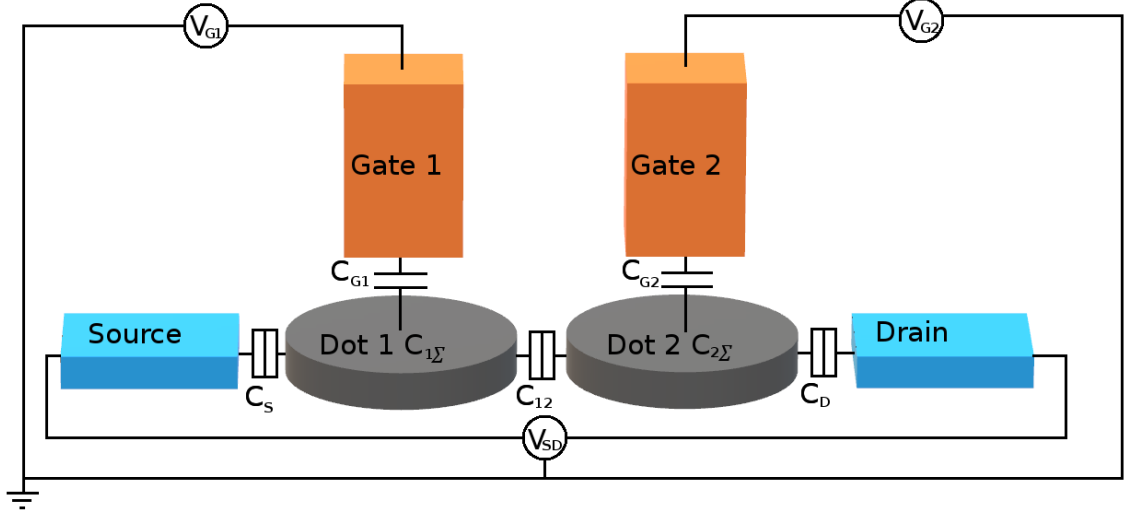


Figure A.1: Double dot Model of a double quantum dot. Here two dots are defined in series, labelled Dot 1 and Dot 2. Each has a self-capacitance $C_{i\Sigma}$. Each is capacitively coupled to a gate with capacitance C_{Gi} and applied voltage V_{Gi} . Each has a tunnel coupling to one of the two electron reservoirs, source and drain, and there is a tunnel junction between them with capacitance C_{12} . An applied voltage V_{SD} forms the potential difference across the two reservoirs.

$$\vec{Q} = \begin{pmatrix} Q_1 + C_S V_S + C_{G1} V_{G1} \\ Q_2 + C_D V_D + C_{G2} V_{G2} \end{pmatrix} \quad (\text{A.3})$$

$$\mathbf{C} = \begin{pmatrix} C_{1\Sigma} & -C_{12} \\ -C_{12} & C_{2\Sigma} \end{pmatrix} \quad (\text{A.4})$$

$$\vec{V} = \begin{pmatrix} V_1 \\ V_2 \end{pmatrix} \quad (\text{A.5})$$

$\mathbf{C}$ is often referred to as the capacitance matrix of the system. In general, the diagonal elements $\mathbf{C}_{ii}$ are the self-capacitance for dot i , while the off-diagonal elements $\mathbf{C}_{ij}$ are the inter-dot capacitance between dots i and j . For the double dot, we obtain:

$$\begin{pmatrix} Q_1 + C_S V_S + C_{G1} V_{G1} \\ Q_2 + C_D V_D + C_{G2} V_{G2} \end{pmatrix} = \begin{pmatrix} C_{1\Sigma} & -C_{12} \\ -C_{12} & C_{2\Sigma} \end{pmatrix} \begin{pmatrix} V_1 \\ V_2 \end{pmatrix} \quad (\text{A.6})$$

The voltage on dot i is given by $\vec{V}_i = C_{i\Sigma}^{-1}(\vec{Q}_i + C_{12}\vec{V}_j)$. The capacitance matrix is invertible, with $\det(\mathbf{C}) = C_1 C_2 - C_{12}^2$, meaning we can rewrite Eqn. A.6 to obtain the on-site potentials V_i :

$$\begin{pmatrix} V_1 \\ V_2 \end{pmatrix} = \frac{1}{C_1 C_2 - C_{12}^2} \begin{pmatrix} C_{1\Sigma} & -C_{12} \\ -C_{12} & C_{2\Sigma} \end{pmatrix} \begin{pmatrix} Q_1 + C_S V_S + C_{G1} V_{G1} \\ Q_2 + C_D V_D + C_{G2} V_{G2} \end{pmatrix} \quad (\text{A.7})$$

As in the single-dot case, we obtain the energy of the system using $U = \frac{Q^2}{2C}$, which in matrix format is:

$$U = \frac{1}{2} \vec{Q} \cdot \mathbf{C}^{-1} \vec{Q} \quad (\text{A.8})$$

For the case where $V_S = V_D = 0$ and $Q_i = -N_i|e|$, we obtain:

$$U(N_1, N_2) = \frac{1}{2} N_1^2 E_{C1} + \frac{1}{2} N_2^2 E_{C2} + N_1 N_2 E_{C12} + f(V_{G1}, V_{G2}) \quad (\text{A.9})$$

Where:

$$\begin{aligned} f(V_{G1}, V_{G2}) = & \frac{1}{-|e|} [C_{G1} V_{G1} (N_1 E_{C1} + N_2 E_{C12}) + C_{G2} V_{G2} (N_1 E_{C12} + N_2 E_{C2})] \\ & + \frac{1}{e^2} \left[\frac{1}{2} C_{G1}^2 V_{G1}^2 E_{C1} + \frac{1}{2} C_{G2}^2 V_{G2}^2 E_{C1} + C_{G1} V_{G1} C_{G2} V_{G2} E_{C12} \right] \end{aligned} \quad (\text{A.10})$$

And we define:

$$E_{C1} = \frac{e^2}{C_{1\Sigma}} \left(\frac{1}{1 - \frac{C_{12}^2}{C_{1\Sigma} C_{2\Sigma}}} \right) \quad (\text{A.11})$$

$$E_{C2} = \frac{e^2}{C_{2\Sigma}} \left(\frac{1}{1 - \frac{C_{12}^2}{C_{1\Sigma} C_{2\Sigma}}} \right) \quad (\text{A.12})$$

$$E_{C12} = \frac{e^2}{C_{12}} \left(\frac{1}{\frac{C_{1\Sigma} C_{2\Sigma}}{C_{12}^2} - 1} \right) \quad (\text{A.13})$$

as the dot 1 charging energy, dot 2 charging energy, and inter-dot electrostatic coupling energy respectively. The coupling energy can be thought of as the change in the energy of dot 1 when an electron is loaded into dot 2 (and vice versa). The individual charging energies E_{Ci} have the same form as seen for the single dot, with an additional correction factor to account for the inter-dot coupling. If we consider the case where the dots are completely decoupled, meaning $C_{12} = 0$ and $E_{C12} = 0$, the energy of the system simplifies to:

$$U(N_1, N_2) = \frac{(-N_1|e| + C_{G1} V_{G1})^2}{2C_{1\Sigma}} + \frac{(-N_2|e| + C_{G2} V_{G2})^2}{2C_{2\Sigma}} \quad (\text{A.14})$$

which is the energy of two individual dots. Conversely, we can increase the inter-dot capacitance C_{12} until there is no meaningful potential barrier between the quantum dots and the system resembles a single dot controlled by two gates. By setting the inter-dot capacitance to be dominant, i.e. $C_{12} \rightarrow C_{i\Sigma}$, the system energy becomes:

$$U(N_1, N_2) = \frac{(-N_1|e| + C_{G1}V_{G1})^2}{2(C_{1\Sigma} + C_{2\Sigma} - 2C_{12})} \quad (\text{A.15})$$

This is analagous to a single-dot energy with modulation via two gates. Therefore by changing the capacitance between the dots C_{12} , we are able to fully tune the system from two isolated dots to a single dot, with various stages of intermediate coupling.

A.2 Double dot transport

Similarly to the single dot situation, it is preferable to express the system in terms of electrochemical potentials when discussing movement of electrons. We define $\mu_1(N_1, N_2)$ as the energy required to add the N_1 th electron to dot 1 with N_2 electrons on dot 2. In contrast to the single dot, we now have to take account of the electron occupancy of the adjacent dot, as this has a capacitive effect on dot 1 due to the inter-dot coupling. Following the same principle as for the single dot, and continuing to consider the purely classical case:

$$\begin{aligned} \mu_1(N_1, N_2) &= U(N_1, N_2) - U(N_1 - 1, N_2) \\ &= (N_1 - \frac{1}{2})E_{C1} + N_2E_{C12} - \frac{1}{|e|}(C_{G1}V_{G1}E_{C1} + C_{G2}V_{G2}E_{C12}) \end{aligned} \quad (\text{A.16})$$

And equivalently for dot 2:

$$\begin{aligned} \mu_2(N_1, N_2) &= U(N_1, N_2) - U(N_1, N_2 - 1) \\ &= (N_2 - \frac{1}{2})E_{C2} + N_1E_{C12} - \frac{1}{|e|}(C_{G1}V_{G1}E_{C12} + C_{G2}V_{G2}E_{C2}) \end{aligned} \quad (\text{A.17})$$

The first term is the single-dot potential due to the number of electrons in the dot. The second term is the additional energy induced by the number of electrons in the adjacent dot. Finally, the third term is the energy induced by each of the two gates V_{G1} and V_{G2} , directly and across the interdot coupling. As before, E_{C1} and E_{C2} are the addition energies of dot 1 and 2, and in the classical regime are identically the charging energies of the dots. E_{C12} is the change in energy experienced by the a dot when the adjacent dot is charged with an additional electron.

This discussion so far has been entirely classical. The quantum part of the addition spectrum is induced by the orbital level splitting, as in the single dot case. We introduce the additional term E_{N_1-1, N_1} to account for the discrete level splitting:

$$\mu_1(N_1, N_2) - \mu_1(N_1 - 1, N_2) = E_{C1} + E_{N_1-1, N_1} \quad (\text{A.18})$$

and similarly for the addition energy of dot 2. In order to analyse the energy spectrum of the quantum dots, the current through the double dot system can be measured for each linear combination of voltages V_{G1} and V_{G2} . This results in a 2D map known as a stability diagram.

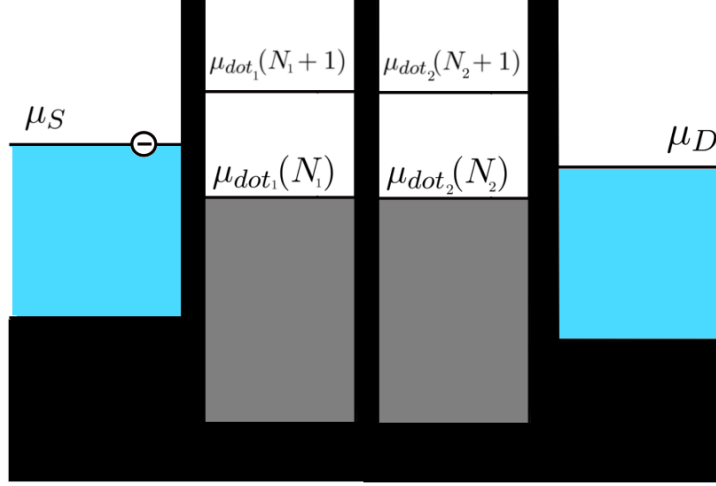


Figure A.2: Double dot potential The potentials of a double dot system in coulomb blockade. In order for blockade to be lifted, a potential level must reside within the bias window for each dot, otherwise transport is prevented through the system. In each dot, the energy levels are filled up to $\mu_1(N_1 + 1)$ $\mu_2(N_2 + 1)$.

The potential landscape of a double dot system in series is indicated in Fig. A.2. For this case we have a new requirement for electron transport through the system, $\mu_S \geq \mu_1(N_1) \geq \mu_2(N_2) \geq \mu_D$. When this condition is fulfilled, current can be seen through the quantum dot. At low temperature and bias, this naturally leads to a blockade pattern similar to that of the single dot case. However, the situation becomes more complex as current is only allowed when the blockade is lifted in both dots simultaneously. As such, signal is seen at the "triple points", where both dots and the bias are in resonance. If we consider the linear regime, where $V_{SD} = 0$, then the condition becomes $\mu_S = \mu_1(N_1) = \mu_2(N_2) = \mu_D$. Similar to the coulomb peaks for a single dot, these triple points are separated by an energy corresponding to the addition of one extra electron. As indicated in Fig. A.3, these triple points can be used to reconstruct the charge configuration of the double dot system at any point in the gate voltage space.

In Fig. A.3a, the situation for two completely decoupled dots is shown. Here $C_{12} = 0$, and the charge configuration of each dot is entirely independent of the other. Inside each "cell" of the grid pattern, the charge number of each dot is constant. However, at a finite C_{12} , the shape of the stability diagram changes to a hexagonal pattern, as in Fig. A.3b. Notably, the triple points, degenerate at $C_{12} = 0$, become separated into two for each electron charge configuration. These two triple points correspond to different charge

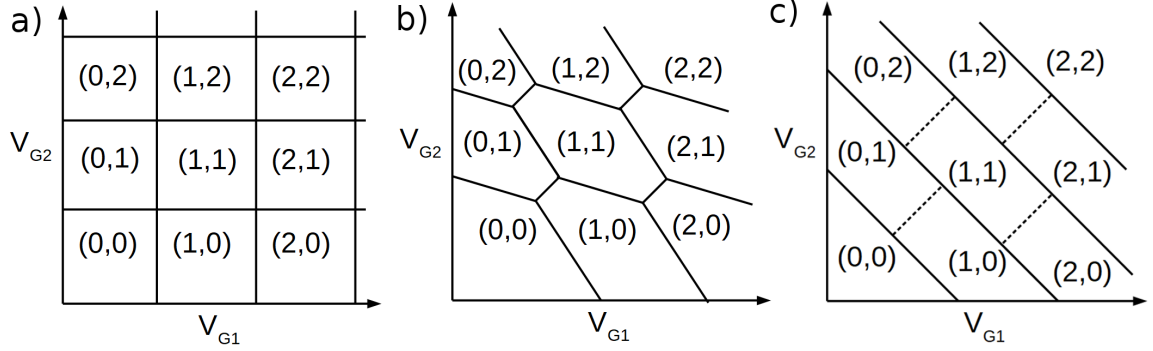


Figure A.3: Double dot stability diagram Stability diagrams for a double dot with different tunnel couplings. The lines indicate the gate voltages at which a change in the electron number is seen. a) For $C_{12} = 0$, the quantum dots are completely decoupled and the stability diagram resembles the single dot case for each dot, with no interaction between them. b) At finite C_{12} , a characteristic hexagonal pattern is seen, arising from the cross-capacitance between the dots. c) At $C_{12} \rightarrow C_{i\Sigma}$, meaning the inter-dot capacitance dominates in the system, a single dot stability diagram is seen, with the chemical potential of the dot being equally modified by each gate.

transfer processes. This results in the characteristic honeycomb pattern associated with coupled quantum dots.

The dimensions of each honeycomb cell can be determined from:

$$\mu_1(N_1, N_2, V_{G1}, V_{G2}) = \mu_1(N_1 + 1, N_2, V_{G1} + \Delta V_{G1}, V_{G2}) \quad (\text{A.19})$$

where we find

$$\Delta V_{G1} = \frac{|e|}{C_{G1}} \quad (\text{A.20})$$

and $\Delta V_{G2} = \frac{|e|}{C_{G2}}$ for dot 2. Since the splitting of the triplet points corresponds to a shift in the potential level for dot 1 when an electron is loaded into dot 2, we can also extract the linear shift in V_{G1} between the triple points:

$$\mu_1(N_1, N_2, V_{G1}, V_{G2}) = \mu_1(N_1, N_2 + 1, V_{G1} + \Delta V'_{G1}, V_{G2}) \quad (\text{A.21})$$

Giving:

$$\Delta V'_{G1} = \frac{|e|C_{12}}{C_{G1}C_{\Sigma}2} = \Delta V_{G1} \frac{C_{12}}{C_2\Sigma} \quad (\text{A.22})$$

and equivalently $\Delta V'_{G2} = \Delta V_{G2} \frac{C_{12}}{C_1\Sigma}$ for dot 2.

At the point A in Fig. A.4a, the cycle of charge transfer is $(N_1, N_2) \rightarrow (N_1 + 1, N_2) \rightarrow (N_1, N_2 + 1) \rightarrow (N_1, N_2)$ corresponding to a single electron moving from the source contact through the system to the drain. At the point B, however, the charge transfer cycle instead involves the $(N_1 + 1, N_2 + 1)$ state, and the cycle instead is $(N_1 + 1, N_2 + 1) \rightarrow$

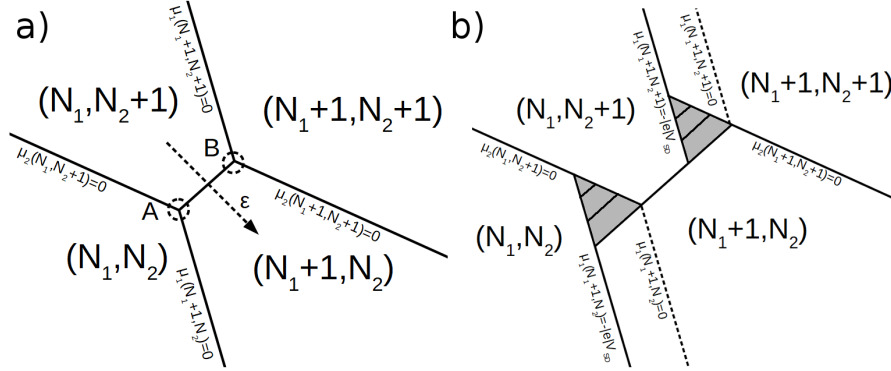


Figure A.4: Triple points At a finite tunnel coupling in the double dot system, two triple points are seen at the intersections between the (N_1, N_2) , $(N_1, N_2 + 1)$ and $(N_1 + 1, N_2)$ states, and the $(N_1, N_2 + 1)$, $(N_1 + 1, N_2 + 1)$, and $(N_1 + 1, N_2)$ states respectively. a) In the zero bias case the triple points are located at the intersections of the electron transition lines for each dot. b) When a finite bias is applied, bias triangles are seen at the triple points. Within the bias triangle, charge transport is allowed through the quantum dots. Resonance lines within the dot are seen when single dot levels are aligned within the double dot system.

$(N_1 + 1, N_2) \rightarrow (N_1, N_2 + 1) \rightarrow (N_1 + 1, N_2 + 1)$. This can be simply interpreted as movement of a hole in the opposite direction. The difference in energy between these two triple points is determined by the inter-dot capacitance C_{12} . The dashed line in Fig. A.4a indicates the detuning axis, which is perpendicular to the degeneracy line between the two triple points where $\mu_1(1,0) = \mu_2(0,1)$.

Fig. A.3c shows the case where $C_{12} \rightarrow C_{i\Sigma}$. Here the coupling between the two dots is strong enough that they have merged to form a single dot with equal coupling to each gate. In this regime, the notion of separated dots is no longer useful as there is only a single potential minimum, and the system can be considered as a single quantum dot.

A.3 Bias triangles

So far we have considered the linear regime, with $V_{SD} \simeq 0$. If a finite bias is applied, we enter the non-linear regime, where the bias window is opened. Then $\mu_S = -|e|V_{SD}$, and $\mu_D = 0$. Since the bias is applied asymmetrically, we have to account for its effect on the electrostatic energy of the system. Taking Eqn. A.9, and substituting $C_{G1}V_{G1}$ with $C_{G1}V_{G1} + C_S V_{SD}$, where C_S is the capacitance to the source from dot 1, we obtain the energy of the system. No correction needs to be made for dot 2 as there is no applied voltage on the drain, and in this model we assume the capacitance from dot 2 to the source is zero.

The condition for transport becomes $\mu_S \geq \mu_1(N_1) \geq \mu_2(N_2) \geq \mu_D$, where $\mu_S = -|e|V_{SD}$. This gives a set of three conditions, $\mu_S \geq \mu_1(N_1)$, $\mu_1(N_1) \geq \mu_2(N_2)$, and $\mu_2(N_2) \geq \mu_D$, which form the boundaries of the conduction region at each triple point. This results in a triangular shape, commonly known as bias triangles, as in Fig. A.4b.

The dimensions of the triangles, given by δV_{G1} and δV_{G2} , are determined by the bias:

$$\alpha_{G1}\delta V_{G1} = |eV_{SD}| \quad (\text{A.23})$$

$$\alpha_{G2}\delta V_{G2} = |eV_{SD}| \quad (\text{A.24})$$

This can be used to determine the lever arm of each gate, α_{Gi} . Once the lever arm is known, then the addition energy E_{add} can be determined for each hexagonal cell by measuring ΔV_{Gi} and using the relation $E_{add} = \Delta V_G \alpha_G$.

When the bias window is large, multiple energy levels can enter the bias window. This means that excited states can play a role in tunnelling through the double dot system. In particular, when the electrochemical potentials of two levels are in resonance, there is a conductance peak, with the background conductance in the bias triangle induced by inelastic processes [Fuj98] and co-tunnelling. This gives rise to resonance lines, as seen in Fig. A.4b. These resonance lines correspond to an excited state matching the chemical potential of another state, allowing fast tunnelling (and thus increased conductance) through the double dots. These resonances can be used to determine the separation between the ground and excited states.