

Optimizing a superconducting transistor to minimize parasitic capacitance and maximize electron mobility

Author: Charlotte Lundquist

Supervisor: Erik Lind

Examiner: Daniel Sjöberg

Master's Thesis of Science

Department of Electrical and Information Technology

Faculty of Engineering | LTH

Lund University

March 2026



LUND
UNIVERSITY

Abstract

The scaling of high-mobility and superconducting transistors places strict demands on electrostatic control, making the gate capacitance a key factor for both performance and operational stability. In hybrid structures that combine semiconductors and superconductors, parasitic capacitances become particularly important, as fringing fields and heterogeneous material interfaces cause the total capacitance to deviate from ideal models. The relationship between capacitance and geometric and dielectric parameters is therefore an active area of research.

This work investigates how variations in gate length, dielectric width, oxide thickness and dielectric permittivity affect the total gate capacitance in a superconducting quantum-well field-effect-transistor (QWFET). A two-dimensional electrostatic COMSOL model is used to extract both intrinsic and parasitic capacitance components through a frequency-domain perturbation analysis. A separate one-dimensional Schrödinger-Poisson model is used to describe quantum confinement and to compute the electron mobility.

The results show that the capacitance increases almost linearly with gate length, with deviations at short lengths where fringing fields dominate. A wider dielectric reduces the capacitance by weakening the lateral electrostatic coupling, while the oxide thickness exhibits a non-monotonic behavior: the capacitance increases initially due to enhanced fringing-fields spreading, but decreases again at larger thicknesses due to a weakened gate coupling. Increasing the dielectric permittivity raises the overall capacitance level while leaving its scaling with gate length essentially unchanged.

Linear extrapolation of the capacitance-gate length relation enables extraction of the parasitic capacitance, which decreases slightly with increasing dielectric width. Oxide thickness shows no clear trend in the extracted parasitic capacitance, while increasing dielectric permittivity leads to a pronounced rise in the parasitic contribution.

The Schrödinger-Poisson model shows strong electron confinement in the quantum well and confirms that impurity scattering is the dominant mechanism affecting the mobility under the studied conditions. Mobility decreases with increasing impurity concentration and reflects the spatial overlap between the confined wavefunction and the doped regions. These results highlight a trade-off between strong electrostatic gate control and high mobility, arising from the opposing requirements on confinement and scattering.

Declaration

I hereby affirm that this Master thesis was composed by myself, that the work herein is my own except where explicitly stated otherwise in the text. This work has not been submitted for any other degree or professional qualification except as specified; nor has it been published. Where other people's work has been used (either from a printed source, internet, or any other source), this has been carefully acknowledged and referenced.

In this thesis, Microsoft Copilot has assisted with creating the MATLAB-scripts that can be found in Appendix B. Furthermore, it has been used to provide alternative phrasings and second opinions on the writing, and to clarify certain technical concepts. Lastly, it was used to help identify relevant literature that I subsequently evaluated and used as sources. All analysis, conclusions, and academic judgments are my own.

Sammanfattning

Moderna elektroniska komponenter bygger på snabba och energieffektiva transistorer, och i takt med att dessa fortsätter att krympa i storlek blir det allt viktigare att upprätthålla god elektrostatisk kontroll. Ett centralt begrepp i detta sammanhang är gate-kapacitansen, vilken avgör hur effektivt styrgrindens spänning kan påverka laddningen i kanalregionen.

Detta blir än mer komplicerat i avancerade komponenter som kombinerar halvledare med supraledande material. Den typen av strukturer har ofta komplexa geometrier och materialgränser som ger upphov till fransfält - elektriska fält som sprider sig från sin ideala riktning - samt parasitiska kapacitanser. Dessa effekter är svåra att beskriva med enkla analytiska modeller, vilket gör det nödvändigt att använda numeriska modeller för att förstå hur geometri och materialval påverkar den totala kapacitansen.

Transistorn som studeras i detta arbete är en kvantbrunns-fälteffekttransistor (QWFET), där elektroner är inneslutna i ett mycket tunt lager-kvantbrunnen. När elektronerna sitter fast på det här sättet i kvantbrunnen kan de bara röra sig i två riktningar, vilket gör att de bildar en tvådimensionell elektrongas. De här förutsättningarna gör att elektronerna kan röra sig väldigt snabbt, vilket är bra för att bygga snabba och energieffektiva komponenter.

För att förstå hur transistorn fungerar byggdes två olika datormodeller i ett verktyg som heter COMSOL. Den första modellen var tvådimensionell och användes för att räkna ut gate-kapacitansen, vilket är ett mått på hur väl gaten kan styra elektronerna i kanalen. Modellen gjorde det även möjligt att skilja på den intrinsiska "inbyggda" kapacitansen och den oönskade, parasitiska, kapacitansen vilken kommer från fransfält. Den andra modellen är en Schrödinger-Poisson-modell, vilken var endimensionell men också byggd i COMSOL. Den beskriver hur elektronerna beter sig inne i kvantbrunnen enligt kvantmekanikens regler. Den ligger till grund för att kunna räkna ut elektronernas mobilitet, och hur mobiliteten påverkas av olika designval.

Fyra olika parametrar undersöktes för att se hur de påverkar kapacitansen. Resultaten visar tydliga trender. När gatelängden ökar ökar kapacitansen nästan linjärt, eftersom en längre gate har större area och därmed kan påverka fler elektroner. När dielektrikumets bredd ökar minskar kapacitansen, eftersom avståndet mellan gate och kanal blir större och kopplingen försvagas.

Oxidtjockleken beter sig mer komplicerat. Teoretiskt borde en tjockare oxid minska kapacitansen, men modellen visade motsatt trend. Detta kan förklaras av att en tjockare oxid ger starkare fransfält, vilket ökar kapacitansen. Materialets dielektriska permittivitet påverkar också kapacitansen: högre permittivitet ger högre kapacitans, vilket stämmer med teorin. Genom att extrapolera kapacitansen som funktion av gatelängd kunde även den parasitiska kapacitansen bestämmas. Den minskar när dielektrikumets bredd ökar, vilket följer samma trend som den vanliga kapacitansen.

Schrödinger-Poisson-modellen visar att elektronerna är fast i kvantbrunnen. Mobilitetsberäkningarna framhäver att mobiliteten främst begränsas av spridning mot dopatomer, detta är på grund av att elektronernas vågfunktion sträcker sig in i de dopade områdena, vilket gör att elektronerna oftare interagerar med dopatomerna. Varje kollision mellan elektroner och dopatomer bromsar elektronerna och gör att mobiliteten sjunker. Således spelar både materialval och hur dopingen placeras en stor roll för hur hög mobilitet som kan uppnås.

Tillsammans ger resultaten från de två modellerna en tydligare förståelse för hur geometri, material och kvantmekanik samverkar i supraledande QWFET-strukturer. Denna kunskap kan användas som vägledning vid design av framtida komponenter där låg parasitisk kapacitans och hög mobilitet är avgörande för snabb och energieffektiv funktion.

Contents

1	Introduction	1
1.1	Motivation	1
1.2	Problem formulation	2
1.3	Scope and delimitation	2
2	Theory	4
2.1	Fundamentals of Quantum Well Field-Effect Transistors	4
2.1.1	An introduction to superconductivity	4
2.1.2	BCS theory and Cooper pair formation	5
2.1.3	Quantum well field-effect transistors (QWFETs)	6
2.2	Capacitance	9
2.2.1	Capacitance and small-signal modeling	9
2.2.2	Fringing fields in two-dimensional geometries	12
2.2.3	Effective fringing length	12
2.2.4	Device electrostatics	13
2.3	Mobility in a 2DEG	13
2.3.1	Electron mobility in quantum wells	13
2.3.2	Combined design considerations: balancing capacitance, mobility and robustness .	17
3	Method	18
3.1	COMSOL Multiphysics	18
3.1.1	Introduction to COMSOL Multiphysics	18
3.1.2	Model 1	18
3.1.3	Model 2	24
3.1.4	MATLAB and COMSOL together: extracting mobility as a function of capacitance	28
3.2	How COMSOL finds solutions	29
4	Results	30
4.1	Extraction of the threshold voltage	30
4.2	Results for Model 1	32
4.2.1	Gate length influence	33
4.2.2	Dielectric width influence	35

4.2.3	Oxide thickness influence	35
4.2.4	Dielectric permittivity influence	36
4.2.5	Table to summarize the capacitive effect of the parameters on Model 1	37
4.2.6	Varying multiple parameters	37
4.2.7	Parasitic capacitance analysis	38
	Dielectric width influence	39
	Oxide thickness influence	40
	Dielectric permittivity influence	41
	Combined influence of geometry and material parameters	42
4.3	Results Model 2	42
4.3.1	Overview of the Schrödinger–Poisson model	42
4.3.2	Eigenstate and wavefunction characteristics	44
4.3.3	Form factor results	47
4.3.4	Screened Coulomb potential	48
4.3.5	Transport time and mobility	49
4.4	Gate dependent mobility and its relation to capacitance	50
4.4.1	Mobility as a function of gate voltage	50
4.4.2	Mobility as a function of capacitance	51
4.4.3	Interpretation and connection to device operation	52
5	Conclusion and Outlook	53
6	Bibliography	55
A	Appendix	57
A	A COMSOL guide on how to get started with COMSOL and set up the geometry	57
A.1	Prerequisites to complete before starting to build the geometry	57
A.2	The geometry of the model	58
A.3	The materials of the model	62
A.4	Boundary conditions	66
A.5	The physics of the model	67
A.6	Meshing	70
A.7	An example of a study used	71
B	MATLAB scripts	73
B.1	compute_Fq_double.m	73
B.2	main_mobility.m	74
B.3	mobility_run_energies.m	76

Chapter 1

Introduction

1.1 Motivation

The development of faster and more energy-efficient electronic components has increased the interest in transistors based on materials with extremely high mobility, and in some cases even superconducting properties. In these kinds of devices, the electrostatic control of charge carriers is crucial for both performance and operational stability. This control is largely governed by the gate capacitance, which determines how effectively the gate voltage can control the charge distribution in the channel.

The total gate capacitance consists of an intrinsic part and an unwanted parasitic contribution. Minimizing the parasitic component is essential, as it directly limits operating speed and power consumption [1]. These effects become even more pronounced in modern nanoscale geometries, where small variations in geometric or material parameters can cause significant changes in device behavior. For transistor structures that combine semiconductors with superconducting materials-such as the hybrid device studied in this work-the trade-off between low parasitic capacitance, device geometry and material constraints becomes even more complex.

To address this difficulty, it is necessary to understand how each part of the device contributes to the parasitic capacitance and how these contributions can be reduced without compromising other key parameters. As analytical models tend to struggle with capturing key effects such as fringing fields and heterogeneous material stacks, numerical modeling using a software such as COMSOL becomes essential to evaluate the impact of geometric variations and to compute the charge distribution inside the structure.

Evidently, there is a clear need for a detailed study that quantifies how geometric and dielectric parameters affect capacitance and charge control in complex transistor structures. Such understanding is essential not only for optimizing existing device designs, and for guiding the development of future components where high mobility and low parasitic capacitance are of central importance [2, 3].

1.2 Problem formulation

Parasitic capacitance in devices that combine semiconductors and superconducting materials is a relatively new research field, and its dependence on geometric and dielectric variations is still not well understood. The structure studied in this work has a complex geometry with heterogeneous material boundaries and significant fringing fields, which makes it difficult for analytical models to capture the capacitance with sufficient accuracy. This becomes particularly evident in nanoscale structures, where small dimensional variations can lead to large changes in the electrostatic environment.

The aim of this work is therefore to quantitatively evaluate how variations in gate length, dielectric width, oxide height, and dielectric permittivity affect the total capacitance and the resulting charge distribution in the studied hybrid structure. To analyze these relationships, a numerical model capable of resolving the complete electrostatics of the model and separating contributions from different parts of the structure is required.

An equally important part of the problem is to investigate how these parameters influence the electron mobility. In this type of transistor, the drain current I_D depends strongly on both the mobility μ and the gate capacitance C_g . In the long-channel approximation for saturation, this relationship is given by

$$I_D = \frac{1}{2} \mu C_g \frac{W}{L} (V_{GS} - V_T)^2$$

which shows that high mobility can only be fully exploited if the gate provides sufficiently strong capacitive control over the channel. This means that the gate can efficiently modulate the carrier density and the channel potential.

At the same time, the total device capacitance includes a parasitic component that does not contribute to channel control but still increases the RC delay. For high-speed operation it is therefore desirable to keep the parasitic capacitance low, while maintaining a sufficiently large intrinsic gate capacitance [1].

These requirements interact with the mobility. A larger separation between the gate and the channel reduces impurity scattering, and thereby increases mobility, but it also reduces the intrinsic gate capacitance. However, a smaller separation strengthens the gate's electrostatic control of the channel potential, but it may lower the mobility. This creates a fundamental trade-off between achieving high mobility, maintaining strong gate control, and minimizing parasitic capacitance.

1.3 Scope and delimitation

This work is limited to studying the electrostatic behavior of a hybrid structure consisting of semiconductor and superconducting materials. The analysis focuses on how variations in gate length, dielectric width, oxide thickness and dielectric permittivity affect the total capacitance and the resulting charge distribution. Other geometric or material variations are not considered. All simulations are performed under stationary and quasi-stationary conditions in COMSOL. This means that only the electrostatic solution is included, complemented by a frequency-domain perturbation step used to extract the capacitance. Time-dependence or transient behavior are not modeled. The bias conditions are kept fixed at $V_{drain} = V_{source} = 0$ V, with the gate voltage set to a chosen operating-point value. How the results change

with varying bias is left for future work.

Current transport, contact resistances and full quantum-mechanical modeling in two-dimensions are not included. Quantum-mechanical effects are instead handled in a one-dimensional model, where the aim is to analyze how the mobility varies with capacitance and with the chosen geometric and dielectric parameters. The mobility model only accounts for impurity scattering, which typically dominates in doped quantum wells at low temperatures [4], while other scattering mechanisms such as phonon scattering and alloy disorder are neglected. Consequently, the mobility should be interpreted in terms of trends and relative changes rather than as an accurate absolute value. Furthermore, the MATLAB implementation used for the mobility calculations relies on simplifying assumptions and COMSOL-derived parameters, which further limits the quantitative accuracy.

A simulation temperature of 50 K is used in COMSOL to improve numerical convergence. This choice may influence the accuracy of the results, as superconducting devices are normally characterized at significantly lower temperatures where the relevant superconducting effects become significant. The temperature dependence of the capacitance and mobility is thus not addressed in this work.

The modeling is limited to two dimensions. Three-dimensional effects, such as variations in width or local defects, are not included. The results should therefore be regarded as indicators of how different parameters affect capacitance and mobility, rather than as exact predictions for a specific physical device.

Chapter 2

Theory

QWFET-Quantum Well Field-Effect Transistor
2DEG-Two-Dimensional Electron Gas
BCS-Bardeen-Cooper-Schrieffer

2.1 Fundamentals of Quantum Well Field-Effect Transistors

The device modeled in this project is a superconducting quantum well field-effect transistor (QWFET). This section outlines the fundamental physical principles that govern its operation, including the material properties, quantum confinement, and electrostatic mechanisms that define its behavior.

2.1.1 An introduction to superconductivity

Superconductivity was first observed in 1911 by Heike Kamerlingh Onnes, one of the pioneers of low-temperature physics. He was the first to liquefy helium, which allowed him to reach temperatures close to absolute zero. His aim was to investigate how the electrical resistivity of materials behaves as the temperature approaches this limit, a question that was still unresolved at the time. When cooling mercury to approximately 4.2 K, Onnes observed that its resistivity did not only decrease, but vanished entirely [5].

This experiment marked the discovery of superconductivity. A material in the superconducting state exhibits two defining properties. Below a material-specific critical temperature, T_c , the electrical resistance drops to zero [6]. In addition, when both the temperature and the applied magnetic field are below their critical values, the material expels magnetic flux from its interior, a phenomenon known as the Meissner effect. This behavior distinguishes a superconductor from an ideal conductor, which would show zero resistance but would not expel magnetic fields [5, 6, 7].

In this project, superconductivity mainly affects the material properties of the contacts and the quantum well. The actual transistor operation is set by electrostatics and quantum confinement rather than by any superconducting switching mechanism. For that reason, the following sections focus on the Schrödinger-Poisson formalism, which describes the band structure, charge distribution, and confinement effects that define the QWFET channel.

2.1.2 BCS theory and Cooper pair formation

The discovery of superconductivity has had many applications in physics and technology, including applications in electronic devices such as transistors. Superconductivity can be explained on a microscopic level by the Bardeen–Cooper–Schrieffer (BCS) theory. This model was developed by the trio Bardeen, Cooper and Schrieffer in 1957, and describes how superconductivity arises because electrons form Cooper pairs through an attractive interaction [8].

Normally, electrons repel each other due to their equal charge. However, when an electron moves through a crystal lattice, it slightly displaces the positively charged ions. This distortion creates a region of increased positive charge density, as illustrated in Fig. 2.1. A second electron can be attracted to this region, and the resulting phonon-mediated interaction can overcome the Coulomb repulsion. Following this, the two electrons form a bound state known as a Cooper pair [9, 10].

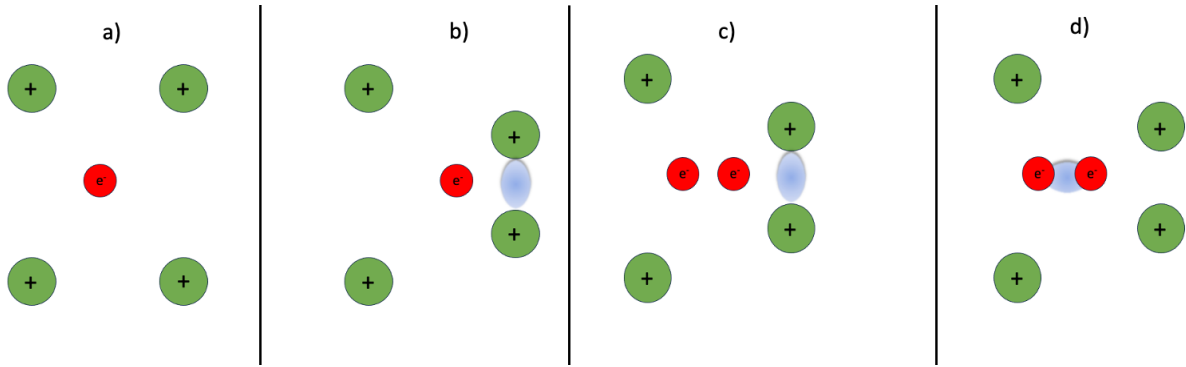


Figure 2.1: Schematic illustration of phonon-mediated electron pairing. (a) An electron moving through the lattice. (b) The electron attracts nearby positive ions, creating a local lattice distortion. (c) The resulting region of increased positive charge can attract a second electron. (d) The two electrons form a bound Cooper pair [10].

Electrons are fermions, but as they pair up in a Cooper pair, they possess an integer spin, and therefore behave as bosons. Moreover, the electrons do not exist as single particles anymore, but move in a shared, ordered quantum state, where all pairs occupy the same macroscopic wavefunction. As electrons form Cooper pairs, the pair occupies a lower energy state than the electrons would as individual particles. The bond between the electrons has a binding energy denoted Δ , and to break the bond an energy of 2Δ is required, and this is referred to as the superconducting gap which is visualized in Fig. 2.2.

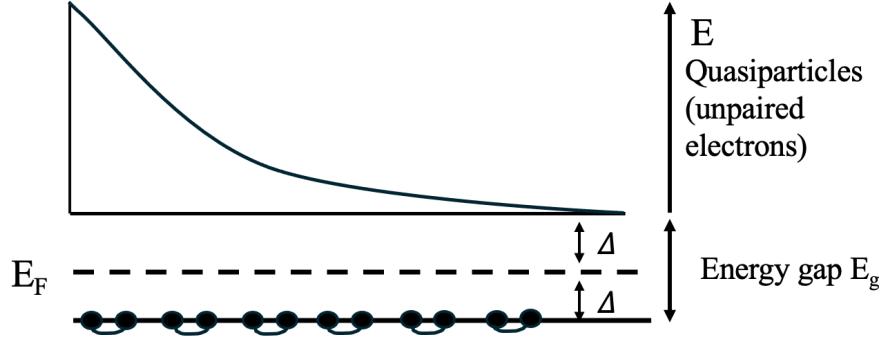


Figure 2.2: Schematic representation of the superconducting energy gap. Paired electrons form the Cooper-pair condensate below the Fermi level, while unpaired quasiparticles occupy states above it. Exciting a quasiparticle requires a minimum energy of Δ , so breaking a Cooper pair and creating two quasiparticles costs 2Δ .

The presence of the gap ensures that low-energy excitations are suppressed, allowing the superconductor to maintain coherence and zero resistance. Breaking a Cooper pair does not produce two free electrons, but two quasiparticles, each with a minimum excitation energy of Δ . The energy gap Δ originates from the pairing of electrons into Cooper pairs, which lowers the energy of the ground state. Since each quasiparticle lies an energy Δ above the ground state, breaking a Cooper pair produces two excitations of that sort, and the system must therefore supply an energy of 2Δ in order for the bond to break [11].

For the purposes of this project, BCS theory is not used to model the transistor operation itself. Instead, it provides the microscopic basis for the material parameters of the superconducting contacts and the quantum well. The behavior of the QWFET is determined by electrostatics and quantum confinement, which are treated in the following sections through the Schrödinger–Poisson formalism.

2.1.3 Quantum well field-effect transistors (QWFETs)

A quantum well field-effect transistor (QWFET) is a type of FET that leverages quantum mechanical effects to enhance device performance. The need for this kind of transistor stems from Moore’s law, which states that the number of transistors that can be implemented in an integrated circuit doubles every two years. This has driven the development of devices that maintain or improve performance while shrinking the dimensions. An example of this is quantum-well-based transistors, which can improve carrier transport even at nanometer scales.

A QWFET is built using several semiconductor layers, almost exclusively III–V materials, stacked on top of each other, where each region serves a specific function. The bottom layer, also known as the substrate, is usually a bulk semiconductor like GaAs or InP and provides mechanical support and lattice matching. On top of this, a barrier region is created with a material like InAlAs which has a larger

bandgap than the quantum well material to create a conduction band offset. This band offset forms a potential barrier that confines electrons to the quantum well. The quantum well is a thin film of a material with a lower bandgap, often InGaAs. Electrons accumulate in the well and become quantum-confined while forming the 2DEG. This confined electron sheet serves as the channel of the transistor, where current flows between source and drain. The source and drain contacts are formed on the sides of the gate region, providing access to the 2DEG in the quantum well.

To improve mobility, an undoped spacer layer is often inserted between the doped barrier and the quantum well. By physically separating the electrons in the well from the ionized dopants in the barrier, the spacer reduces Coulomb scattering and improves transport properties. Finally, a metal gate is deposited on top of the barrier to control the electron density in the quantum well and thereby modulate the conductivity of the channel. The layered structure of the QWFET is illustrated in Fig. 2.3, highlighting the role of each region in forming and controlling the quantum well channel [1, 12].

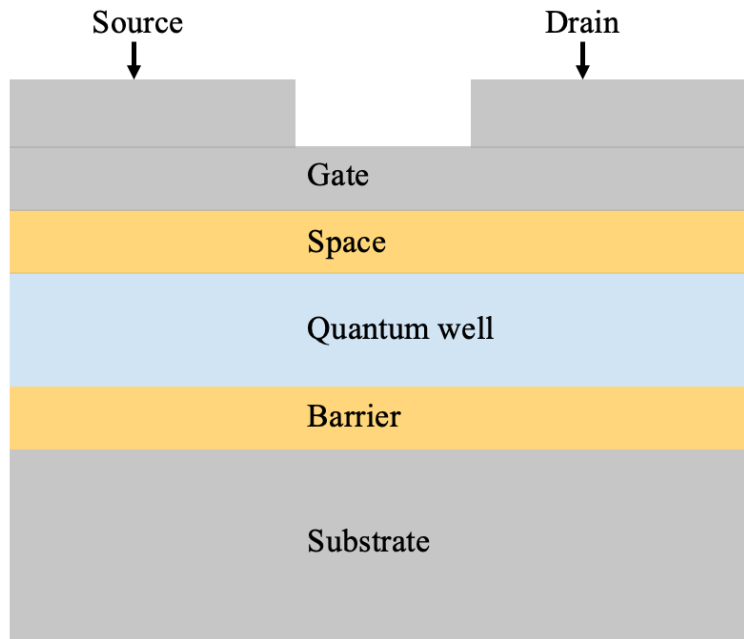


Figure 2.3: Schematic cross-section of a QWFET showing the layered structure and the formation of the 2DEG in the quantum well.

The heterostructure visualized in Fig. 2.3 yields a conduction band profile $V(z)$ in the growth direction, where the difference in band gap between the barrier and the quantum well material creates a potential well which confines electrons. This has a significant impact on the electronic states in the well, which divide into discrete subbands in accordance with the one-dimensional Schrödinger equation

$$-\frac{\hbar^2}{2} \frac{d}{dz} \left(\frac{1}{m(z)} \frac{d}{dz} \right) \psi_n(z) + V(z) \psi_n(z) = E_n \psi_n(z) \quad (2.1)$$

where $m^*(z)$ is the position-dependent effective mass, $V(z)$ is the conduction-band edge, and E_n are the quantized subband energies, $n = 1, 2, \dots$, each associated with one $\psi_n(z)$. Because electrons are confined in the growth direction z but free in the $x - y$ plane, the total wavefunction Ψ separates as

$$\Psi(x, y, z) = \psi(z)e^{i(k_x x + k_y y)} \quad (2.2)$$

where $\psi(z)$ describes the quantized motion in the well, k_x, k_y are the in-plane components of the electron wave vector, and $e^{i(k_x x + k_y y)}$ describes the free propagation in the $x - y$ -plane. This variable separation creates a subband structure where each E_n forms the bottom of a 2D subband, above which electrons can move freely in the x - y plane. The kinetic energy in this plane remains parabolic, so within a semiconductor quantum well system the total energy of an electron is equal to $E(k_x, k_y)$

$$E(k_x, k_y) = E_n + \frac{\hbar^2}{2m^*}(k_x^2 + k_y^2) \quad (2.3)$$

where m^* is the in-plane effective mass (assumed constant). This captures that electrons behave as free particles in the $x - y$ -plane, but are confined in the z -direction. The electrons occupying the quantized subband states (E_n, ψ_n) form the 2DEG, which is a thin sheet of carriers with high mobility in the quantum well. Since the electrons are confined to a narrow region, the number of allowed energy states is reduced, and they preferentially occupy the lowest subbands to minimize their energy. Because the electrons occupy states at the bottom of the well, far from the ionized dopants in the barrier, impurity scattering is strongly reduced. Together, the strong confinement and reduced impurity interaction result in QWFETs having a higher electron mobility than a conventional bulk device [13, 14].

To accurately model the conduction band profile $V(z)$, it is not sufficient to consider only the material band offsets. The presence of mobile carriers, ionized dopants and interface charges all modify the electrostatic potential $\phi(z)$, which is described by Poisson's equation

$$\frac{d}{dz}(\varepsilon(z)\frac{d\phi(z)}{dz}) = -\rho(z) \quad (2.4)$$

where $\varepsilon(z)$ is the position-dependent permittivity and $\rho(z)$ is the total charge density. The conduction-band edge $E_c(z)$ is therefore given by

$$V(z) = E_c(z) - q\phi(z) \quad (2.5)$$

which accounts for both the electrostatic band bending and the material band offsets, and is essential for determining the actual potential well that traps the two-dimensional electron gas [1]. The resulting conduction-band profile and charge distribution are illustrated in Fig. 2.4, which visualizes the material band offsets, electrostatic band bending, and the spatial extent of the ground-state wavefunction $|\psi_1(z)|^2$. The ionized donors in the barrier and the resulting charge density $\rho(z)$ in the well are also indicated, showing the relation between the donor layer, the confined electron state, and the resulting electrostatic potential.

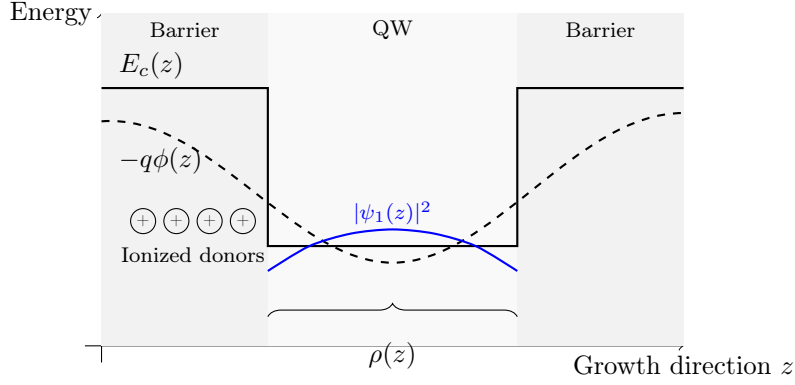


Figure 2.4: Schematic conduction band edge $E_c(z)$ and electrostatic potential $-q\phi(z)$ in a quantum well heterostructure, including the ground-state wavefunction $|\psi_1(z)|^2$ representing the 2DEG distribution.

2.2 Capacitance

2.2.1 Capacitance and small-signal modeling

The capacitance of a transistor is one of the most important parameters to consider when designing and fabricating devices used in modern technology. Understanding what determines the capacitance and how it influences device performance begins with its relation to charge. The continuity equation,

$$\frac{\partial Q}{\partial t} + \frac{\partial J}{\partial x} = 0 \quad (2.6)$$

where Q represents charge and J represents current density, states that any temporal change in charge must be balanced by a divergence in current density. Since capacitance describes how charge responds to changes in voltage, this equation highlights the strong coupling between charge and current that underlies capacitive behavior.

For a general transistor with three contacts—source, channel, and drain—the schematic in Fig. 2.5 illustrates the associated currents. The interaction between the gate and the channel determines how effectively charge is modulated, while the source and drain regions contribute to parasitic capacitances that influence device performance [15].

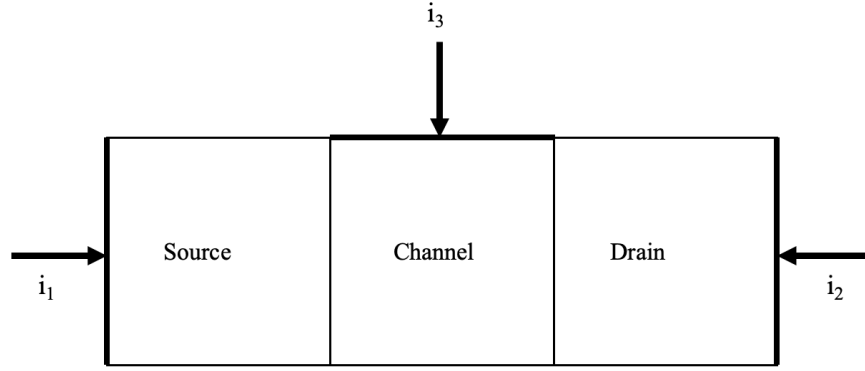


Figure 2.5: The schematic illustrates the currents i_1 , i_2 and i_3 in a transistor with a source, drain and channel region. Electrons enter through the source (i_1), propagate through the channel (i_3), and exit through the drain (i_2).

To describe capacitance more explicitly, a small-signal model is commonly used. In this representation, the capacitive couplings between the gate and source, and between the gate and drain, capture how a small voltage change at one terminal alters the charge distribution in the channel [1]. The small-signal network in Fig. 2.6 includes the elements C_{13} and C_{23} which represent the gate–source and gate–drain capacitances, respectively. In this small-signal model, C_{33} is also included, albeit not seen in the figure, and represents the self-capacitance of the channel. This means the charge stored in the channel when its own potential changes. However, it does not affect the parasitic capacitance, and is therefore not discussed further here in this work [15].

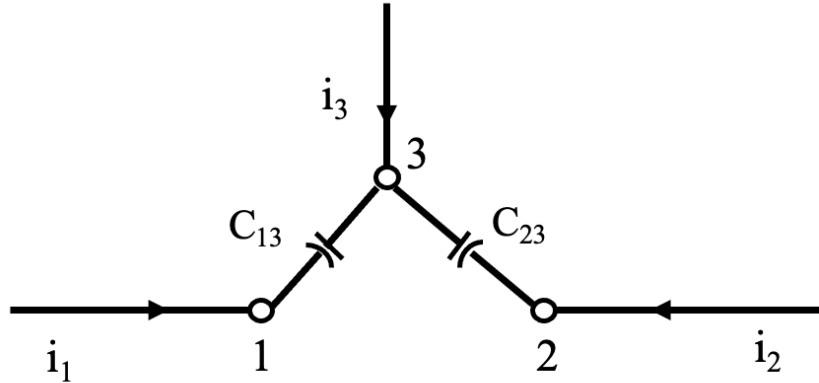


Figure 2.6: Small-signal capacitive network used to model the electrostatic coupling between the source (node 1), drain (node 2), and channel (node 3). The elements C_{13} and C_{23} represent the gate–source and gate–drain capacitances, respectively, and the currents i_1 , i_2 and i_3 represent the corresponding displacement currents at each node.

Time-dependent voltages across the capacitors cause the electric fields between the terminals to vary, redistributing charge on the capacitor plates and generating displacement currents. A displacement current arises from a time-varying electric field rather than from the physical transport of charge carriers, and is determined by the rate of change of the electric displacement field [16]. According to the charge-based definition of the small-signal capacitances, the displacement current associated with each terminal is given by $i_k = \partial Q_k / \partial t$. In the small-signal derivation, the source and drain terminals are assumed to be AC-grounded, that is $V_1 = V_2 = 0$. Therefore, the only time-varying potential is V_3 , and $V_3 = V_{31}$. Under this assumption, the charge relation in [15] reduces to

$$0 = -C_{13} \frac{\partial V_{31}}{\partial t} - C_{23} \frac{\partial V_{31}}{\partial t} + C_{33} \frac{\partial V_{31}}{\partial t} \equiv i_1 + i_2 + i_3, \quad (2.7)$$

which expresses Kirchhoff's current law at the channel node in the charge-based formulation for the special case where only the gate-channel voltage varies, as discussed in [15]. This relation shows that the total capacitive current flowing into the channel is the sum of the contributions from each coupling.

The channel of a QWFET is formed by the two-dimensional electron gas confined within a quantum well, typically created by an InGaAs/InAlAs heterostructure. The conduction-band profile restricts electron motion in the growth direction, quantizing it into discrete subbands, while allowing free motion in the plane of the well. In such a structure, the gate capacitance quantifies how the charge stored in the quantum well responds to changes in the gate voltage V_{gate} . The total gate capacitance C_g is composed of an intrinsic part, C_i , and a parasitic part, C_p

$$C_g = C_i + C_p, \quad (2.8)$$

where

$$C_p = C_{gs} + C_{gd}. \quad (2.9)$$

In the small-signal notation of Fig. 2.6, C_{gs} and C_{gd} correspond to C_{13} and C_{23} , respectively. Here, C_{gs} and C_{gd} arise from electrostatic coupling between the gate and the access regions. Fig. 2.7 illustrates the intrinsic and parasitic contributions.

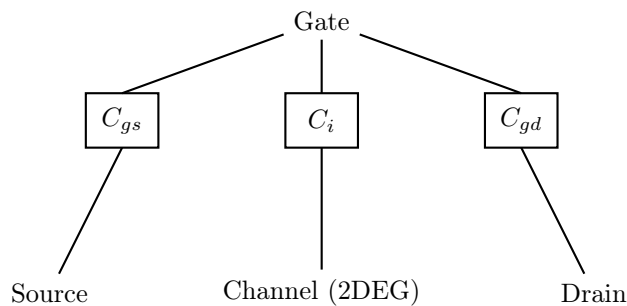


Figure 2.7: Conceptual illustration of intrinsic and parasitic capacitances in a QWFET.

Having defined the parasitic contribution, it is useful to relate it to the intrinsic capacitance in order to understand the total gate capacitance. The intrinsic part corresponds to the channel-related small-signal capacitances introduced earlier, while the parasitic part originates from the gate's coupling to the access

regions [17].

2.2.2 Fringing fields in two-dimensional geometries

In an ideal one-dimensional MOS structure, the gate capacitance is determined primarily by the vertical oxide thickness, t_{ox} and behaves as $C \propto 1/t_{ox}$. However, for certain device geometries, the capacitance does not follow this simple relationship. Instead, the capacitance is highly dependent on lateral fringing fields that spread between the gate, access regions and surrounding dielectric.

A fringing field is the portion of the electric field that extends outward from the edges of a conductor rather than remaining confined beneath it. Fringing fields occur because the electric field cannot remain strictly vertical near the gate edges, causing the field lines to spread laterally into the surrounding dielectric. Each fringing field contributes an additional capacitance beyond the ideal parallel-plate value, which deviates from the simple $1/t_{ox}$ dependence [18]. This phenomenon is illustrated in Fig. 2.8, which highlights how a thicker oxide causes a larger fringing-field region to form.

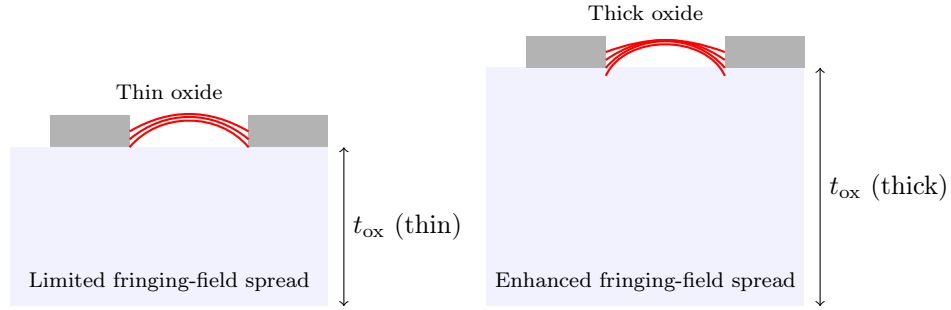


Figure 2.8: Schematic illustration of lateral fringing fields in a 2D geometry. An increased oxide thickness enlarges the region in which fringing fields can spread, which in turn increases the effective gate capacitance.

2.2.3 Effective fringing length

In many FET structures, the total gate capacitance can be expressed as a linear function of the gate length

$$C(L_{gate}) = aL_{gate} + 2C_p, \quad (2.10)$$

where a is the intrinsic gate-capacitance per unit length and $2C_p$ is the length-independent parasitic contribution due to the fringing fields between the gate-drain and the gate-source. The slope a gives the intrinsic capacitance scaling with gate length, and the intercept yields twice the parasitic capacitance.

The influence of fringing fields can be quantified through the effective fringing length L_{fringe} through the ratio

$$L_{fringe} = \frac{2C_p}{a} \quad (2.11)$$

which represents the equivalent gate length that would produce the same capacitance as the parasitic fringing fields alone. A large L_{fringe} implies strong parasitic capacitance and significant lateral spreading of the fringing fields, whereas a small value corresponds to weak parasitic contributions [15, 19].

The effective fringing length is particularly useful when evaluating device scalability. As the gate length is shortened, the parasitic capacitance becomes a larger fraction of the total gate capacitance, leading to deviations from ideal device scaling. By expressing the parasitic contribution in terms of an equivalent length L_{fringe} , this provides an intuitive measure of how far the device is from ideal electrostatic control. Maintaining a small L_{fringe} is therefore crucial to uphold strong gate control and ensure that the intrinsic capacitance is the primary contributor to the total capacitance.

This linear decomposition is commonly used as it separates the length-dependent intrinsic capacitance from the non-length-dependent parasitic part. Although approximate and somewhat difficult to implement on real devices, it provides a useful understanding for how fringing fields influence the total gate capacitance, especially in short-channel devices where the parasitic effects are progressively more significant [20].

2.2.4 Device electrostatics

In a semiconductor, the charge density ρ depends on the mobile carriers (p, n) and the fixed ionized dopants (N_D^+, N_A^-)

$$\rho(x, y) = q(p - n + N_D^+ - N_A^-), \quad (2.12)$$

where q is the elementary charge. This charge distribution serves as the source term in Poisson's equation,

$$\nabla^2 \phi(x, y) = -\frac{\rho(x, y)}{\varepsilon_r \varepsilon_0}, \quad (2.13)$$

where ε_r is the relative permittivity and ε_0 is the vacuum permittivity. Solving Poisson's equation provides the electrostatic potential distribution $\phi(x, y)$ throughout the device. In this work, the electrostatics are solved in a two-dimensional cross-section of the device.

The potential distribution is then used to obtain the channel charge by integrating the charge density over the quantum-well region

$$Q = \int \rho(x, y) dV, \quad (2.14)$$

which gives the channel charge $Q(V_{gate})$. The intrinsic gate capacitance is defined as

$$C_i = \frac{\partial Q}{\partial V_{gate}}, \quad (2.15)$$

and the derivative is evaluated numerically using finite differences between successive bias points. The capacitance is obtained by sweeping the gate voltage in COMSOL and extracting the corresponding charge [21, 1].

2.3 Mobility in a 2DEG

2.3.1 Electron mobility in quantum wells

The term mobility generally refers to both electron mobility μ_e and hole mobility μ_h . In this work, the focus is on electron mobility which will be denoted as simply μ . μ describes how quickly an electron

can move through a solid when driven by an electric field E . Mobility is a crucial parameter that strongly influences the operational speed and high-frequency response of devices like transistors or logic gates. Given this, maximizing the mobility is a key objective in the design and optimization of electronic devices [22].

A standard way to define mobility is

$$\mu = \frac{e\tau}{m^*} \quad (2.16)$$

where e is the elementary charge, m^* is the effective mass, and τ is the momentum relaxation time. The relaxation time is the average time an electron travels between two scattering events, and is highly influenced by mechanisms such as phonon interactions and ionized impurities. In high-quality materials which are clean and have smooth interfaces, the relaxation time is longer. A longer relaxation time allows the electron to travel farther without being interrupted, which in turn results in a higher mobility [1].

At low temperatures where few phonons are present, impurity scattering, such as that caused from ionized donors and acceptors, is what ultimately limits the mobility of electrons. To quantify how impurity scattering limits the relaxation time, it is fundamental to take into account not only how frequently electrons scatter, but also how much momentum electrons lose during it.

The directionality of the scattering is very important, and two different scenarios for this can be seen in Fig. 2.9. Forward scattering events, where the deflection angle θ is very small, result in the electron continuing to move in almost the same direction as before the collision. These types of events have minimal impact on the transport and hence the mobility, while scattering events like those on the right where the deflection angle θ is large results in substantial momentum loss. When an electron is scattered by an impurity, its initial wave vector $\mathbf{k}$ transforms to a new wave vector $\mathbf{k}'$. These two can be used to form the difference

$$\Delta\mathbf{k} = \mathbf{k}' - \mathbf{k}, \quad (2.17)$$

where $\Delta\mathbf{k}$ quantifies the momentum transferred from the electron to the impurity potential. The magnitude and direction of this momentum change depend on the scattering angle θ between $\mathbf{k}$ and $\mathbf{k}'$. This is illustrated in Fig. 2.9.

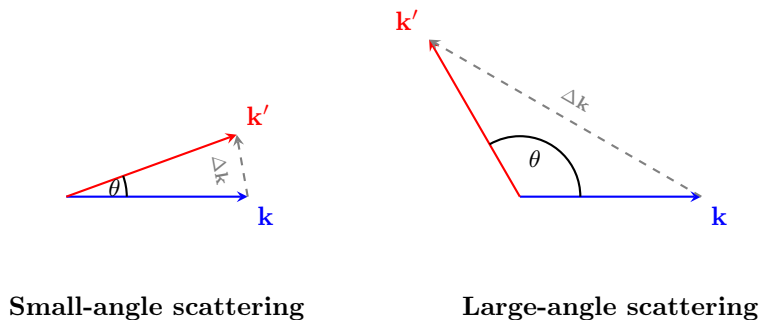


Figure 2.9: A comparison of small and large-angle scattering. The momentum loss due to the scattering is represented by the dashed vector $\Delta\mathbf{k} = \mathbf{k}' - \mathbf{k}$. For small-angle events, $\Delta\mathbf{k}$ is very small and subsequently also the momentum loss associated with them, while large-angle events result in substantial momentum loss.

Consequently, small-angle scattering events, which barely deflect the electron, have little effect on transport, whereas large-angle events have a significant role in momentum relaxation. This important contrast can be captured by introducing the transport relaxation time τ_{tr} , where each scattering event is weighted by the factor $(1 - \cos \theta)$ which measures how much momentum is lost. Thus, in a two-dimensional system, τ_{tr} can be obtained from the scattering rate via

$$\frac{1}{\tau_{tr}} = \sum_{\mathbf{k}'} W_{\mathbf{k} \rightarrow \mathbf{k}'} (1 - \cos \theta), \quad (2.18)$$

where $W_{\mathbf{k} \rightarrow \mathbf{k}'}$ is the scattering rate between the states $\mathbf{k}$ and $\mathbf{k}'$. The factor $(1 - \cos \theta)$ ensures that only momentum-relaxing events contribute to the transport. In order to evaluate equation 2.18 explicitly, the scattering rate $W_{\mathbf{k} \rightarrow \mathbf{k}'}$ must be determined for the leading scattering mechanism. Elastic scattering from ionized impurities like dopants N_A^- , N_D^+ dominate at low temperatures, and this is what often limits the mobility in a system the most given this model. The basic idea behind equation 2.18 comes from Fermi's golden rule which states that an electron with initial state $\mathbf{k}$ can be scattered into many possible final states $\mathbf{k}'$, where each possible transition has a probability per unit time to occur of

$$W_{\mathbf{k} \rightarrow \mathbf{k}'} = \frac{2\pi}{\hbar} \left| \langle \mathbf{k}' | U(\mathbf{r}) | \mathbf{k} \rangle \right|^2 \delta(E_{\mathbf{k}} - E_{\mathbf{k}'}) \quad (2.19)$$

where $U(\mathbf{r})$ is the impurity scattering potential, and $E_{\mathbf{k}}$ and $E_{\mathbf{k}'}$ are the electron energies of the initial and final quantum states with wave vectors $\mathbf{k}$ and $\mathbf{k}'$ respectively. The Dirac delta function, δ , makes sure that the energy is conserved during the scattering process by requiring the relationship $k' = k$ to be fulfilled. Since the impurities do not exchange energy with the electrons, the scattering is elastic and the initial and final electron energies are equal. To evaluate the matrix elements in 2.19, the impurity potential $U(\mathbf{r})$ must be rewritten in a form that can be used directly with plane-wave wavefunctions. Since electrons in a 2DEG are described by plane-wave states with wave vectors $\mathbf{k}$ and $\mathbf{k}'$, the relevant quantity in the scattering matrix element is not the real-space potential $U(\mathbf{r})$, but its Fourier transform $\tilde{U}(\mathbf{q})$. The momentum transfer during the scattering event is defined by the relationship

$$\mathbf{q} = \mathbf{k}' - \mathbf{k}, \quad q = |\mathbf{q}|. \quad (2.20)$$

and as the magnitude of momentum is conserved, and from the geometry in figure 2.9, it can be inferred that

$$q = 2k_F \sin \frac{\theta}{2} \quad (2.21)$$

Bringing these relationships into equation 2.18 yields the following expression

$$\frac{1}{\tau_{tr}} = n_{imp}^{(2D)} \frac{m}{\pi \hbar^3} \int_0^\pi \left| \tilde{U}(2k_F \sin \frac{\theta}{2}) \right|^2 (1 - \cos \theta) d\theta \quad (2.22)$$

where $n_{imp}^{(2D)}$ is the average impurity density per unit area. The next step is finding an expression for the potential $\tilde{U}(\mathbf{q})$ which describes how strongly an impurity scatters an electron by a momentum q . This follows directly from the two-dimensional Fourier transform of the real-space Coulomb potential $U(\mathbf{r})$. Consider a charged impurity located a distance d from the 2DEG, here the three-dimensional potential

will be

$$U(\mathbf{r}, z) = \frac{e^2}{4\pi\epsilon_0\epsilon_r} \frac{1}{\sqrt{r^2 + (z-d)^2}}. \quad (2.23)$$

Performing a 2D Fourier transform with respect to $\mathbf{r}$ results in

$$\tilde{U}(\mathbf{q}, z) = \frac{e^2}{2\epsilon_0\epsilon_r} \frac{e^{-\mathbf{q}|z-d|}}{q}, \quad (2.24)$$

which expresses how strongly the impurity scatters an electron by a momentum $\mathbf{q}$. However, equation 2.24 does not account for screening caused by polarization of the semiconductor dielectric. It is therefore necessary to modify equation 2.24 with the Thomas-Fermi factor with the result

$$\tilde{U}(\mathbf{q}, z) = \frac{e^2}{2\epsilon_0\epsilon_r q} \frac{e^{-\mathbf{q}|z-d|}}{(1 + \frac{q_{TF}}{q})} \quad (2.25)$$

where q_{TF} is the Thomas-Fermi screening wave number in 2D defined by

$$q_{TF} = \frac{m^* e^2}{2\pi\hbar^2 \epsilon_0 \epsilon_r} \quad (2.26)$$

Due to the electron wavefunction $\psi(z)$ being spread out in the z -direction, the potential $\tilde{U}(q)$ needs further adjustment. This is done by averaging the impurity potential over the probability density $|\psi(z)|^2$, which then accounts for how electrons experience the impurity field. From this, it is needed to define the form factor $F(q)$ to account for the averaged potential. For remote impurities distributed over a finite region, the form factor is calculated using a double integral:

$$F(q) = \int dz |\psi(z)|^2 \int dz' e^{-q|z-z'|} \quad (2.27)$$

which captures the fact that neither electrons nor impurities are located at a single point, but instead spread out in space. Due to this, the electron will feel an effectively weaker potential, since it averages the impurity potential across its complete wavefunction. Ultimately, the expression for $F(q)$ modifies equation 2.25 to an effective potential

$$\tilde{U}_{eff}(\mathbf{q}) = \tilde{U}(\mathbf{q})F(q). \quad (2.28)$$

Replacing $\tilde{U}(\mathbf{q})$ by $\tilde{U}_{eff}(\mathbf{q})$ in equation 2.22, allows the relaxation time τ -and thus the mobility-to be computed. As established above, the transport relaxation time depends highly on the screened impurity potential and its interplay with the electron wavefunction. The full expression has been evaluated explicitly for an electron wave function $\psi(z)$, which allowed for the effective potential $\tilde{U}_{eff}(\mathbf{q})$ to be computed. This provides a solid foundation for evaluating how impurity scattering limits momentum relaxation and hence the mobility in the system [14, 4].

2.3.2 Combined design considerations: balancing capacitance, mobility and robustness

Even though capacitance and mobility originate from different physical mechanisms, they are fundamentally linked through the way the gate voltage controls the electron distribution in the quantum well [1]. The gate capacitance determines how efficiently charge is induced in the 2DEG, while the mobility reflects how easily the confined electrons can move. Both quantities are determined by the spatial position of the wavefunction and the overall electrostatics of the device, which is why the device geometry and material parameters influence them simultaneously.

A lower capacitance reduces the electrostatic coupling between the gate and the channel, and therefore the amount of charge induced in the quantum well for a given gate voltage. In addition to this, it affects the shape and position of the wavefunction by altering the band bending. When the wavefunction is shifted deeper into the well and closer to the doped top barrier, the overlap with ionized donors increases. This enhances impurity scattering and reduces the mobility, so in this regime a reduction in capacitance directly leads to a reduction in mobility [23].

Moreover, device designs that aim to maximize mobility typically rely on pushing the wavefunction away from the dopants. To achieve this, a larger separation between the gate and the channel is required, which inevitably weakens the gate's ability to control the channel potential and therefore reduces the intrinsic gate capacitance. These two situations are not contradictory. Capacitance and mobility are both governed by the same gate-channel separation, but the trade-off plays out differently depending on whether the design prioritizes strong gate control or maximizing mobility [24].

Consequently, capacitance and mobility cannot be optimized simultaneously. They are coupled through the electrostatics, and the design goal is therefore to find a configuration where improvements in one quantity do not excessively degrade the other. On top of this, both quantities are sensitive to variations in parameters such as gate length, dielectric width and barrier composition [25]. A practical device must therefore not only achieve low parasitic capacitance and high mobility, but also maintain both quantities robust to small variations in geometry or material properties.

An optimization strategy is therefore required to identify a region in the design space where capacitance, mobility and robustness are simultaneously satisfied. Since these quantities depend non-linearly on geometry, material composition and electrostatics, this task is far from straightforward. A systematic optimization approach is therefore crucial to balance these trade-offs and detect regions with stable device characteristics.

Chapter 3

Method

3.1 COMSOL Multiphysics

3.1.1 Introduction to COMSOL Multiphysics

COMSOL Multiphysics was used as the primary simulation tool for this project. All simulations were performed using the Semiconductor Module in a two-dimensional geometry. The software solves the coupled Poisson and continuity equations using the finite-element method, which makes it suitable for modeling electrostatics, charge transport and capacitance in semiconductor devices. The following sections describe how the device geometry, material parameters, physics interfaces and boundary conditions were implemented. These are supposed to serve as a guide of how COMSOL works, and give tips that can help other people starting to work in COMSOL [26].

3.1.2 Model 1

This is a condensed version consisting of fewer steps on how to build the model in COMSOL. To get the extended version, please see A COMSOL guide on how to get started with COMSOL and set up the geometry

The goal was to build a model of a superconducting transistor in COMSOL, and it was crucial that the model was designed to closely resemble a real transistor, since this would impact the accuracy of the measurements. This included both semiconductor domains, as well as source, drain and gate-potential. The first step was to set up COMSOL correctly. For this model, the Semiconductor Module was used and a two-dimensional model was chosen. After establishing this, the model was built step by step using rectangles under the geometry node as well as geometry parameters. All geometries were parametrized to simplify and allow for changes in the model, and to account for any geometry criteria that may arise. The finished model can be seen in Fig. 3.1 below.

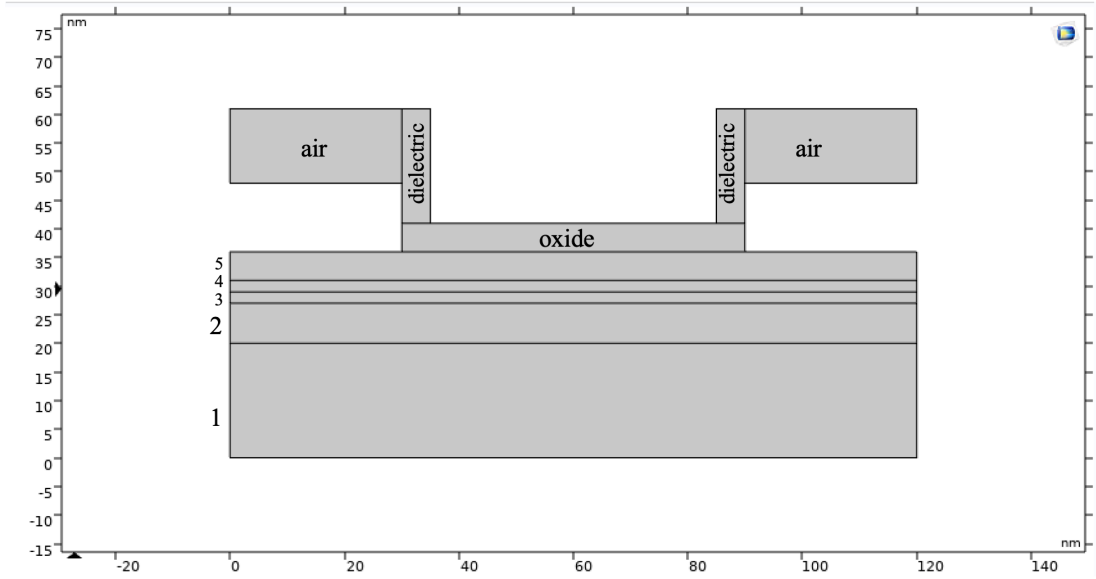


Figure 3.1: Overview of the finished model that was used in this thesis.

The semiconductor stack was implemented as five parametrized rectangular regions which are illustrated in Fig. 3.1, each corresponding to one layer of the heterostructure. The materials, doping concentrations and functional roles of these regions are summarized in Table 3.1.

Table 3.1: Summary of semiconductor and dielectric regions used in the COMSOL model.

Region	Material	Doping	Function
1	$\text{In}_{0.5}\text{Al}_{0.5}\text{As}$	10^{21} m^{-3} p-type	Buffer / lower barrier
2	$\text{In}_{0.7}\text{Ga}_{0.3}\text{As}$	10^{21} m^{-3} n-type	Lower channel
3	$\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$	n^+ , 10^{21} m^{-3}	Access region
4	$\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$	n^{++} , 10^{25} m^{-3}	Contact region
5	$\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$	Undoped	Spacer / confinement control
Oxide	SiO_2	–	Gate dielectric
Dielectric	SiO_2	–	Isolation
Air	Air	–	Surrounding domain

Since the goal was to model a superconducting transistor, the model had to reflect not only the geometry of a conventional transistor, but also capture the electrostatic effect of the superconducting contacts. Under the semiconductor physics node, both a Semiconductor Material Model and a Charge Conservation Model were included, which can be seen in Fig. 3.2.

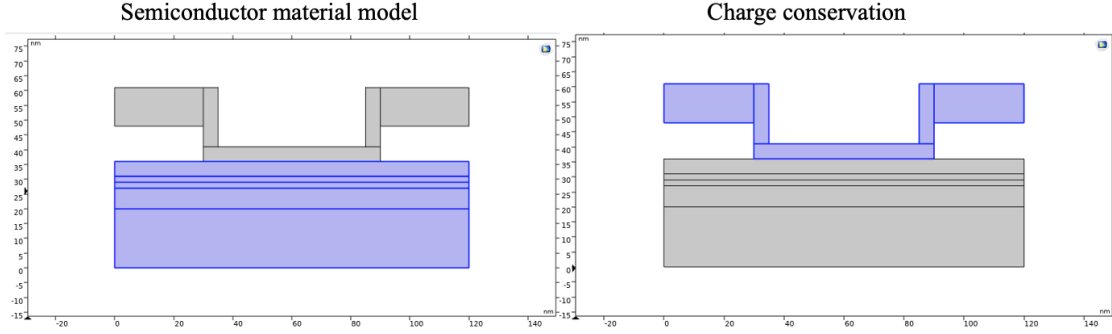


Figure 3.2: Domains assigned to the Semiconductor Material Model and Charge Conservation Model.

The Semiconductor Material Model was placed on all semiconductor domains, and it defined the electronic properties of the domains, such as carrier mobilities and doping levels. The Charge Conservation Model was used for the remaining, non-semiconductor, domains including the oxide, dielectric and air domains. These regions do not support mobile carriers, and therefore only the electrostatic Poisson equation is solved. This guarantees correct field distribution in insulating regions without introducing unnecessary transport equations.

Using both models was essential for simulating realistic device behavior, since the combination of the two ensured that both intrinsic material properties and charge transport are captured in a reliable way. The integration of both allowed the simulation to account for both steady-state operations, like current-voltage characteristics, and dynamic responses, such as the frequency dependence of the capacitance.

To implement the electrical contacts, aluminum was added to selected boundaries. In the model, the aluminum contacts were treated as ideal superconducting boundaries, meaning they were set to a fixed potential and assumed to have negligible resistance. The boundaries were chosen to mimic the placement of metal contacts in an actual superconducting device, and the boundaries corresponding to the source and drain terminals were assigned aluminum, as shown in Fig. 3.3.

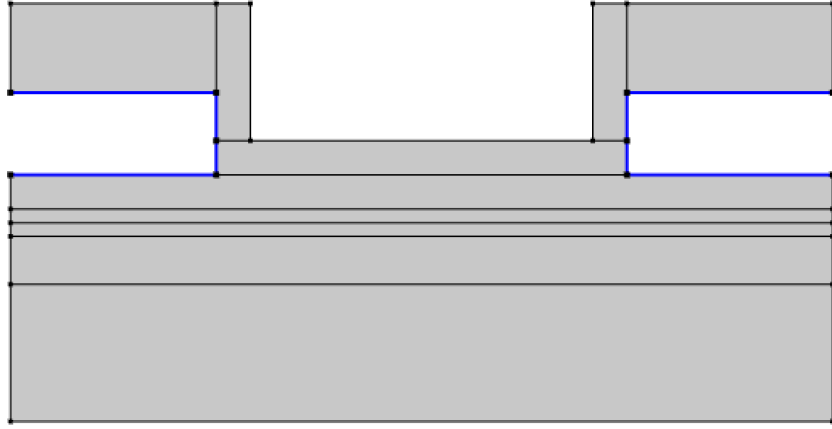


Figure 3.3: Aluminum boundaries defining the source and drain contacts.

Additionally, aluminum was added to three boundaries to define the gate, visualized in Fig. 3.4.

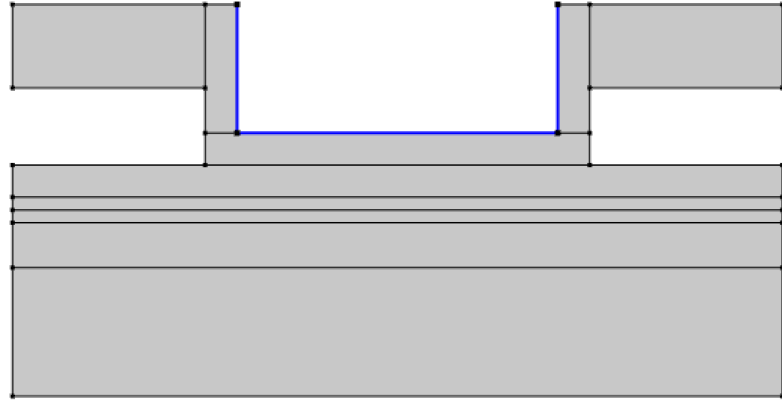


Figure 3.4: Aluminum assigned to the boundaries defining the gate terminal.

Once all materials had been defined in the model, the next step was to implement the electrical contacts of the device. COMSOL provides two types of electrical contacts: metal contacts (ohmic or Schottky) and terminals. In this model, both ohmic contacts and terminals were used. Ohmic contacts were applied to boundaries between semiconductor domains, while terminals were placed on the boundaries of conductive domains, such as metals.

For the source contact, an ohmic boundary condition was assigned on the semiconductor side, and a

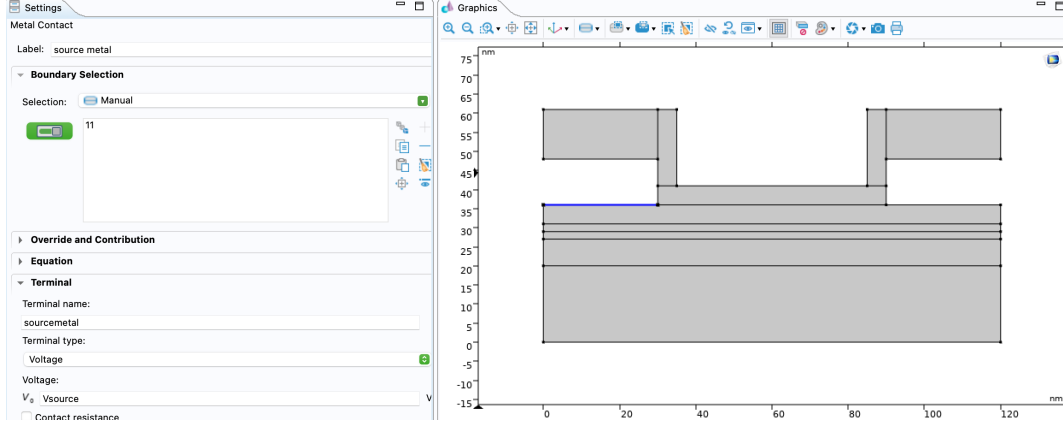


Figure 3.5: Ohmic metal contact with voltage V_{source} for the source.

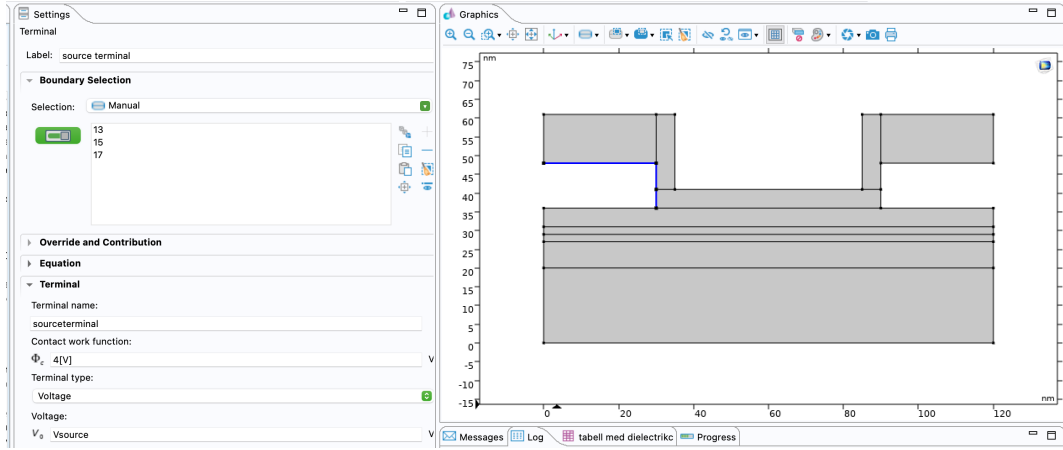


Figure 3.6: The source terminal with voltage V_{source} .

terminal was placed on the corresponding aluminum boundary. The voltage of the source contact was set to V_{source} , which was varied depending on the simulation settings. Its implementation is shown in Fig. 3.5 and Fig. 3.6. Similarly, the drain region required both an ohmic contact and a terminal. The drain contacts are shown in Fig. 3.7 and Fig. 3.8, and the voltage of the drain contact was V_{drain} , which was also varied.

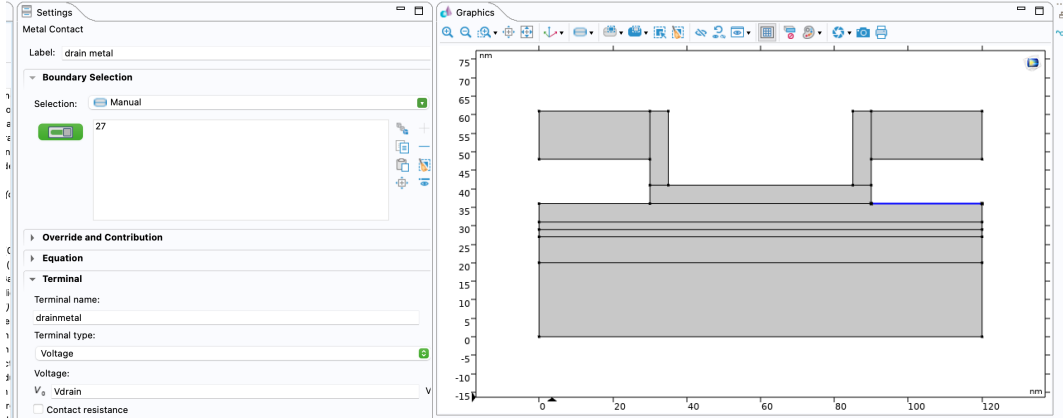


Figure 3.7: Ohmic metal contact with voltage V_{drain} for the drain.

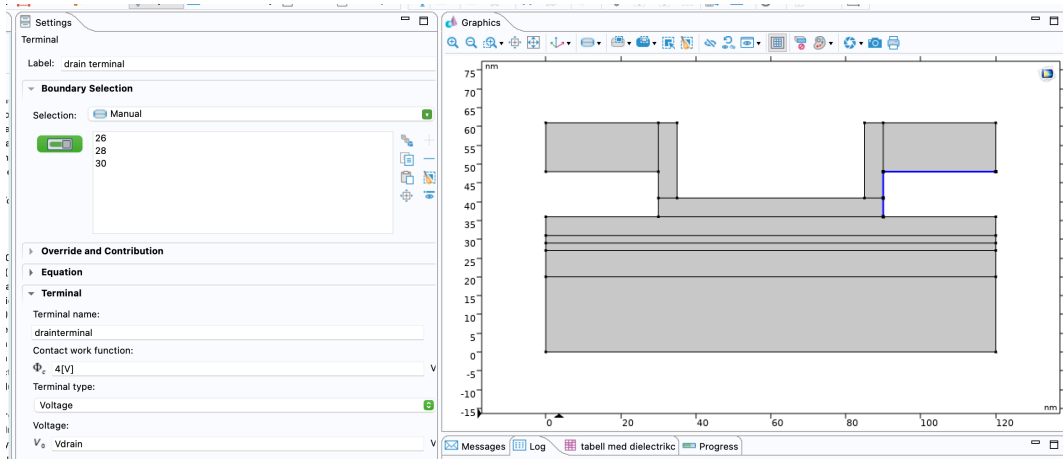


Figure 3.8: The drain terminal with voltage V_{drain} .

The gate contact was applied solely to aluminum boundaries, and therefore required only a terminal at the voltage $V_{gatearbete}$, as shown in Fig. 3.9. The voltage $V_{gatearbete}$ corresponded to the transistor's operating point.

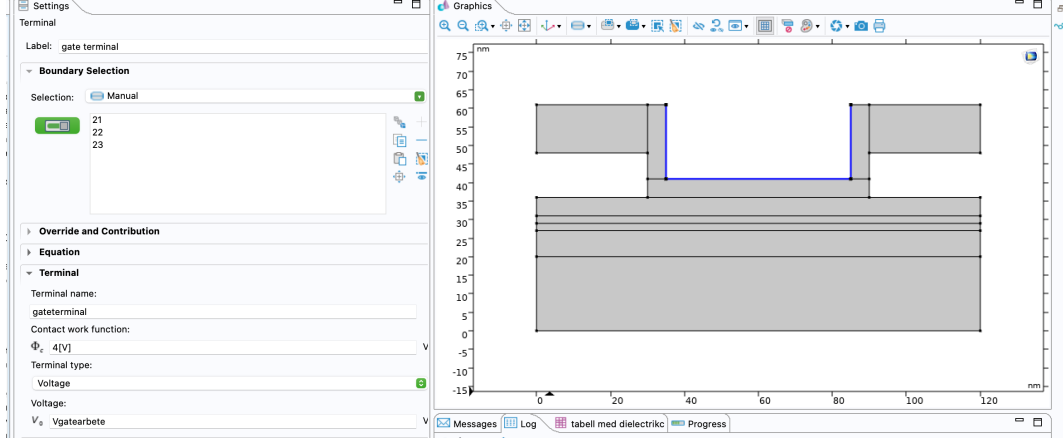


Figure 3.9: The gate terminal with voltage $V_{gatearbete}$.

To simulate the transistor's response to an alternating current (AC) signal, the harmonic perturbation shown in Fig. 3.10, with an amplitude of 0.001 V, was applied to the gate contact. This perturbation represented a small sinusoidal perturbation which was added on top of the DC gate bias. This allowed the model to account for the device's dynamic behavior under high-frequency excitation. A voltage amplitude of 0.001 V was small enough to force linear response, while still capturing the frequency-dependent capacitance.

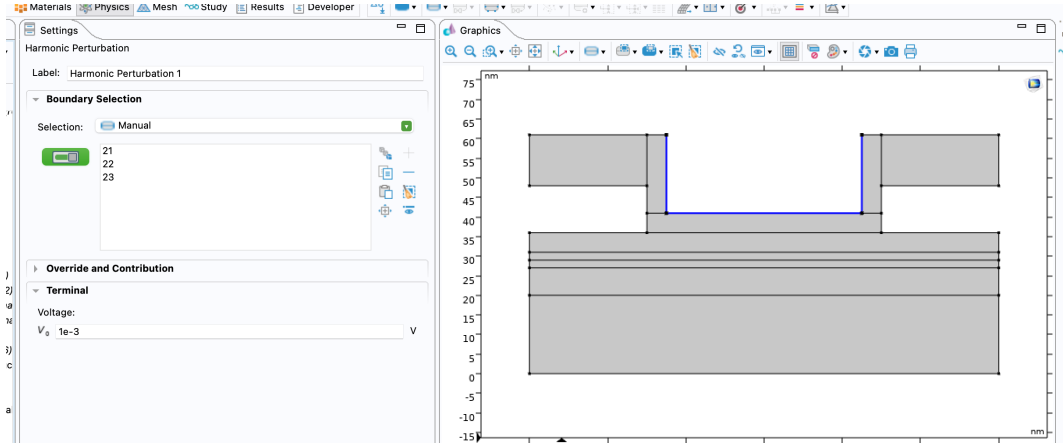


Figure 3.10: A harmonic perturbation placed on the gate terminal.

Model 1 therefore provided the device-level electrostatic response of the full transistor structure, including the capacitance and potential distribution under realistic boundary conditions. These results were later combined with the mobility obtained from Model 2 presented below, which ensured that the geometry optimization was restricted by the transport properties of the quantum well.

3.1.3 Model 2

While Model 1 captured the full transistor geometry and its electrostatic behavior, Model 2 was developed to isolate the quantum well region and resolve the quantum-mechanical confinement with higher

accuracy. The second model focused exclusively on the semiconductor layers where quantum confinement is essential. By reducing the geometry and moving from a 2D to a 1D representation, numerical accuracy was increased while computational cost was reduced. This model could be used to solve the Schrödinger–Poisson system self-consistently, which in turn yielded information on subband energies, wavefunctions and charge distribution with high accuracy. These quantities form the foundation for the mobility calculations and make it possible to study how different geometric choices influence the transport properties in a clear and controlled way.

The second model was also created using COMSOL, this time as a 1D model consisting of six line segments denoted A-F, as shown in Fig. 3.11. Each segment represented a specific part of the transistor. Segment A, with a length of 5 nm, was silicon dioxide (SiO_2). This layer electrically isolates the semiconductor from the gate while setting the boundary conditions for the electrostatic potential in the quantum well.

To the right of the oxide, between 5 and 10 nm, Segment B consisted of an undoped spacer layer of $\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$, which separated the quantum well from the doped region. The purpose of the spacer layer was to reduce ionized impurity scattering and ensure a well-defined confinement potential. Segment C, a 10 nm layer of $\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$, formed the left barrier of the well region, setting the band offset and defining the confinement on the gate side. This barrier was n -doped with a donor concentration of 10^{25} m^{-3} , providing the electron supply for the quantum well.

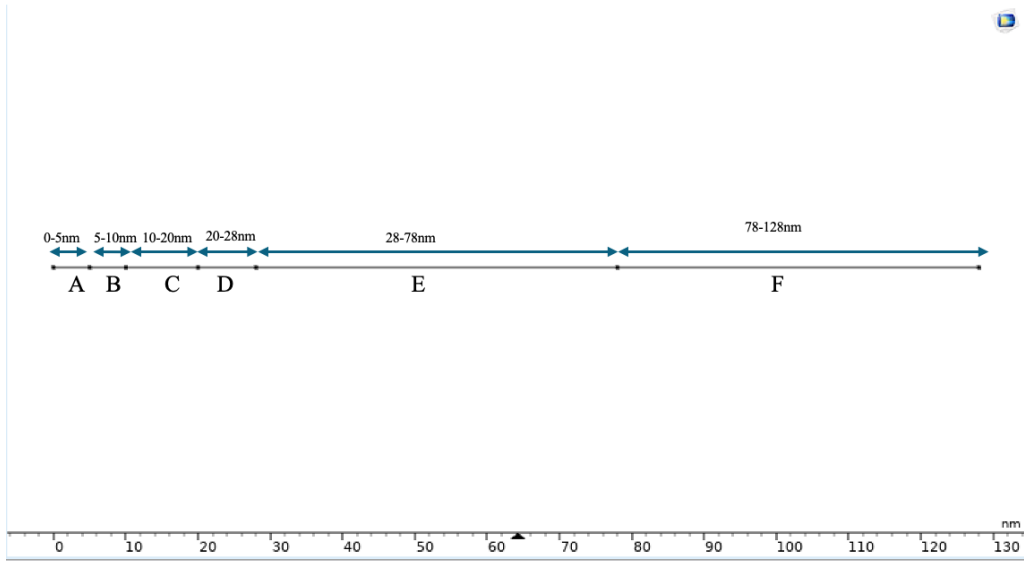


Figure 3.11: An overview of the second model consisting of six line segments denoted A-F, each representing different parts of the transistor.

The quantum well, represented by segment D, was composed of an 8 nm layer of $\text{In}_{0.7}\text{Ga}_{0.3}\text{As}$ that was n -doped to 10^{21} m^{-3} , providing the electron population in the confined region. Segments E and F formed the right barrier, each a 50 nm layer of $\text{In}_{0.5}\text{Al}_{0.5}\text{As}$. Segment E was left undoped to maintain a well-defined confinement potential, whereas Segment F was p -doped to 10^{21} m^{-3} to supply the fixed acceptor charge on the drain side of the structure [1].

For this model, four different physics interfaces were used in COMSOL: Schrödinger Equation, Semiconductor, Schrödinger–Poisson Coupling, and Electrostatics. The Electrostatics and Semiconductor interfaces were used to solve the Poisson equation based on the specified doping profiles and fixed charges. The Schrödinger Equation was used to obtain the eigenenergies and corresponding bound-state wavefunctions, which could then be plotted for analysis. Lastly, the Schrödinger–Poisson coupling node used the user-defined inputs and created the variables and equations necessary for the bidirectional coupling between the Electrostatics and Schrödinger Equation interfaces, which made it possible to model the behavior of charged carriers inside the confined region [27].

To complete the model, a set of variables and operators were defined to supply the Schrödinger and Poisson solvers with the necessary physics. Different effective masses were assigned to the well and barrier domains to account for the underlying material-dependent band structure. Since the effective mass of the charge carriers depends on the curvature of the conduction band, each material has its own characteristic value. It was important to use domain-specific effective masses, since this ensured that the Schrödinger equation correctly described carrier confinement in the well and tunneling through the barriers. In the quantum-well domain (Segment D), the effective mass m^* was therefore set to

$$m^* = m_{well} m_0 \quad (3.1)$$

where m_{well} is a material-specific parameter and m_0 is the free-electron mass. The conduction-band edge was fixed to (chosen as reference energy)

$$E_c = 0 \quad (3.2)$$

In the barrier domains C, E, and F the corresponding variables were defined as

$$m^* = m_{barrier} m_0 \quad (3.3)$$

and

$$E_c = \Delta E_{c,b} \quad (3.4)$$

where $m_{barrier}$ is a material-specific parameter that differs between the barrier domains, and $\Delta E_{c,b}$ is the conduction-band offset of the barrier material, ensuring that each region used the correct effective mass and band offset [27]. Auxiliary global variables were defined to compute normalized wavefunctions and charge densities. These included

$$\psi_{\text{norm}} = \Psi / \max(|\Psi|) \quad (3.5)$$

where $\int |\Psi|^2 = 1$ and ψ_{norm} is a visualization-normalized wavefunction obtained by dividing ψ by its maximum magnitude so that its peak value becomes unity. Additionally, the normalized probability density was defined as $\rho_{\text{norm}} = |\Psi|^2 / \int |\Psi|^2$, where ρ_{norm} is the physically normalized probability density whose integral equals 1, as required by quantum mechanics.

To substantiate these expressions, two integration operators, `intPsi` and `intopwell`, and one maximum operator `maxop1` were introduced. `intPsi` was used to compute the total probability over the barrier and confined region, which was then used to evaluate the normalized probability density ρ_{norm} . The operator `intopwell` restricted the integration to the quantum well domain only, enabling evaluation of the probability of locating the electron inside the well. `maxop1(abs(schr.Psi))` was used to find

the maximum value of the wavefunction magnitude over the quantum domains (C–F). This value was subsequently used in ψ_{norm} to normalize the amplitude of the wavefunction for clearer visualization and plotting [27, 1].

To connect the quantum-mechanical model to the transport analysis, the output from Model 2, specifically the subband energies and wavefunctions, was exported to MATLAB. Two MATLAB scripts were developed to compute the mobility based on these quantities. The code used the screened Coulomb potential $\tilde{U}(q)$ obtained from the Thomas-Fermi expression, together with the form factor $F(q)$ derived from the COMSOL-computed wavefunction $\psi(z)$, to construct the effective impurity potential

$$\tilde{U}_{\text{eff}}(q) = \tilde{U}(q) F(q). \quad (3.6)$$

This effective potential was then inserted into the angular integral for the transport relaxation time τ , from which the mobility follows directly. Model 2 therefore provided a quantum-mechanically consistent description of the confined electrons and their energies, enabling an accurate evaluation of the mobility and corresponding conductivity. These quantities acted as constraints in the subsequent geometry optimization performed for Model 1, ensuring that the capacitance is minimized only for the geometries that remain compatible with the transport properties of the quantum well.

Although Model 2 captured the dominant confinement physics, it relied on standard assumptions such as one-dimensional potential variation, material-dependent effective masses, and the exclusion of several scattering mechanisms. The mobility model included only ionized impurity scattering, which typically dominates in doped quantum wells at low temperatures, whereas other scattering processes such as phonon scattering and alloy disorder were neglected [4, 14]. As a result, the computed mobility should be interpreted as an approximation limited by the included scattering mechanism, but it remains accurate enough for the comparative and trend-based analysis carried out in this work.

To conclude, the workflow used in this thesis combines two complementary COMSOL models that are used in conjunction with post-processing in MATLAB. Model 1 represented the full transistor geometry, and was used to evaluate physical quantities such as capacitance and voltage. Model 2 focused on the quantum well region and solved the Schrödinger–Poisson system self-consistently, providing information about quantum-mechanical quantities with high accuracy. These quantities were exported to MATLAB, where the mobility was computed using a screened impurity-scattering model. The computed mobility and conductivity served the role of physical constraints to optimize the geometry for Model 1.

Combined, this workflow enabled a coherent connection between quantum-mechanical confinement, electrostatics and transport properties, establishing a consistent framework for the analysis presented in the following sections. However, since the reduced 1D model does not directly provide a gate-dependent sheet density, the effective gate-dependent sheet density in the quantum well $n_{2D}(V_{gate})$ was extracted from the normalized probability density $\psi(x)^2$, which created the link between the transport model and the quantum-mechanical description.

The quantum-mechanical quantities obtained from Model 2 therefore form the foundation for the transport analysis. To connect these results to the electrostatic behavior of the full device, the next section describes how COMSOL and MATLAB were combined to extract the gate-dependent mobility and capacitance.

3.1.4 MATLAB and COMSOL together: extracting mobility as a function of capacitance

To link the quantum-mechanical results from Model 2 with the electrostatic behavior obtained from Model 1, a combined COMSOL-MATLAB workflow was developed. The purpose of this workflow was to determine how the electron mobility varies with the device capacitance and thereby establish a direct connection between transport and electrostatic properties.

The process began by sweeping V_{gate} in Model 1 between -0.25 V and 0.25 V in steps of 0.05 V. The cross section was taken vertically through the center of the gate, along the growth direction z , that is, normal to the semiconductor layers shown in Fig.3.1. For each of the eleven gate-bias points, the electrostatic potential $V(z)$ was calculated along the growth direction z and exported from COMSOL. These potential profiles were used as input to the Schrödinger Equation interface in Model 2, where they were used to determine the bound-state energies and wavefunctions. Consequently, each of the gate voltages produced a unique confinement potential and therefore a ground-state wavefunction $\psi(z)$ with a corresponding eigenenergy.

For every gate-bias point, the computed wavefunction was exported from COMSOL. This text file contained the spatial coordinate x and the corresponding wavefunction amplitude along the confinement direction, yielding eleven files of the form **z vs.psi**. Each of these files hence represented a state of the well at a specific gate voltage. These files were then imported into MATLAB with their associated eigenenergies.

A MATLAB script (Appendix B) was developed to process all wave-function files and compute the mobility as a function of gate voltage. The script

- imported the gate-dependent wavefunctions $\psi(z)$ and eigenenergies from COMSOL,
- computed the probability density $|\psi(z)|^2$ and used it to evaluate the form factor $F(q)$ which accounts for the finite width of the 2DEG,
- constructed the screened Coulomb potential using the Fourier-transformed impurity potential together with Thomas–Fermi screening,

$$\tilde{U}(q) = \frac{e^2}{2\varepsilon_0\varepsilon_r q(1 + q_{TF}/q)}$$

- combined the screened potential with the form factor to obtain the effective impurity potential

$$\tilde{U}_{\text{eff}}(q) = \tilde{U}(q) F(q)$$

- evaluated the transport relaxation time by inserting $\tilde{U}_{\text{eff}}(q)$ into the angular scattering integral

$$\frac{1}{\tau} = n_{\text{imp}}^{(2D)} \frac{m^*}{\pi\hbar^3} \int_0^\pi \left| \tilde{U}_{\text{eff}}\left(2k_F \sin \frac{\theta}{2}\right) \right|^2 (1 - \cos \theta) d\theta$$

- and finally computed the mobility from

$$\mu(V_{\text{gate}}) = \frac{e\tau}{m^*}$$

This workflow produced a consistent set of gate-dependent mobilities and capacitances, forming the basis for the subsequent analysis of the connection between mobility and capacitance.

3.2 How COMSOL finds solutions

COMSOL Multiphysics is a finite-element simulation platform that solves physics problems. This is done by converting the equations involved, such as the Schrödinger equation or Poisson's equation, to a numerical form that can be solved by a computer. COMSOL does this by dividing the geometry of the constructed model into many small elements. The physical fields of each element are then approximated with simple basis functions, for example piecewise linear or piecewise quadratic. This transforms continuous differential equations into a huge system of algebraic equations that can be solved numerically.

Each physics interface in COMSOL was represented by a set of governing equations. The semiconductor interface uses Poisson's equation among others, while the Schrödinger interface solves the eigenvalue problem for quantized states. COMSOL took all the different parts of the model, such as geometry and material parameters and boundary conditions, and created a model that accounts for all this. To solve the equations for a given model, the solver can use different approaches. There are a number of "studies" in COMSOL which one can add to solve its problem, these include stationary solutions, time-dependent solutions and frequency-perturbed solutions [27].

An extension to this, which is commonly used for more complicated problem, is to use the multiphysics part of COMSOL. The Schrödinger–Poisson coupling is an example of this, where COMSOL uses an iterative scheme. One equation is solved using the current estimate of the other, and the results are sent in between the two physics interfaces until the solution stabilizes around a value. Consequently, the result was a correct solution for all equations [28].

With this as a background, the way COMSOL finds solutions can be summarized as

- representing the geometry with a mesh consisting of a finite number of elements. COMSOL can't calculate at every infinitely small point, so it divides the structure into many small pieces.
- translating the governing equations into a numerical system
- using the material parameters and boundary conditions
- being able to solve complex systems using numerical solvers
- iterating to find a solution when multiple physics are coupled

This method makes COMSOL a leading software for handling complex semiconductor structures, for instance quantum wells and solar cells, and for obtaining reliable solutions to challenging problems [29, 30].

Chapter 4

Results

This section presents the simulation results obtained from the two COMSOL models described in the method chapter. The goal with these simulations is to quantify how variations in the device geometry and selected material parameters influence the total gate capacitance of the superconducting QWFET. To do this, it is vital to separate the intrinsic and parasitic contributions to the gate capacitance, and to evaluate how the capacitance components scale with parameters such as gate length and gate voltage.

The theoretical framework from the theory section serves as the background for interpreting these results. The following subsections will present the results for simulations on Model 1 and Model 2, which were done using a combination of COMSOL and MATLAB.

4.1 Extraction of the threshold voltage

Before analyzing the capacitive behavior of the device, it was necessary to determine an appropriate operating point for the simulations. The drain and source voltages, V_{drain} and V_{source} were set to 0V when performing the capacitance simulations, but it was needed to find a reasonable gate voltage, V_{gate} for the simulations. The gate bias used throughout this work was chosen to be

$$V_{gate} = V_{th} + 0.1 V \quad (4.1)$$

where V_{th} is the threshold voltage of the transistor. To extract V_{th} , a gate-voltage sweep was performed in COMSOL with $V_{gate} = \text{range}(-0.5, 0.01, 0.5) V$. For this sweep, the drain voltage was fixed at 50 mV and the source voltage at 0 V. The resulting drain current I_D was first plotted as a function of V_{gate} . To identify the threshold region more clearly, the square root of the drain current, $\sqrt{I_D}$, was then plotted as a function of V_{gate} , which can be seen in Fig. 4.1.

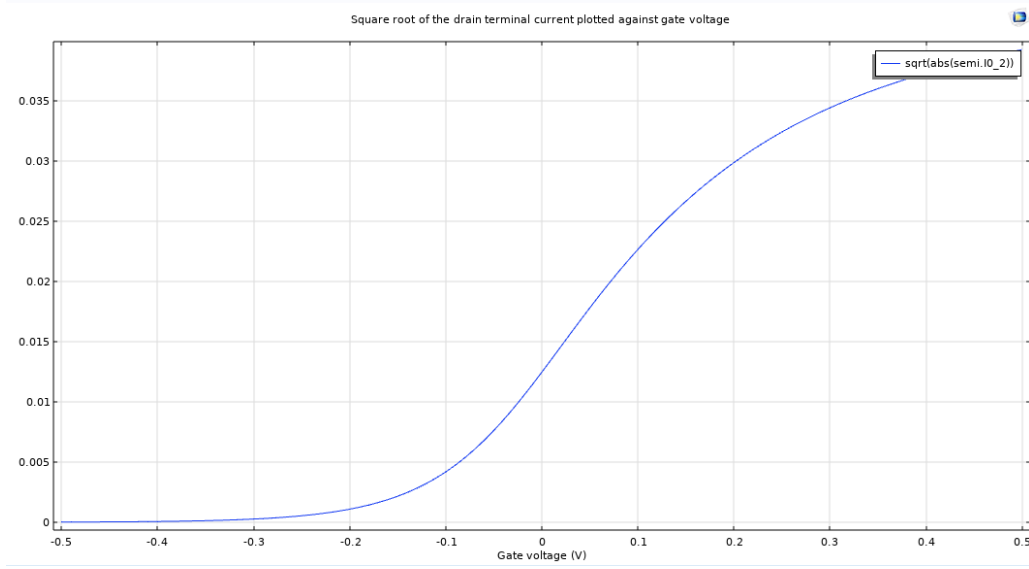


Figure 4.1: Square root of the drain current, $\sqrt{I_D}$, plotted against V_{gate} with the purpose of helping to determine the threshold voltage, V_{th} .

Although the square-root transformation originates from long-channel MOSFET theory, where the saturation current follows a quadratic dependence on $(V_{gate} - V_{th})$, the present simulations are performed at a low drain bias of $V_{drain} = 50$ mV. At such a small drain voltage the device operates between the linear and saturation regimes, resulting in a $\sqrt{I_D}$ - V_{gate} characteristic that is not perfectly linear. Nevertheless, the curve still exhibits a sufficiently linear interval to allow an approximate extraction of the threshold voltage.

A linear region was identified in the $\sqrt{I_D}$ - V_{gate} curve in Fig. 4.1, and a linear fit was applied to the chosen interval of $V_{gate} \in [0.05, 0.17]$ V. This interval was chosen because it exhibited the most linear behavior while avoiding numerical noise at low current levels, and its linear behavior in this region can be seen in Fig. 4.2.

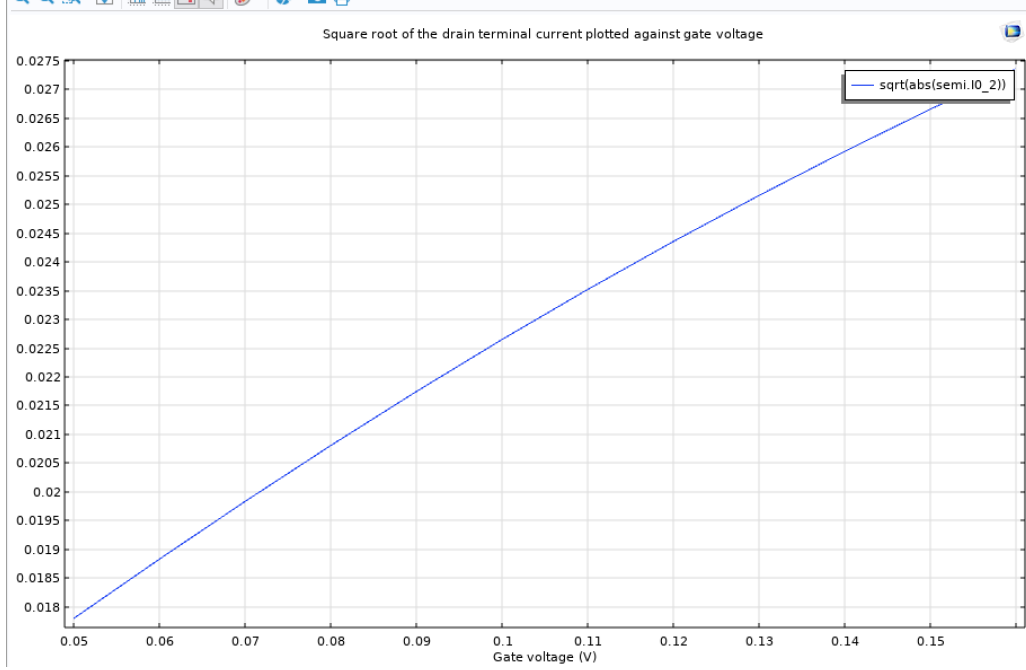


Figure 4.2: Linear region of the $\sqrt{I_D}$ - V_{gate} characteristic used for threshold voltage extraction. The highlighted interval shows the portion of the curve exhibiting approximately linear behavior and used for the linear fit.

The extracted fit exhibited a high degree of linearity, where the coefficient of determination R^2 was calculated to be as high as 0.9972. The fitted line was extrapolated to $\sqrt{I_D} = 0$, which in turn yielded a threshold voltage of $V_{th} = -0.18$ V. This threshold voltage is consistent with expectations for this kind of device structure, as discussed in standard MOSFET threshold-voltage theory [1]. Even though the exact value of V_{th} may shift with material parameters and doping levels are varied, this serves as an accurate enough estimation to base the following measurements on. The operating point for V_{gate} was defined as

$$V_{gate} = V_{th} + 0.1 \text{ V} \approx -0.08 \text{ V} \quad (4.2)$$

This offset ensures that the device operates above threshold while avoiding a high electron density in the channel. Hence, many of the subsequent capacitance measurements were performed at a gate bias of $V_{gate} = -0.08$ V using a harmonic perturbation of 1 mV applied to the gate, while both the source and drain terminals were kept at 0 V.

4.2 Results for Model 1

Model 1 provides the full two-dimensional representation of the transistor cross section, allowing the electrostatic response of the full device geometry to be evaluated. The purpose of this model is to quantify how the total gate capacitance changes with four key parameters. Three of these are geometrical: the gate length L_{gate} , the width of the surrounding dielectric regions $W_{dielectric}$ and the oxide thickness t_{ox} which are illustrated in Fig. 4.3. The fourth parameter investigated was the dielectric permittivity, ϵ_r

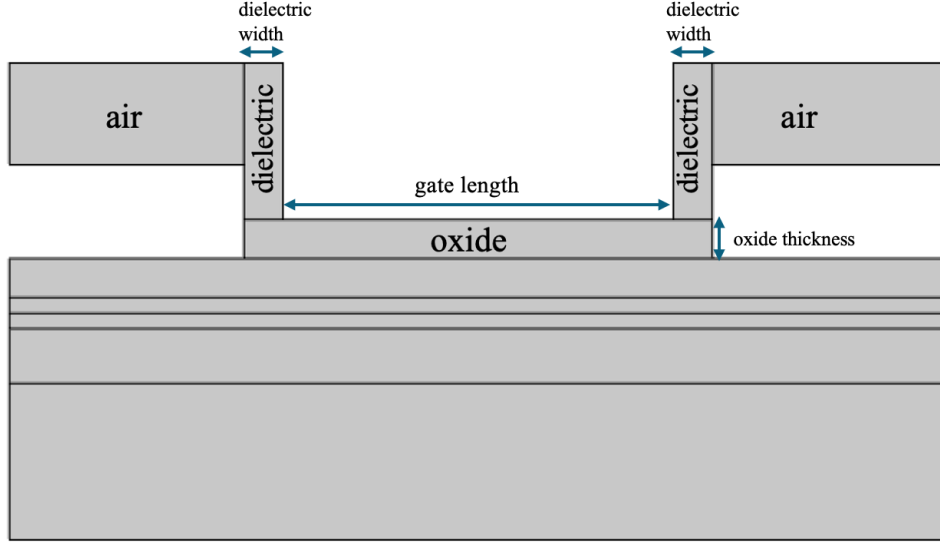


Figure 4.3: Illustration of the placement of the geometrical parameters swept in Model 1 to evaluate their influence on the total gate capacitance.

of the dielectric layers. To make the analysis clear, the influence of each parameter on the total gate capacitance is first examined individually before considering their combined effect.

4.2.1 Gate length influence

From Fig. 4.3 it is clear that the gate length is defined by the oxide length L_{ox} and the dielectric width through the geometric relation

$$L_{oxide} = L_{gate} + 2 W_{dielectric}$$

This follows directly from the 2D geometry, where the oxide spans the gate region and the two adjacent dielectric sections. To investigate the influence of the gate length more thoroughly, a parametric sweep over the gate length was performed for six different gate voltages: -0.28, -0.18, -0.08, 0.02, 0.12 and 0.22 V. These voltages were chosen in steps of 0.1 V around the calculated operating point of -0.08 V. For each of these gate biases, the gate length was swept from 2 to 175 nm, and the resulting capacitances can be seen in Fig. 4.4

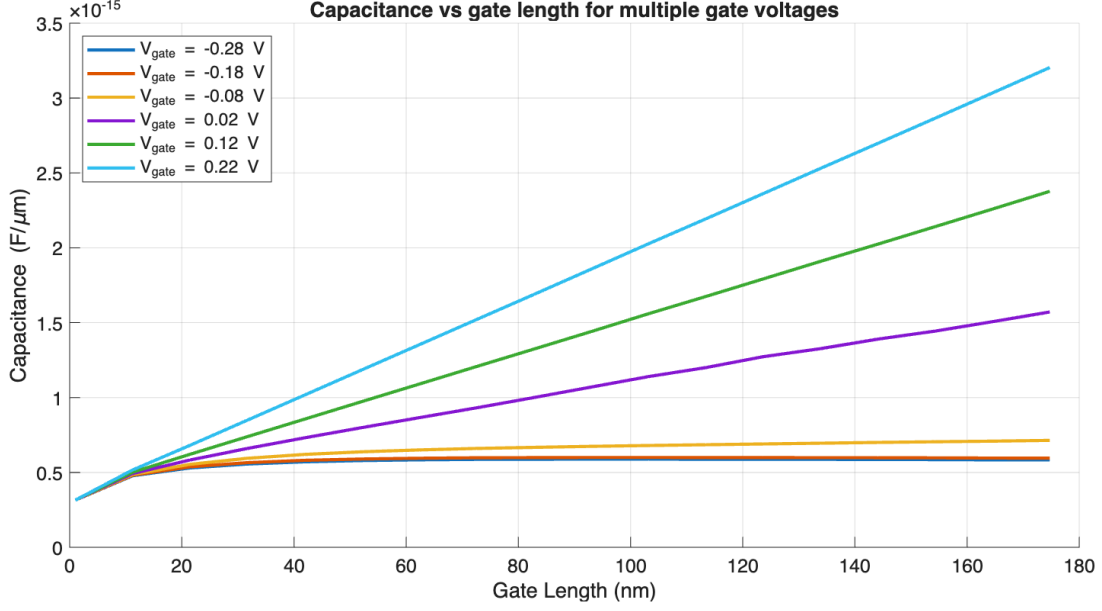


Figure 4.4: Capacitance as a function of gate length and gate voltage.

The capacitance visualized in Fig. 4.4 corresponds to the gate capacitance that was extracted from the small-signal AC analysis, where a harmonic perturbation is applied to the gate voltage and the resulting charge response is measured using COMSOL. Therefore this quantity includes both the intrinsic contribution from charge modulation in the quantum well, and the parasitic components created by fringing fields in the surrounding dielectric regions.

From the figure it is clear that the capacitance increases almost linearly with gate length. However, for very short gate lengths (below about 10 nm), the capacitance increases more rapidly due to short-channel effects, where fringing fields dominate and the gate loses electrostatic control of the channel. For longer channel lengths, the near-linear behavior follows from the capacitance definition $C = \partial Q / \partial V_{gate}$ introduced in equation 2.15 (although equation 2.15 defines the intrinsic part, the same relation applies to the total gate capacitance extracted in Model 1), and is a confirmation that the model behaves in accordance with its expected electrostatic behavior. When the gate voltage increases, more electrons are pulled into the quantum well, resulting in stronger electron accumulation. Furthermore, this increases the charge density per unit length, so an increase in gate length adds a larger amount of charge at high voltages than it does at low voltages. Consequently, the capacitance grows more rapidly with increasing L_{gate} , yielding a steeper $C - L_{gate}$ relation at high gate bias. At low gate bias, the channel contains only a small electron population, so changes in gate length have a much smaller effect on the total charge, and the capacitance therefore varies only weakly with L_{gate} . Around the operating point ($V_{gate} = -0.08$ V), the slope is moderate, which tells us that the device is close to the onset of significant electron accumulation in the quantum well.

4.2.2 Dielectric width influence

To examine how the width of the two dielectrics influences the capacitance, a set of three simulations was performed. The dielectric width was set to 3, 4 and 5 nm while the gate length was swept between 2 and 175 nm. As previously established, a longer gate length yields a higher capacitance, but another important trend can be seen in Fig. 4.5.

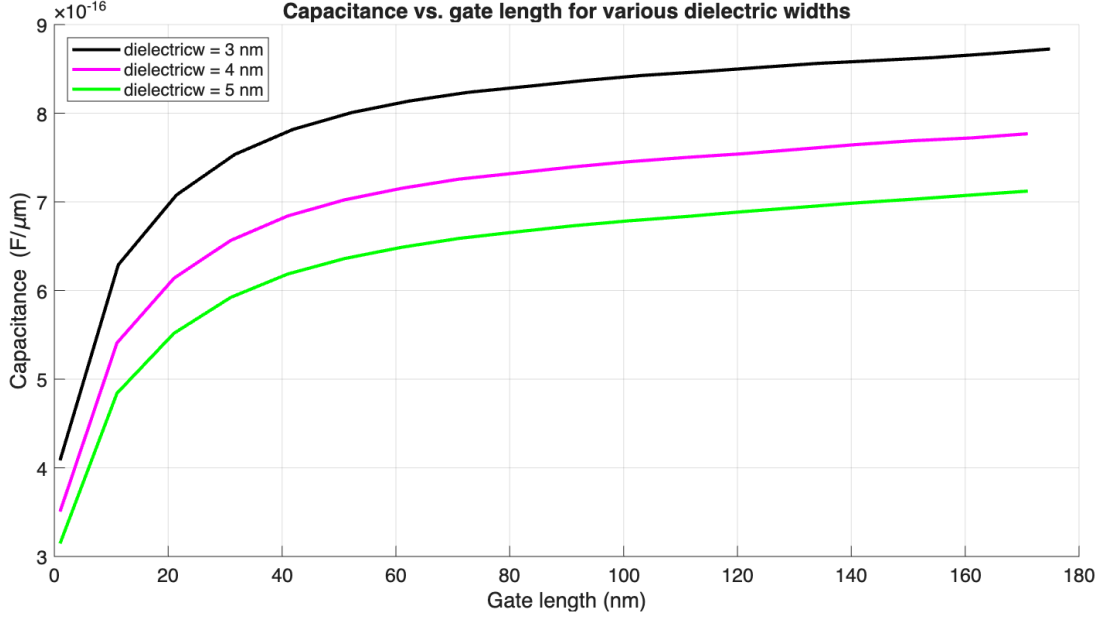


Figure 4.5: Capacitance varying with gate length and dielectric width for three different widths of the dielectrics.

It is clear that the thinner the dielectric, the higher the capacitance. This is because reducing the lateral separation strengthens the electrostatic coupling between the gate and the channel. A smaller dielectric width shortens the path taken by the fringing fields, which allows the electric field from the gate to reach the channel more effectively. This means that more charge is induced in the channel per unit gate voltage, thereby increasing the effective gate capacitance. However, the nearly parallel curves indicate that the dielectric width primarily shifts the overall capacitance level, while the dependence on gate length remains largely unchanged. This is consistent with the fringing-field behavior discussed in the theory, which highlights the importance of lateral field spreading in the 2D model.

4.2.3 Oxide thickness influence

The capacitive influence of the oxide thickness, t_{ox} , was evaluated for the same gate length-range as the previous studies, and the resulting capacitance curves are presented in Fig. 4.6.

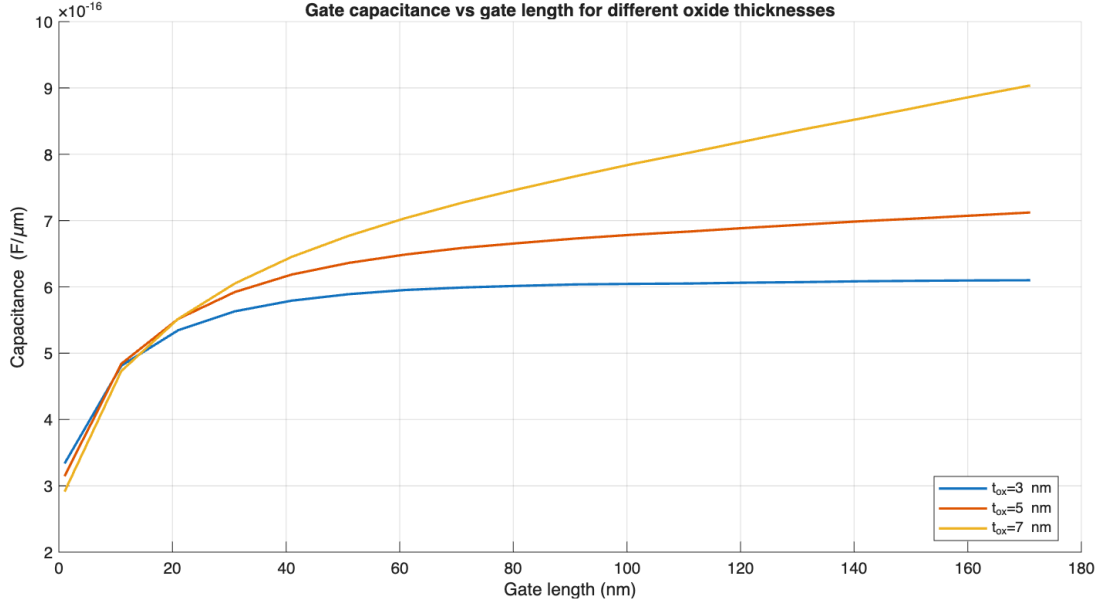


Figure 4.6: Capacitance vs gate length for three different oxide thicknesses: 3, 5 and 7 nm.

The figure shows that the capacitance increases with gate length for all oxide thicknesses, which is consistent with the geometric scaling observed earlier. However, it is clear that a thicker oxide consistently yields a higher capacitance. This behavior is the opposite of what is expected from an ideal MOSFET, where a thicker oxide layer usually decreases the capacitance due to the proportionality relation $C \propto 1/t_{ox}$. For the present 2D geometry, the behavior of the capacitance differs from the ideal case. The extracted capacitance is dominated by lateral fringing fields, and a thicker oxide provides a larger dielectric volume for these fields to spread into, thereby enhancing the capacitive coupling. Due to this, the total capacitance increases with the oxide thickness for this model, despite deviating from the ideal MOSFET trend.

However, it should be noted that another factor that may contribute to this behavior is the chosen operating point for the gate voltage, which was set to -0.08 V. This operating point was originally calculated for the geometry with an oxide thickness of 5 nm and further analysis shows that it shifts when the oxide thickness is changed. Therefore, the corresponding gate voltages were recalculated for the geometries with 3 nm and 5 nm oxide thickness, and the capacitance was reevaluated. Despite this adjustment, the same trend in how the capacitance varies with oxide thickness remained, and only minor changes in capacitance were observed when changing the operating point.

4.2.4 Dielectric permittivity influence

There is a direct proportionality relation between the capacitance and the dielectric permittivity, $C \propto \epsilon_r$, so as the permittivity of the dielectrics increase, so does the capacitance. As shown in Fig. 4.7, the three curves remain nearly parallel across the entire gate-length range, indicating that the dielectric permittivity primarily shifts the capacitance upward without altering how the capacitance scales with L_{gate} .

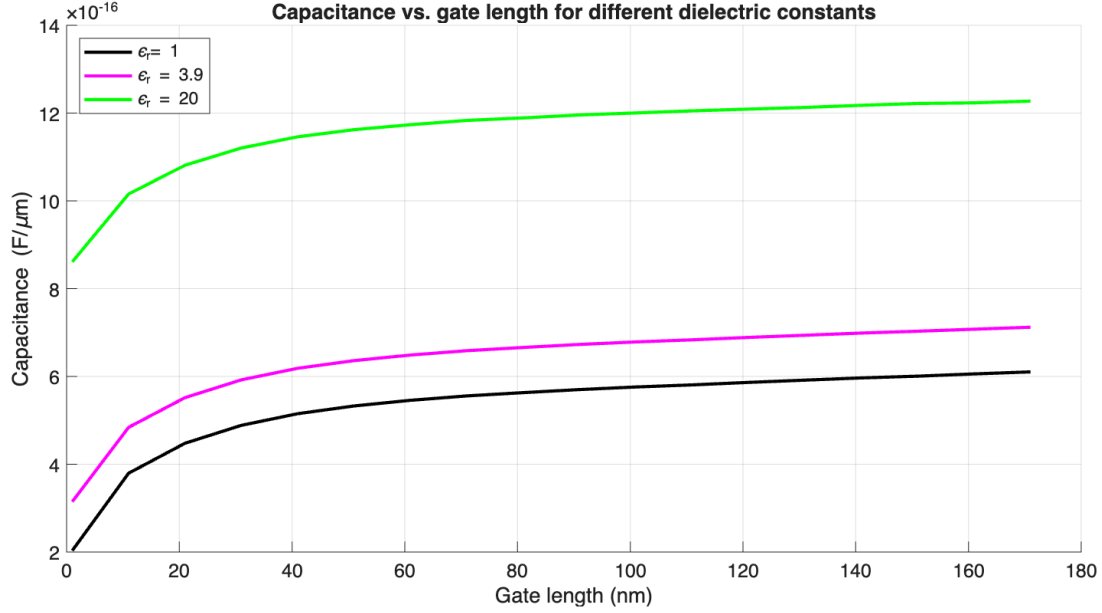


Figure 4.7: Capacitance varying with gate length and dielectric permittivity for three different values of the dielectric permittivity: 1, 3.9 and 20.

This shows that ϵ_r shifts the capacitance upward but does not modify how the capacitance scales with gate length. An additional factor, which is not visualized in Fig. 4.7, is the technological motivation behind using high-permittivity dielectrics in modern MOSFETs. While the simulations focus on how the capacitance behaves as a function of gate length, it is also relevant to consider why high-permittivity materials are used in real MOSFET technology. As the gate length is reduced, maintaining electrostatic control over the channel becomes increasingly challenging, which traditionally required thinning the gate oxide. However, as the oxide becomes too thin, direct quantum tunneling yields excessive leakage current. To avoid this while still achieving a high gate capacitance, materials with a higher dielectric permittivity are used. Historically SiO_2 with $\epsilon_r \approx 3.9$ was used, which is the dielectric used in Model 1 and Model 2, but it has been replaced with high permittivity materials such as HfO_2 with $\epsilon_r \approx 25$, enabling both increased capacitance and reduced leakage current.

4.2.5 Table to summarize the capacitive effect of the parameters on Model 1

Since several parameters influence the capacitance, the overview table 4.1 is included to clarify the main trends.

4.2.6 Varying multiple parameters

After examining each parameter individually, it is useful to consider how the different geometric and material properties interact. The capacitance in this 2D model is governed by a balance between vertical and lateral electric-field coupling, and the relative importance of each parameter depends on how the fringing fields redistribute themselves in the device cross-section. Furthermore, a full parametric sweep varying all four parameters at once was performed in COMSOL where the gate length was swept between

Table 4.1: Table to summarize the capacitive effect of the parameters on Model 1.

Parameter	Capacitive effect on Model 1
Gate length L_{gate}	A longer gate increases the capacitance by extending the region over which the gate modulates the channel charge.
Dielectric width $W_{dielectric}$	A wider dielectric region reduces the capacitance by increasing the lateral distance over which the fringing fields must extend.
Oxide thickness t_{ox}	A thicker oxide increases the capacitance in this 2D geometry since the fringing fields can spread through a larger dielectric volume.
Dielectric permittivity ϵ_r	A higher permittivity strengthens the electric-field coupling between gate and channel, and therefore increases the capacitance.

10-80 nm, the dielectric width and the oxide thickness between 1-15 nm, and the dielectric permittivity between 1-20.

Across all simulations, the relative permittivity ϵ_r has the strongest and most consistent influence. Increasing ϵ_r always leads to a higher capacitance, since a larger dielectric constant enhances the electrostatic coupling between the gate and the channel. This trend appears clearly regardless of the geometric configuration. The gate length L_{gate} also increases the capacitance in a predictable way. A longer gate provides a larger region over which charge can be modulated, and the total capacitance therefore grows approximately linearly with L_{gate} . This behavior is robust and remains visible even when other parameters are varied simultaneously.

The geometric parameters that control the separation between the gate and the channel are the dielectric width $W_{dielectric}$ and the oxide thickness t_{ox} . A thicker oxide pushes the gate further from the channel and reduces the capacitance, whereas a smaller dielectric width has the opposite effect: it shortens the lateral distance the fringing fields must travel, which is analogous to reducing the oxide thickness in the vertical direction. As a result, decreasing the dielectric width increases the capacitance. However, their impact is more sensitive to the specific geometry than the effects of ϵ_r or L_{gate} . For some parameter configurations, their influence on the capacitance is pronounced, while in others it is less visible and partially masked by the stronger influence of the other parameters.

No contradictory behavior appears in the dataset, and the lowest capacitance values arise for the combination of low permittivity, large dielectric width, small oxide thickness, and short gate length. This is consistent with the expected electrostatics, since all of these conditions weaken the gate-channel coupling. However, to fully understand the capacitive behavior of the device, the parasitic contribution must be isolated and analyzed separately.

4.2.7 Parasitic capacitance analysis

The total gate capacitance obtained from Model 1 contains two contributions: a length-dependent intrinsic component due to charge modulation in the quantum well, and a length-independent parasitic component arising from the fringing fields. To separate these two contributions in a consistent way, the

Table 4.2: Extracted parasitic capacitance parameters for different dielectric widths.

Dielectric width (nm)	Slope a (F/nm/ μm)	$2C_p$ (F/ μm)	L_{fringe} (nm)
3	2.426×10^{-19}	7.664×10^{-16}	3160.6
4	1.950×10^{-19}	6.668×10^{-16}	3414.6
5	1.500×10^{-19}	6.003×10^{-16}	4001.8

linear decomposition introduced in the theory section is applied.

$$C(L_{gate}) = a L_{gate} + 2C_p \quad (4.3)$$

where the slope a represents the intrinsic capacitance per unit length. The parameter C_p corresponds to the length-independent parasitic capacitance associated with the gate-source and gate-drain regions. Since these parasitic components do not scale with gate length, their influence increases as the channel is shortened.

Dielectric width influence

To extract the parasitic contribution, the capacitance-length curves were fitted linearly for different dielectric widths, as illustrated in Fig. 4.8. The marked points at $L_{gate}=0$ correspond to $2C_p$, obtained from linear extrapolation since the intercept represents the length-independent part of the capacitance.

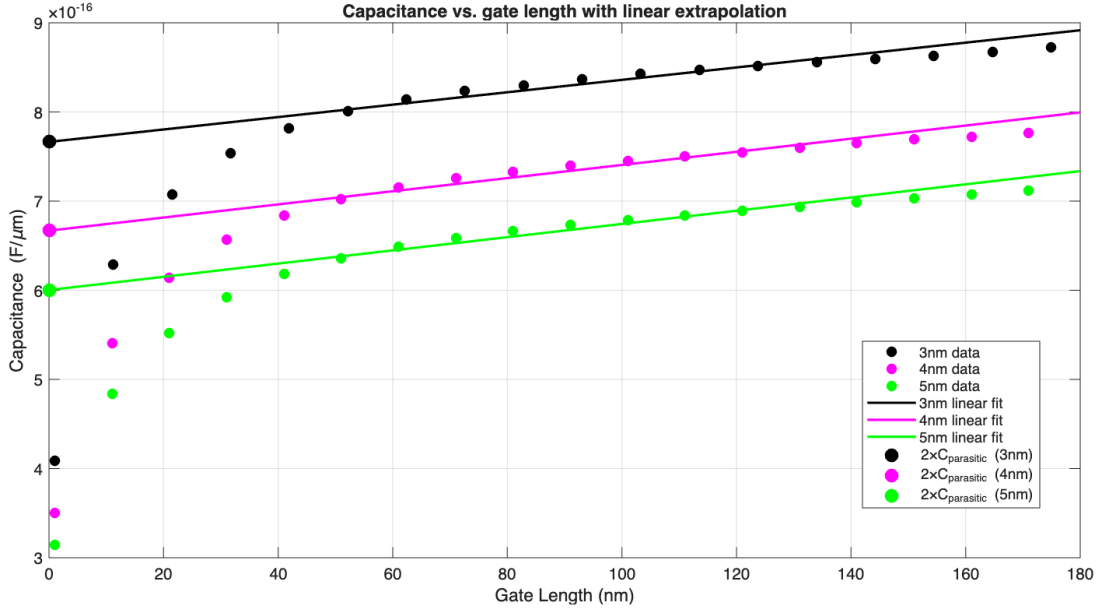


Figure 4.8: Capacitance vs gate length for three different dielectric widths: 3, 4 and 5 nm. This also includes a linear fit and an extrapolation to $L_{gate} = 0$, which allows for the parasitic capacitance to be extracted for each curve which are the values seen as Y in the boxes to the left.

The results are summarized in Table 4.2, where a weakly decreasing trend for the parasitic capacitance with dielectric width is observed. This downward trend can be explained by the fact that increasing the dielectric width pushes the fringing-field region farther from the gate, slightly reducing the strength of the

gate-source and gate-drain coupling. The slope of each fitted line changes more noticeably, indicating that the intrinsic capacitance is strongly affected by the dielectric geometry. Ultimately, this highlights that the parasitic components' contribution has a larger impact for shorter gate lengths.

Furthermore, as the dielectric width increases, L_{fringe} increases as well. Although the parasitic capacitance decreases slightly, the intrinsic capacitance decreases even more strongly, which is why the ratio $2C_p/a$ and hence L_{fringe} grows. In other words, widening the dielectric weakens the intrinsic gate control faster than it suppresses the parasitic coupling, leading to a larger effective fringing length. This demonstrates that the dielectric width strongly influences device scalability, since a large L_{fringe} causes the parasitic part of the capacitance to dominate in short-channel devices.

Oxide thickness influence

The same linear-extrapolation analysis can be performed for the oxide thickness. By extrapolating the linear region of the C - L_{gate} characteristics to $L_{gate}=0$ from Fig. 4.6, it is possible to obtain $2C_p$ for this parameter set. The corresponding extracted parasitic capacitance can be seen in Fig. 4.9.

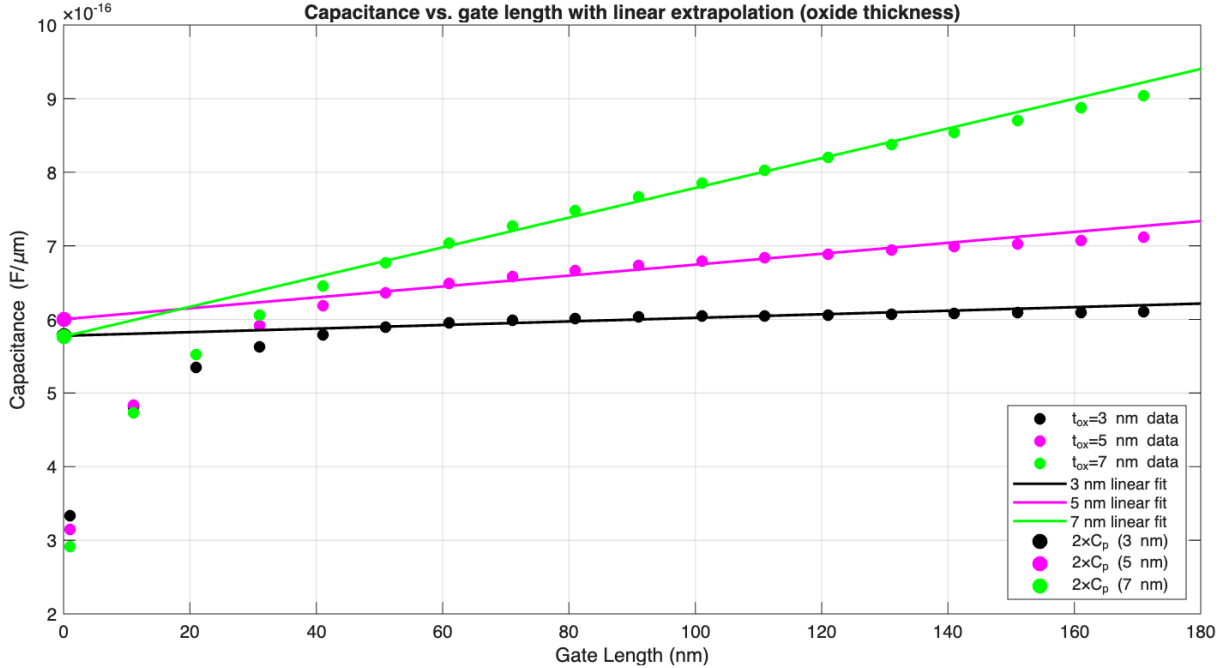


Figure 4.9: Capacitance as a function of gate length for different oxide thicknesses, including linear fits in the quasi-linear region and extrapolated intercepts used to extract the parasitic capacitance $2C_p$.

A summary of the findings from this graph can be seen in Table 4.3. For these values of t_{ox} , no specific trend can be seen for the parasitic capacitance. The lack of a clear trend in $2C_p$ indicates that oxide thickness mainly affects how far the fringing fields extend, rather than the size of the parasitic capacitance itself.

However, the fringing length L_{fringe} decreases rapidly with increasing oxide thickness. This suggests that the thickest oxide has the best gate control because the fringing fields become more localized in thicker

Table 4.3: Extracted parasitic capacitance parameters for different oxide thicknesses.

t_{ox} (nm)	Slope a (F/nm/ μm)	$2C_p$ (F/ μm)	L_{fringe} (nm)
3	2.426×10^{-19}	5.779×10^{-16}	2381.63
5	7.411×10^{-19}	6.003×10^{-16}	809.93
7	2.020×10^{-18}	5.767×10^{-16}	285.56

Table 4.4: Extracted parasitic capacitance parameters for different dielectric permittivities.

ε_r	Slope a (F/nm/ μm)	$2C_p$ (F/ μm)	L_{fringe} (nm)
1	7.531×10^{-19}	4.964×10^{-16}	659.08
3.9	7.411×10^{-19}	6.003×10^{-16}	809.93
20	6.573×10^{-19}	1.131×10^{-15}	1720.11

oxides. This yields a larger ratio of intrinsic capacitance to parasitic capacitance.

Dielectric permittivity influence

To finish the work, it is needed to investigate the effect of the dielectric permittivity. This is done in the same way, where the corresponding results can be seen in Fig. 4.10 and Table 4.4.

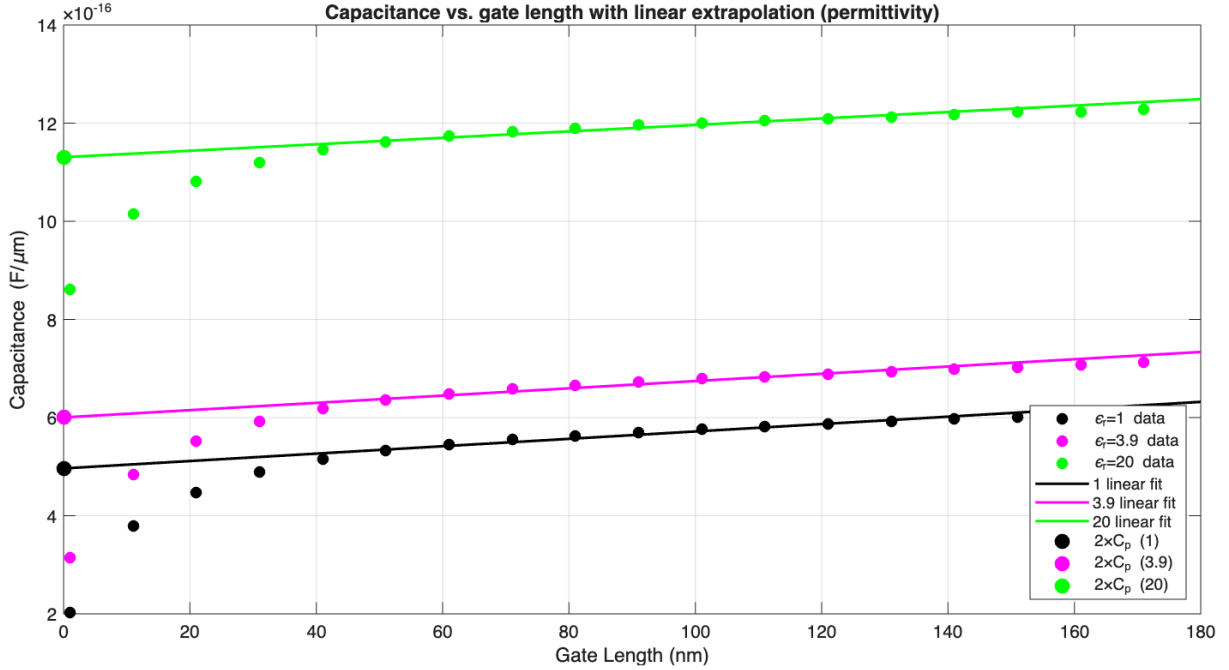


Figure 4.10: Capacitance as a function of gate length for different dielectric permittivities, including linear fits in the quasi-linear region and extrapolated intercepts used to extract the parasitic capacitance $2C_p$.

The parasitic capacitance increases with ε_r , and the effective fringing length grows as well. Moreover, a higher dielectric permittivity strengthens the fringing-field coupling, which weakens the gate's electrostatic control. This effect is particularly pronounced at short gate lengths.

Combined influence of geometry and material parameters

To summarize the influence of the four parameters on the intrinsic and parasitic capacitance components, an overview is given in Table 4.5.

Table 4.5: Influence of geometrical and material parameters on intrinsic, parasitic, and fringing-field capacitance components.

Parameter	Intrinsic capacitance (slope a)	Parasitic capacitance ($2C_p$)	Fringing length (L_{fringe})	Physical explanation
Gate length L_{gate}	Increases linearly with L_{gate}	No effect (length-independent)	No effect	Intrinsic charge modulation scales with gate area; fringing fields do not depend on gate length.
Dielectric width $W_{\text{dielectric}}$	Decreases with increasing width	Slightly decreases with increasing width	Increases with increasing width	Wider dielectric spacing weakens intrinsic gate control more than it reduces parasitic coupling.
Oxide thickness t_{ox}	Strong increase with thickness	No clear trend	Strong decrease with thickness	Thicker oxide confines fringing fields better, improving gate control.
Dielectric permittivity ε_r	Slight decrease with increasing ε_r	Strong increase with increasing ε_r	Strong increase with increasing ε_r	High- ε_r materials strengthen fringing-field coupling, increasing parasitics.

Since $L_{\text{fringe}} = 2C_p/a$, it indicates how large the parasitic contribution is compared to the intrinsic part at short gate lengths. A larger fringing length therefore means that the parasitic part makes up a bigger share of the total capacitance. Among the studied parameters, the dielectric permittivity has the strongest increasing effect on this share, since both $2C_p$ and L_{fringe} increase clearly with ε_r . In contrast, oxide thickness has the strongest reducing effect on L_{fringe} , even though its trend in $2C_p$ is less clear. The dielectric width has a more moderate influence, with only small changes in $2C_p$ but a clearer increase in L_{fringe} due to the reduced intrinsic capacitance.

4.3 Results Model 2

4.3.1 Overview of the Schrödinger–Poisson model

Model 2 focuses solely on the quantum well region and solves the Schrödinger–Poisson system self-consistently using COMSOL. This approach makes it possible to extract key quantum-mechanical quantities such as eigenenergies, wavefunctions and carrier concentrations. Because the model is one-dimensional, which significantly reduces computational cost, it captures only the potential variation along the growth direction x . Note that in Model 2 the growth direction is denoted x , corresponding to the same physical direction previously referred to as z .

The model enables accurate visualization of the potential profile, bound states, and charge distribution inside the quantum well. By solving the Schrödinger–Poisson system self-consistently, the potential, charge distribution, and eigenstates become mutually consistent at the chosen operating point. These results reflect the physics captured by the model, and serve as the basis for the upcoming mobility calculations, all evaluated at the chosen operating point of $V_{\text{gate}} = -0.08$ V.

Fig. 4.11 shows the conduction- and valence-band edges across the heterostructure along the growth direction x , obtained from the self-consistent Schrödinger–Poisson solution. The blue curve represents

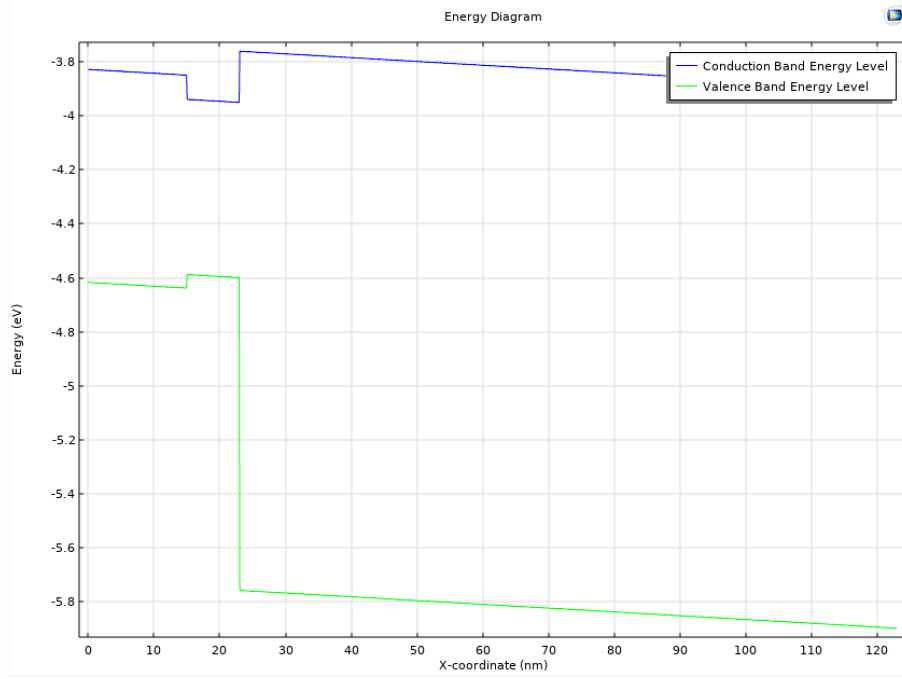


Figure 4.11: Energy diagram showing the conduction- and valence-band edges along the heterostructure in the growth direction x , obtained from the self-consistent Schrödinger–Poisson solution.

the conduction-band edge E_c . E_c exhibits distinct steps at the material interfaces of the heterostructure. The slope of the curve within a material corresponds to the electrostatic band bending originating from ionized dopants. The green curve, E_v , shows the valence-band edge, which follows the same material transitions as E_c .

With this understanding of the band alignment, the results in Fig. 2.2 are consistent with theory. E_c reaches its minimum value in the quantum well, while the surrounding barriers are higher to confine the electrons. The downward slope in the n -doped region drives electrons toward the well, establishing the expected confinement profile. Based on this band alignment, the corresponding potential energy profile used in the Schrödinger equation is shown in Fig. 4.12. This potential can be interpreted as what the electrons “see” when traveling along the growth direction.

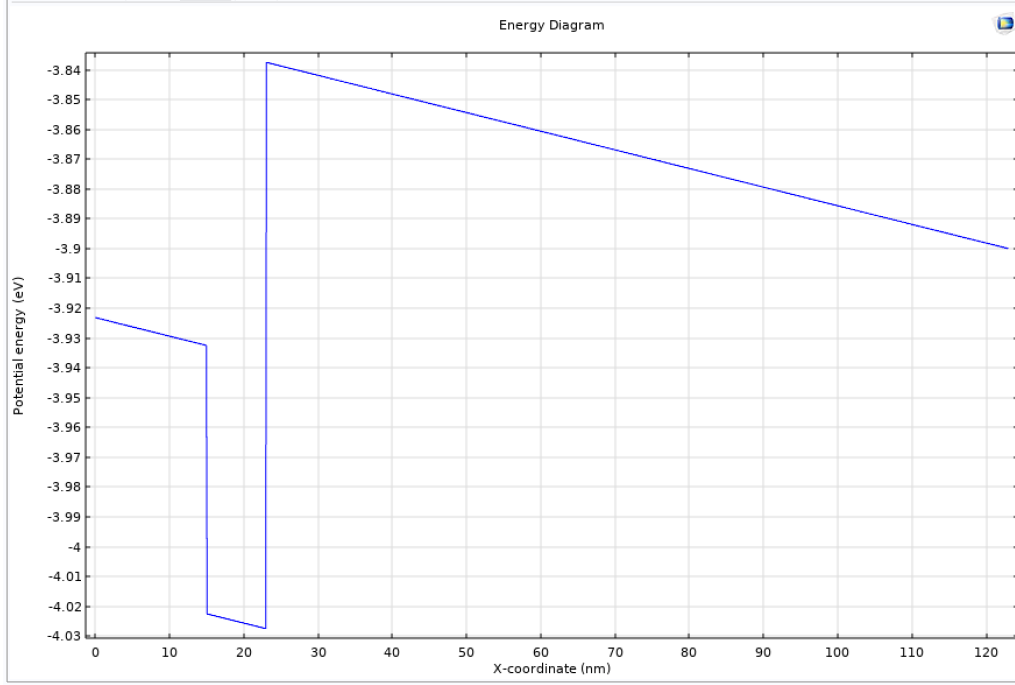


Figure 4.12: Potential energy V plotted along the growth direction x of the heterostructure.

The potential energy, $V(x)$, reaches its minimum value inside the quantum well. On each side of the well, the two high plateaus represent the barriers, which have a higher band edge and therefore a higher potential energy. These barriers restrict the electron motion and confine them within the well. The potential energy follows the shape of the conduction-band edge diagram, but expressed as an energy potential for the Schrödinger equation. Given the shallow and narrow well, only a small number of bound states are expected, and the modest barrier height implies weaker confinement and increased tunneling into the surrounding regions. Taken together, this potential therefore forms the basis for the subsequent calculation of eigenenergies and wavefunctions.

4.3.2 Eigenstate and wavefunction characteristics

The probability density $|\psi(x)|^2$ for the ground state is shown in Fig. 4.13. It is obtained by squaring the magnitude of the wavefunction and represents the likelihood of finding an electron at a given position. The ground state exhibits a single peak centered in the quantum well, consistent with the lowest-energy states without nodes.

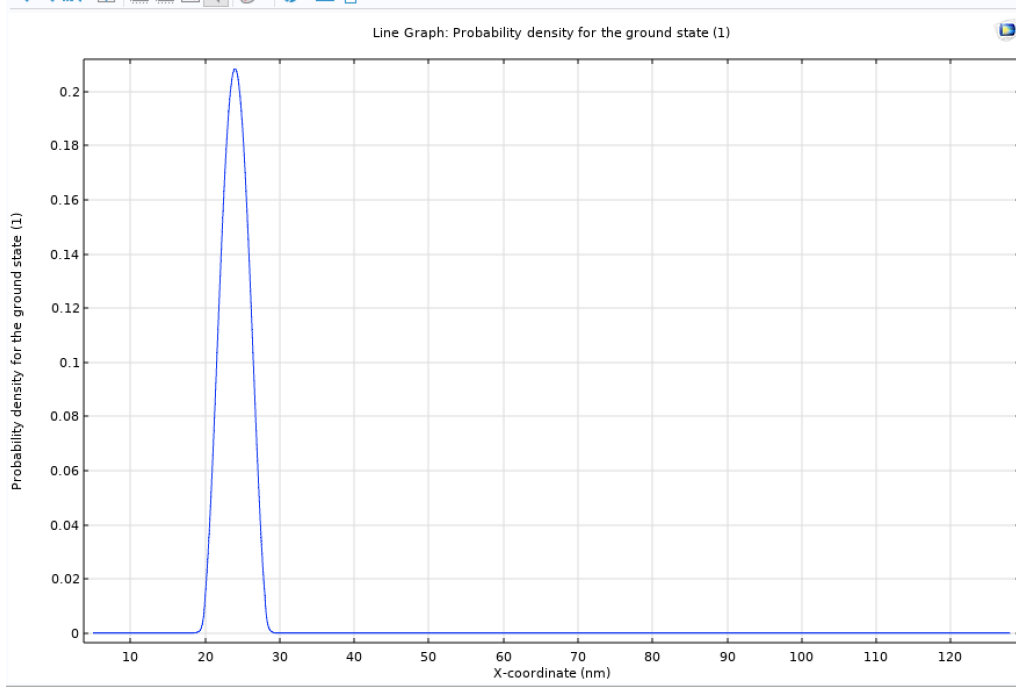


Figure 4.13: Probability density $|\psi(x)|^2$ for the ground state. The distribution shows a single peak in the well region, corresponding to the eigenenergy of 0.030991 eV. The apparent localization is enhanced by the squaring of the wavefunction; the actual confinement is weak due to the shallow potential.

Due to the shallow and narrow quantum well, there is only one bound state, which is visible in Fig. 4.14. Its eigenenergy lies below the barrier height, confirming that the state is bound, albeit weakly. This is expected because the well depth is small and the barrier height is modest, which limits the number of allowed quantized energy levels. The wavefunction oscillates inside the well and decays slowly into the surrounding barriers, and due to the slow nature of this decay this results in significant tunneling. Although the probability density appears strongly localized, this is a consequence of the squaring of the wavefunction, which suppresses the tunneling tails. The underlying wavefunction exhibits significant penetration into the barriers, indicating weak confinement consistent with the shallow potential.

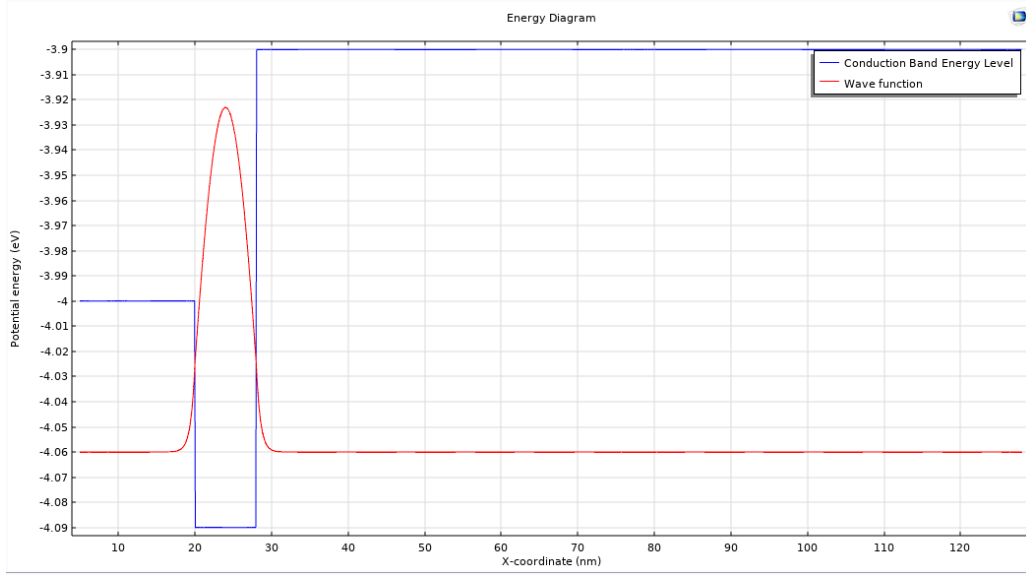


Figure 4.14: The single bound eigenstate supported by the shallow quantum well. The wavefunction is weakly confined and exhibits significant tunneling into the barriers.

How this eigenstate influences the carrier concentrations in the structure is seen in Fig. 4.15, where the electron and hole concentrations are plotted as a function of the growth direction x .

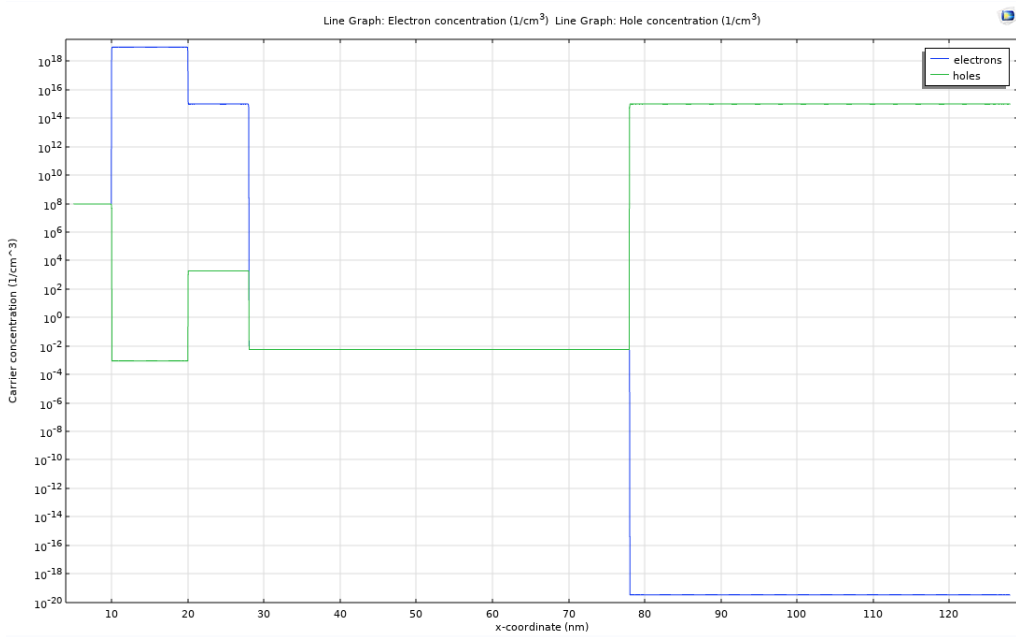


Figure 4.15: Electron and hole concentration along the growth direction.

The hole concentration (green curve) is highest on the right-hand side of the barrier, where the material is p-doped. Moving into the quantum well, the hole density drops by many orders of magnitude and remains low throughout the well and the left-hand side of the device. This behavior reflects the valence-band

profile and the strong p-type doping on the right side, which pins the Fermi level close to the valence band only in that region.

The electron concentration (blue curve) shows the opposite trend. It is strongly suppressed in the p-type region on the right, but increases sharply inside the quantum well and remains high in the n-type region on the left, and increases even more in the heavily n-doped domains.. The electron density therefore follows the doping profile rather than directly mirroring the conduction-band edge.

With the ground state wavefunction and carrier distribution established, the next step is to look at how the spatial extent of the wavefunction influences scattering. This is captured by the form factor, which describes how strongly the electron wavefunction responds to spatial variations of different length scales in the impurity potential. The form factor therefore provides the link between the quantum-mechanical confinement and the upcoming calculation of screened Coulomb scattering and mobility.

4.3.3 Form factor results

The form factor $F(q)$ quantifies how strongly the ground-state wavefunction in the quantum well couples to an impurity located at a given distance and varying with spatial frequency q . Using the wavefunction $\psi(x)$ obtained from the Schrödinger-Poisson solution in COMSOL, $F(q)$ is evaluated numerically in MATLAB for each doped region, and the scripts used for this are provided in Appendix B. Fig. 4.16 shows the resulting form factors for dopants in the left and right barriers.

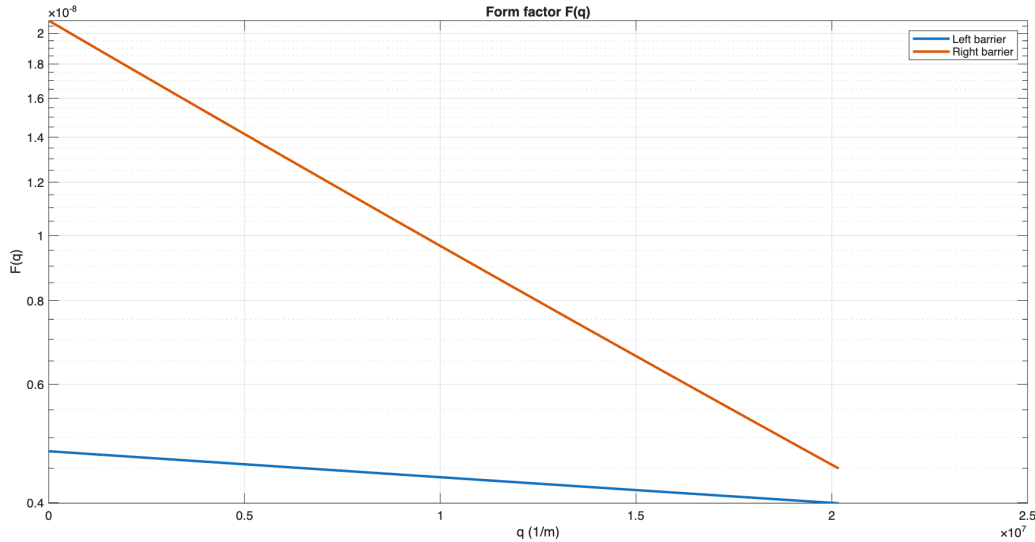


Figure 4.16: Form factor $F(q)$ for dopants in the left and right barriers.

The values for $F(q)$ are several orders of magnitude below unity because the dopants lie outside the quantum well, which gives a very small spatial overlap with the confined wavefunction. The right barrier form factor is larger than the left barrier one due to the much wider doped region (78-128 nm), compared to the narrow doped region in the left barrier (10-20 nm). Since the form factor is a double integral over both the wavefunction and the width of the doped region, a wider region produces a larger value even if it is located farther away.

Both curves decrease monotonically with increasing q . This follows directly from the exponential factor $e^{-q|z-z'|}$ in the definition of the form factor: as q increases, the exponential decays more rapidly, which reduces the contribution from impurities located at positions z' outside the well. The resulting form factors determine how strongly each doped region contributes to the screened Coulomb potential, which is the next step of the analysis.

4.3.4 Screened Coulomb potential

Using the numerically obtained form factors, the screened Coulomb potentials for the left barrier, right barrier, and well dopants were evaluated in MATLAB. The results are shown in Fig. 4.17, where the y -axis is logarithmic and $U(q)$ is plotted as a function of the wavenumber q .

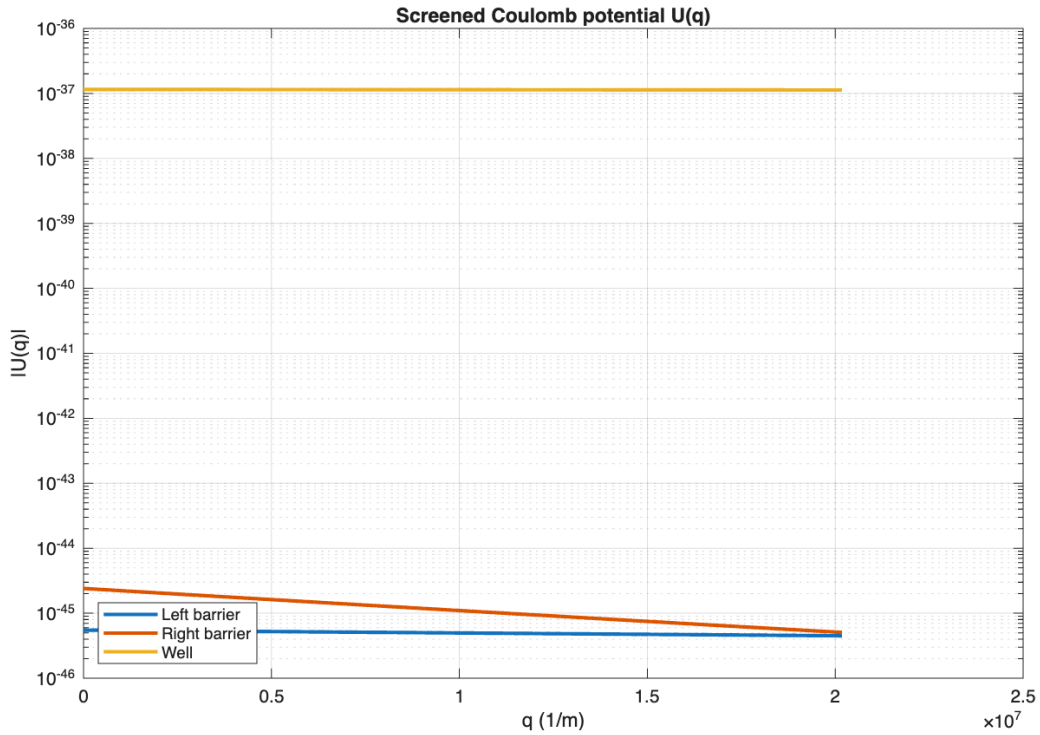


Figure 4.17: Screened Coulomb potential $U(q)$ for dopants in the left barrier, right barrier, and inside the well.

To interpret these results, it is useful to recall the relation between the form factor and the screened potential. Combining equations 2.24–2.27 gives the proportionality

$$U(q) \propto \frac{F(q)}{q + q_{TF}},$$

which shows that the potential is determined by both the wavefunction overlap and the Thomas–Fermi screening. Since the scattering rate depends on $U(q)^2$, a larger potential leads to stronger scattering and therefore a lower mobility.

The yellow curve, corresponding to dopants inside the well, has the highest magnitude and is nearly

constant over the entire q -range. This reflects that $F(q) \approx 1$ for in-well impurities and that $q \ll q_{TF}$ in the relevant interval, meaning that the screening is dominated by the constant Thomas-Fermi term q_{TF} rather than by q itself. As a result, $U(q)$ becomes almost independent of q .

Below this, the right-barrier curve lies above the left-barrier curve, which is consistent with the form factor results from Fig. 4.16. Although the right-barrier dopants are located farther from the well, their width produces a larger integrated contribution to the form factor and therefore to $U(q)$.

All curves decrease monotonically with increasing q , which follows directly from the exponential factor $e^{-q|z-z'|}$ and the screening term $1/(q + q_{TF})$. Since $U(q)_{well}$ is several orders of magnitude larger than the barrier contributions, impurities inside the well dominate the Coulomb scattering and result in the largest decrease in carrier mobility. This is consistent with the transport integral, where large-angle (large q) scattering events contribute most strongly to momentum relaxation. In contrast, as the dopants in the barriers are located far away from the wavefunction, $U(q)_{barriers}$ are extremely small and the electrons experience only a weak potential. Due to this, the scattering is weak and the mobility remains high when dopants are placed in the barriers rather than inside the well.

4.3.5 Transport time and mobility

Equation 2.22 shows that the transport time is governed by the scattering potential through the proportionality

$$U(q)(1 - \cos \theta) \propto \frac{1}{\tau_{tr}}.$$

As $U(q)$ represents the Fourier component of the screened Coulomb interaction from a single dopant, while the factor $(1 - \cos \theta)$ weights the contribution from different scattering angles. Since $|U(q)|$ is largest for dopants located inside the well, this case yields the strongest scattering and therefore the shortest transport time. This is because impurities inside the well contribute strongly at large q , where the factor $(1 - \cos \theta)$ enhances momentum-relaxing scattering. For dopants in the barriers, $|U(q)|$ is reduced by several orders of magnitude due to the spatial separation from the wavefunction, resulting in extremely long transport times. For the two barriers, the right-barrier dopants have a slightly shorter transport time, consistent with their marginally larger value of $|U(q)|$.

From the mobility expression in equation 2.16

$$\mu \propto \tau_{tr}$$

it follows directly that the mobility exhibits the same trend as the transport time. The lowest mobility is therefore obtained when dopants are placed inside the well, and the highest mobility occurs for dopants in the left barrier.

The same spatial separation that suppresses Coulomb scattering from barrier dopants also influences the electrostatic response of the structure. Consequently, the dopant position affects not only the mobility but also the capacitance. The connection between these two quantities and how to tune both for improved device performance is therefore the next topic of discussion.

4.4 Gate dependent mobility and its relation to capacitance

Having established the electrostatic behavior of the device in Model 1 together with the quantum-mechanical confinement in Model 2, the next step is to make use of the combined COMSOL-MATLAB workflow. This allows the electron mobility to be evaluated as a function of gate voltage, and also makes it possible to examine its correlation with the total gate capacitance. However, it is important to note that this model is a simplification and does not take all scattering mechanisms into account, and should therefore serve as a trend-level interpretation rather than exact quantitative values.

4.4.1 Mobility as a function of gate voltage

Fig. 4.18 shows the extracted mobility μ as a function of the applied gate voltage V_{gate} between -0.25V and 0.25V .

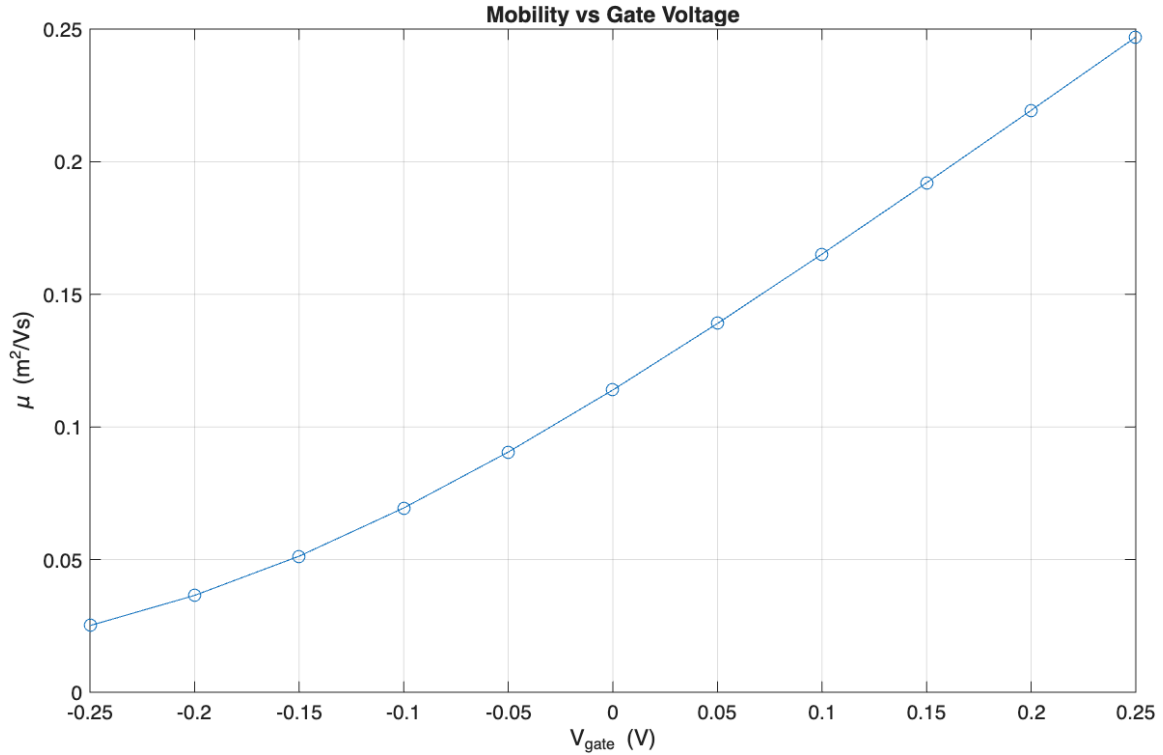


Figure 4.18: Mobility μ as a function of gate voltage V_{gate} .

The graph shows that μ increases clearly and monotonically with the applied gate voltage. This behavior is expected in a QWFET, where the formation of a 2DEG depends on how strongly the gate confines electrons in the quantum well. At lower gate voltages, the electron density in the well is small and the wavefunction lies closer to the doped barrier. As a result, impurity scattering from dopants increases, resulting in low mobility.

As the gate voltage increases, the wavefunction shifts further into the center of the well, reducing its overlap with ionized donors and therefore weakening impurity scattering. The near-linear trend suggests that Coulomb impurity scattering is the dominant mechanism in this regime. In a more complete model,

the screening of impurities would also increase with electron density. However, in the simplified model used in this work, the Thomas–Fermi screening wave vector q_{TF} is treated as a constant, so screening does not contribute to the mobility variation. The magnitude and trend of the extracted mobility are consistent with impurity-limited transport in a modulation-doped quantum well.

4.4.2 Mobility as a function of capacitance

When the mobility is instead plotted directly against the total gate capacitance, the same physical trend becomes visible from a different perspective, as illustrated in Fig 4.19.

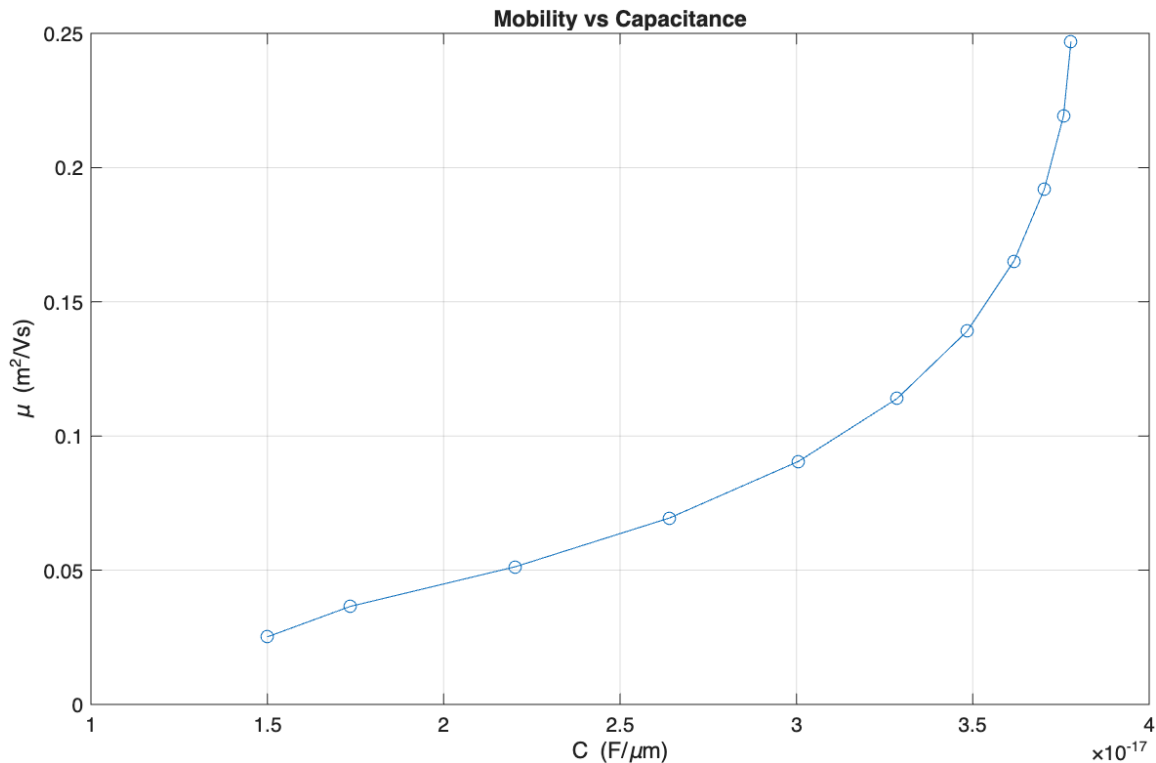


Figure 4.19: Mobility μ as a function of capacitance.

A higher capacitance means more charge induced into the quantum well, which corresponds to a higher electron density and a stronger electrostatic coupling between the gate and the channel. Since both mobility and capacitance depend on the spatial distribution of the wavefunction in the quantum well, both quantities follow a clear positive correlation.

The weakly convex shape illustrates that the mobility increases more strongly at higher capacitance values. This is in line with the screening mechanism: as the electron density increases, the effective impurity potential becomes weaker. This yields a longer relaxation time and thus a higher mobility. For low capacitance values, impurity scattering is strong and the mobility is low, consistent with the screening effect becoming stronger as the electron density increases.

4.4.3 Interpretation and connection to device operation

To conclude, the mobility trends, both as a function of gate voltage and as a function of gate capacitance, show a consistent physical behavior. The results indicate that mobility and capacitance are linked through the electrostatics of the quantum well, which determine the carrier density and the position of the wavefunction. The observed trends reflect a shift from impurity-limited transport at low carrier densities to more efficient transport as the gate voltage increases. The latter coincides with increased carrier density and thus enhances screening and reduces scattering.

In addition to these trends, there are also regimes where both mobility and capacitance change very little or remain almost constant. These regions appear when the wavefunction and subband occupation in the quantum well are relatively insensitive to small variations in gate voltage or device geometry. For these cases, the carrier distribution and scattering mechanisms remain nearly unchanged, resulting in nearly invariant transport and electrostatic properties. This in turn suggests that the device can maintain stable electrical performance even without ideally optimized parameters. From a design and fabrication perspective, this is particularly useful, as it highlights that some parts of the design space are less sensitive to variations and consequently more robust.

Chapter 5

Conclusion and Outlook

This work has investigated how geometrical and dielectric parameters affect capacitance and mobility in a superconducting QWFET-structure. This was done by analyzing two supplementary models: one two-dimensional electrostatic COMSOL-model to extract intrinsic and parasitic capacitance, and one one-dimensional Schrödinger-Poisson model to describe quantum confinement and calculate electron mobility.

Results show that gate-length almost linearly increases the intrinsic capacitance, with deviations at short lengths where fringing fields dominate. The dielectric width reduces the capacitance by weakening the lateral electrostatic coupling, while the oxide thickness experiences a non-monotonic behavior where thin oxides amplify fringing fields, but thicker oxides weaken the connection between the gate and the channel. Dielectric permittivity has an overall increasing effect on the capacitance, and also increases the parasitic component. The extracted parasitic capacitance is reduced weakly with increased dielectric width, lacks a trend for the oxide thickness, and increases visibly with higher permittivity.

The Schrödinger-Poisson model shows that the electrons are weakly confined in the quantum well, and that impurity scattering is the dominant mechanism limiting the mobility. Mobility decreases with increasing dopant concentrations, which follows from the wavefunction's overlap with the doped regions. Altogether, the results highlight a clear trade-off. Parameters that improve the electrostatic control tend to increase scattering and hence reduce the mobility. Consequently, this work could not identify any clear optimal dimensions, but instead trends that can be used as guidelines when designing future QWFET devices.

Modeling in this work is limited to two dimensions, stationary conditions, and a mobility model where only impurity scattering is accounted for. The results should therefore be interpreted as qualitative trends rather than exact quantitative values.

Future work could include three-dimensional modeling of fringing fields and parasitic capacitance, more comprehensive mobility models incorporating phonon and alloy scattering, and a systematic analysis of gate-bias-dependent effects. It would also be valuable to investigate how the doping profile can be optimized to balance gate control and mobility, and to compare the numerical results with experimental measurements for further validation. Finally, this work establishes a coherent workflow in which COMSOL

and MATLAB are used in combination, providing a foundation for more advanced simulation and analysis in future studies.

Chapter 6

Bibliography

- [1] S. M. Sze and K. K. Ng, *Physics of Semiconductor Devices*. Wiley, 3 ed., 2007.
- [2] Y. Taur and T. H. Ning, *Fundamentals of Modern VLSI Devices*. Cambridge, UK: Cambridge University Press, 2 ed., 2013.
- [3] R. Chau, B. Doyle, S. Datta, J. Kavalieros, and M. Zhang, “A 50 nm depletion-mode cmos transistor and the role of parasitic capacitances in nanoscale mosfets,” *IEEE Electron Device Letters*, vol. 25, no. 6, pp. 408–410, 2004.
- [4] T. Ando, A. B. Fowler, and F. Stern, “Electronic properties of two-dimensional systems,” *Reviews of Modern Physics*, vol. 54, no. 2, pp. 437–672, 1982.
- [5] R. Combescot, *Superconductivity: An Introduction*. Cambridge University Press, 2022.
- [6] J. S. Galsin, “Superconductivity,” in *History of Solid State Physics*, Springer, 2025. Accessed: 2025-12-18.
- [7] R. Nave, “Meissner effect for superconductors.” <http://hyperphysics.phy-astr.gsu.edu/hbase/Solids/meis.html>. Accessed: 2025-02-18.
- [8] J. Bardeen, L. N. Cooper, and J. R. Schrieffer, “Theory of superconductivity,” *Physical Review*, vol. 108, no. 5, pp. 1175–1204, 1957.
- [9] DoITPoMS, University of Cambridge, “Cooper pairs.” <https://www.doitpoms.ac.uk/tlplib/superconductivity/cooper.php>. Accessed: 2026-01-06.
- [10] A. Höglund and C. Lundquist, “Coupling a quantum dot to a superconductor.” Project report, FFFN35, Lund University, 2026. Accessed: 2026-01-06.
- [11] C. Yu, “Superconductivity lecture notes.” <https://ps.uci.edu/~cyu/p115B/LectureNotes/superconductivity.pdf>. Accessed: 2026-01-06.
- [12] M. Fu and Y. Lu, *Fabrication and Characterization of Quantum-well Field Effect Transistor*. PhD thesis, Lund University, 2022. Clear description of QWFET layer stack, spacer, and contacts.

- [13] P. Harrison, *Quantum Wells, Wires and Dots: Theoretical and Computational Physics of Semiconductor Nanostructures*. John Wiley & Sons, 2 ed., 2005. Covers Schrödinger–Poisson modelling and quantum well simulations.
- [14] J. H. Davies, *The Physics of Low-Dimensional Semiconductors: An Introduction*. Cambridge, UK: Cambridge University Press, 1998.
- [15] D. L. Pulfrey, *Understanding Modern Transistors and Diodes*. Cambridge University Press, 2012.
- [16] R. Fitzpatrick, “The displacement current.” <https://farside.ph.utexas.edu/teaching/em/lectures/node46.html>, 2023. Accessed: 2026-01-12.
- [17] Y. Hanlunmyuang and P. Sharma, “Quantum capacitance: A perspective from physics to nanoelectronics,” *JOM*, vol. 66, no. 4, pp. 660–663, 2014.
- [18] A. Sonanvane, B. N. Joshi, and A. M. Mahajan, “Analysis of capacitance across interconnects of low-k dielectric used in a deep sub-micron cmos technology,” *Progress In Electromagnetics Research Letters*, vol. 1, pp. 189–196, 2008.
- [19] J. A. del Alamo, *Nanometer-Scale Transistors: Physics, Modeling, and Simulation*. Cambridge, MA: MIT Press, 2013.
- [20] S. V. Suryavanshi, C. D. English, H.-S. P. Wong, and E. Pop, “Scaling theory of two-dimensional field-effect transistors,” *arXiv preprint arXiv:2105.10791*, 2021.
- [21] J. A. del Alamo, “Iii–v fets for future digital applications,” *IEEE International Electron Devices Meeting (IEDM)*, pp. 1–4, 2011.
- [22] M. J. Schulz, V. N. Shanov, and Z. Yin, eds., *Nanotube Superfiber Materials: Changing Engineering Design*. Micro & Nano Technologies, Oxford, UK: William Andrew, 2014.
- [23] S. Das Sarma, E. H. Hwang, S. Kodiyalam, L. N. Pfeiffer, and K. W. West, “Transport in two-dimensional modulation-doped semiconductor structures,” *Physical Review B*, vol. 91, no. 20, p. 205304, 2015.
- [24] W. T. Masselink, “Electron mobility in delta-doped quantum well structures,” in *Negative Differential Resistance and Instabilities in 2-D Semiconductors*, pp. 83–98, Springer, 1990.
- [25] W. Yi, A. A. Kiselev, J. Thorp, *et al.*, “Gate-tunable high mobility remote-doped insb/in_{1-x}al_xsb quantum well heterostructures,” *Applied Physics Letters*, vol. 106, no. 14, p. 142103, 2015.
- [26] COMSOL AB, “COMSOL Company Overview.” <https://www.comsol.com/company>. Accessed: 2025-11-19.
- [27] COMSOL AB, Stockholm, Sweden, *Semiconductor Module User’s Guide*, 2024. Version 6.2.
- [28] COMSOL AB, Stockholm, Sweden, *COMSOL Multiphysics Reference Manual*, 2024. Version 6.2.
- [29] COMSOL AB, Stockholm, Sweden, *COMSOL Multiphysics Reference Manual*, 2024.
- [30] O. C. Zienkiewicz and R. L. Taylor, *The Finite Element Method*. Butterworth-Heinemann, 2000.

Appendix A

Appendix

This appendix contains the technical details that are required to reproduce both the COMSOL-model for Model 1, and the MATLAB-based calculations that are used in this study. As COMSOL is a difficult software to learn on your own, this serves as a practical walkthrough on how to begin with a model. The complete MATLAB-scripts to calculate the form factor, the screened Coulomb potential and the mobility are also included. Together, the appendix is a complete basis to recreate and further develop the model that forms the basis for the results in the main text.

A A COMSOL guide on how to get started with COMSOL and set up the geometry

Model 1 was created using the semiconductor module in COMSOL. This section will guide a first-time user on how to get started in COMSOL, and how to establish a model like Model 1 correctly using geometry parameters.

A.1 Prerequisites to complete before starting to build the geometry

When opening COMSOL, one gets the option to choose either a blank model or a model wizard, as seen in Fig. A.1. The correct option here is to choose model wizard.

Following this, one gets the option to choose the number of space dimensions for the model. For this model, the 2D-option seen in Fig. A.2 was chosen.

To simulate the behavior of a transistor, as seen in Fig. A.3, one can choose either the electrostatics (es) module or the semiconductor (semi) module. The electrostatics module is easier to use, but it does not account for mobile charge carriers or current flow. The semiconductor module is based on both drift-diffusion and Poisson's equations, and is essential when modeling a semiconductor device.

To be able to perform measurements on a model, a study must be chosen. For this purpose, it is convenient to start with a simple stationary study step, which can be seen in Fig. A.4. Additional studies and study steps can be added later.

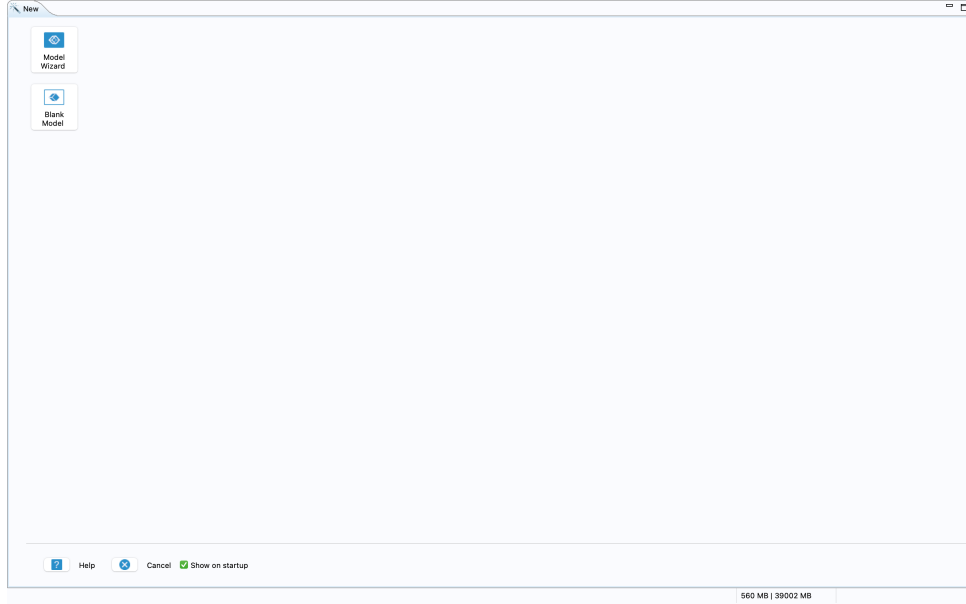


Figure A.1: First choice when opening COMSOL to create a new model.

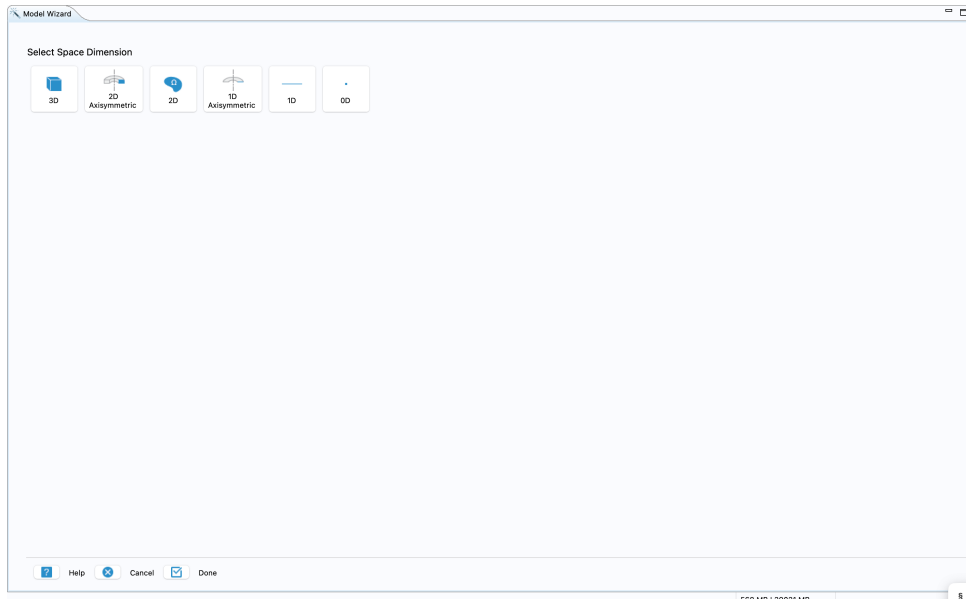


Figure A.2: Choosing space dimension for the model.

To begin building the geometry, one right-clicks the geometry node and chooses from several geometric shapes, as seen in Fig. A.5.

A.2 The geometry of the model

Building the geometry in COMSOL is most efficient when all components are parametrized. In this thesis, maintaining a fixed gate length regardless of other geometric changes was crucial. The final and suggested way of building a geometry in COMSOL is presented below. The model consisted of five semiconductor

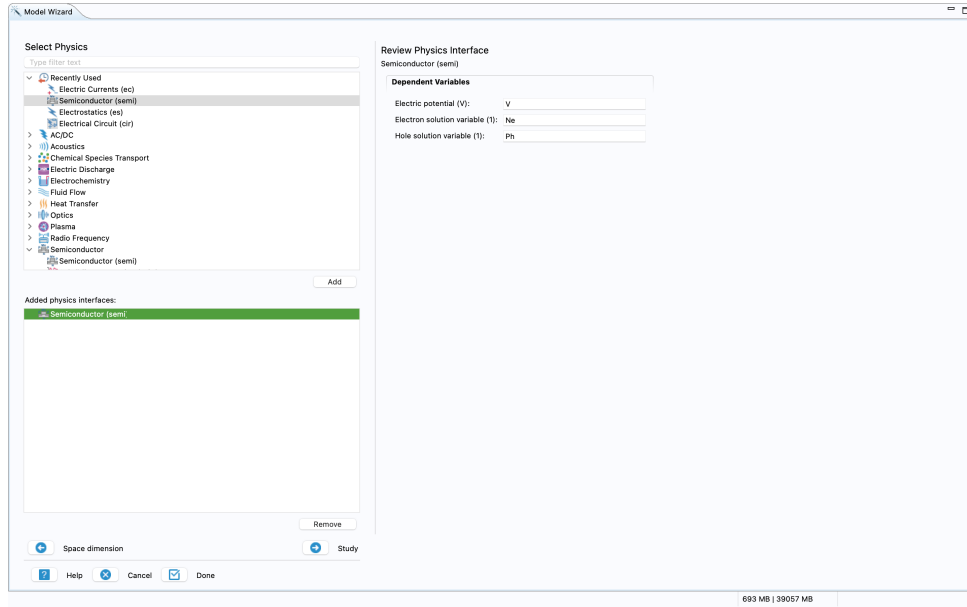


Figure A.3: Choosing which physics to use.

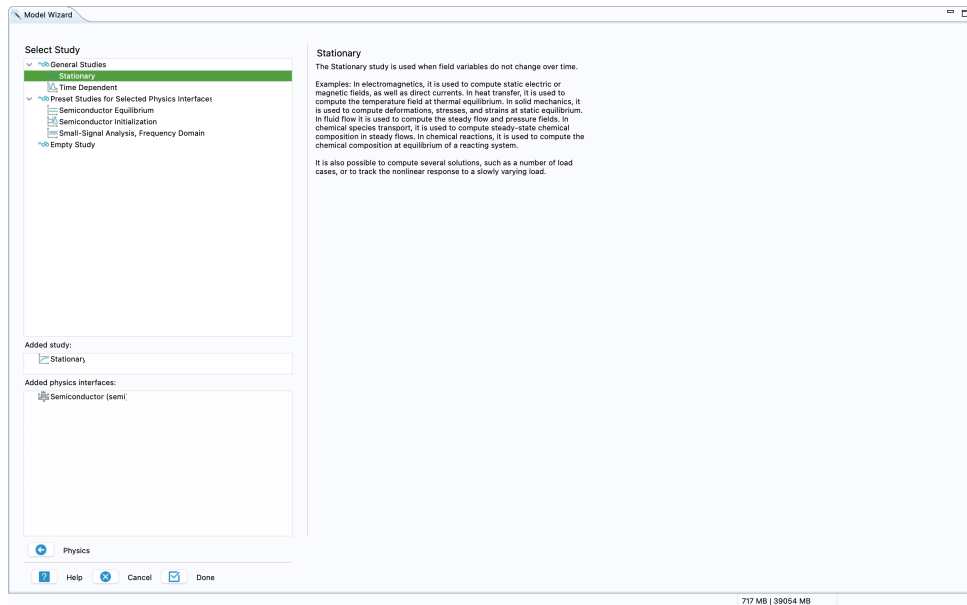


Figure A.4: Adding a study and the options for different types of study.

domains, followed by an oxide layer, three gate boundaries, two dielectrics and two air regions. The geometric parameters used can be seen in Fig. A.6.

The bottom layer, shown in Fig. A.7, was first given a width named *layer1w*, but after adjusting the model to be invariant to changes in gate length, it was given the width $2 \cdot \text{superconductorwidth} + 2 \cdot \text{dielectricw} + \text{gatelength}$. This became the full width of the model and was used as the width for all bottom layers. The positions x and y represent where in the geometry the rectangle should be located;

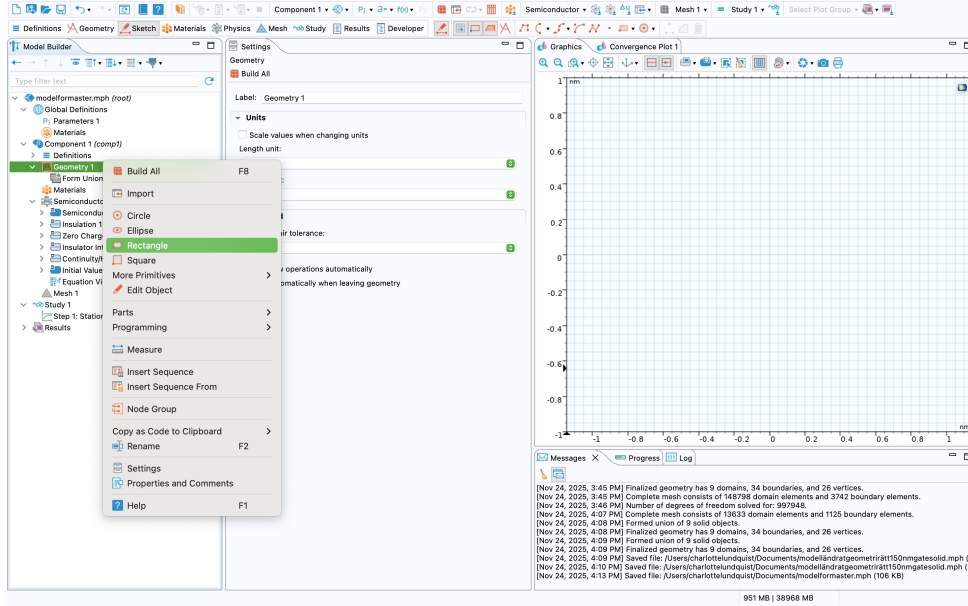


Figure A.5: How to start building the geometry of the model.

Name	Expression	Value	Description
layer1h	20 [nm]	2E-8 m	
layer2h	7 [nm]	7E-9 m	
layer3h1	2 [nm]	2E-9 m	
oxideheight	5 [nm]	5E-9 m	
oxidewidth	60[nm]	6E-8 m	
gateheight	20 [nm]	2E-8 m	
dielectric1h	20 [nm]	2E-8 m	
superconduc12	12 [nm]	1.2E-8 m	
superconduc30	30 [nm]	3E-8 m	
superconduc12	12 [nm]	1.2E-8 m	
dielectricw	5 [nm]	5E-9 m	
dielectric2h	20 [nm]	2E-8 m	
layer3h	4 [nm]	4E-9 m	
layer3h2	2 [nm]	2E-9 m	
gatelength	50 [nm]	5E-8 m	

Figure A.6: Geometry parameters in COMSOL.

for the bottom rectangle, this is $(0,0)$.

The second layer, shown in Fig. A.8, should have the same y -position as the first layer. This method is used for the upcoming layers, seen in Fig. A.9, A.10 and A.11.

The oxide region, shown in Fig. A.12, was built directly on top of the semiconductor layers. Its width equals the distance between the dielectric regions plus the gate length. Its x -position was set to the superconductor width, placing the oxide at the center of the semiconductor domain.

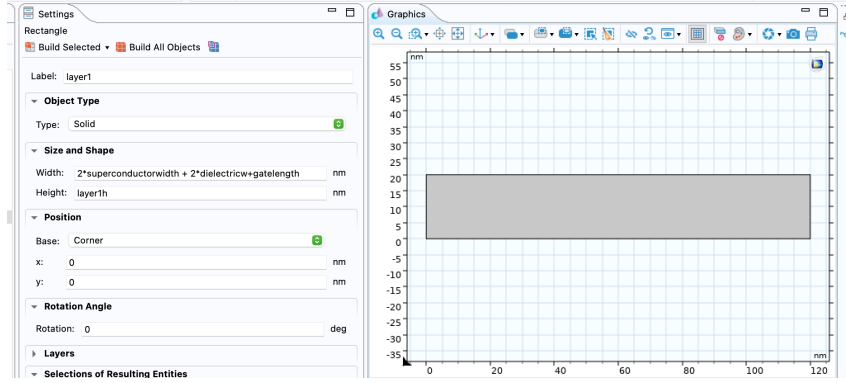


Figure A.7: The bottom layer of the model.

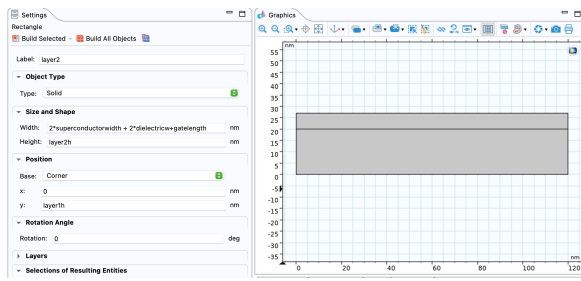


Figure A.8: The second layer of the model.

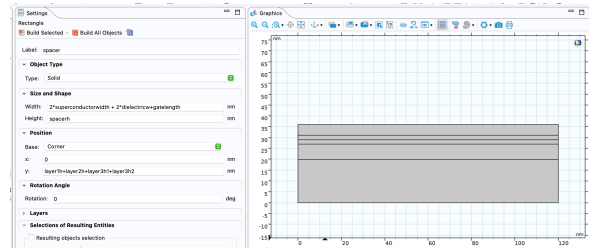


Figure A.9: The third layer of the model.

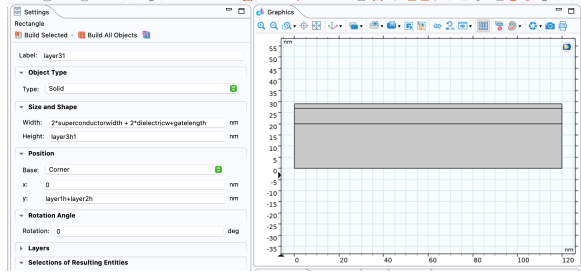


Figure A.10: The fourth layer of the model.

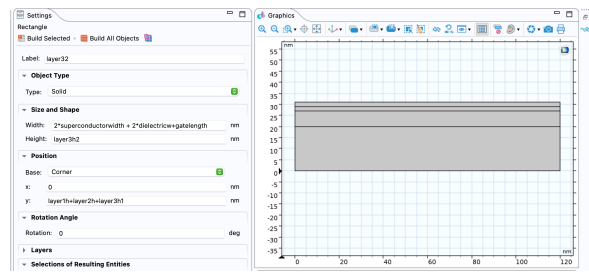


Figure A.11: The fifth layer of the model.

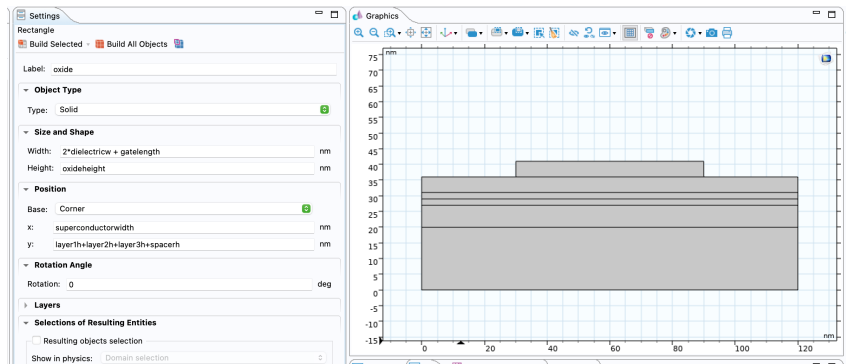


Figure A.12: The oxide layer.

Two regions of air were built, shown in Fig. A.13 and A.14. These were previously used for other purposes during model development, which explains their names.

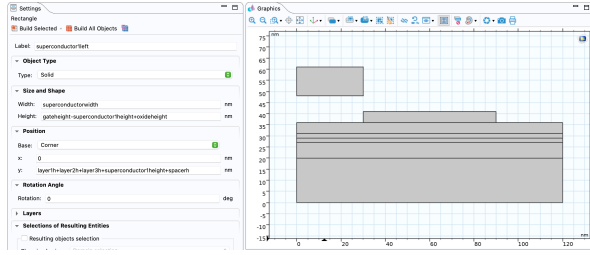


Figure A.13: The left air region.

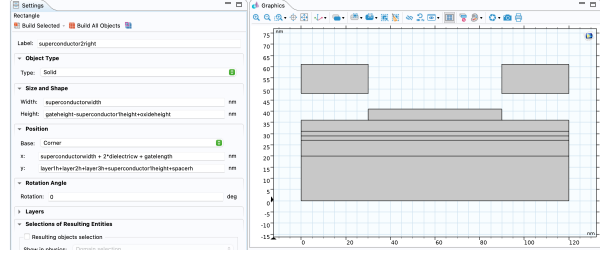


Figure A.14: The right air region.

The two dielectric regions, shown in Fig. A.15 and A.16, were placed between the gate and the air regions, anchored in the oxide. With these widths and heights, the model became easy to adjust and rebuild for different parameter values.

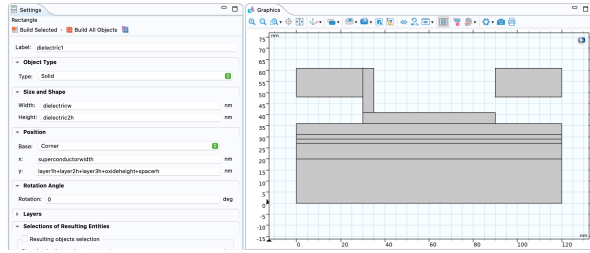


Figure A.15: The left dielectric.

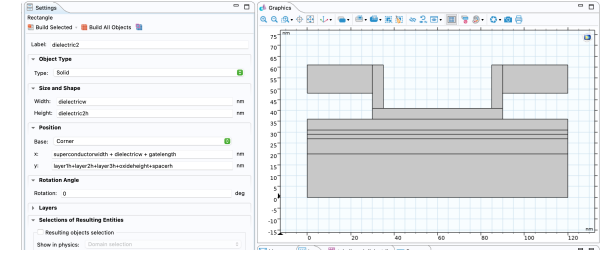


Figure A.16: The right dielectric.

A.3 The materials of the model

Several materials were defined in COMSOL. The semiconductor materials used in the five-layer stack are shown in Fig. A.17–A.19. The oxide region consisted of SiO_2 (Fig. A.20), with its permittivity increased from 3.9 to 20. The air regions (Fig. A.21) used the predefined air material. The dielectric regions (Fig. A.22) also used SiO_2 , but with a tunable permittivity parameter. Aluminum was assigned to the superconducting contacts (Fig. A.23) and to the gate boundaries (Fig. A.24), where it was treated as a perfect electric conductor.

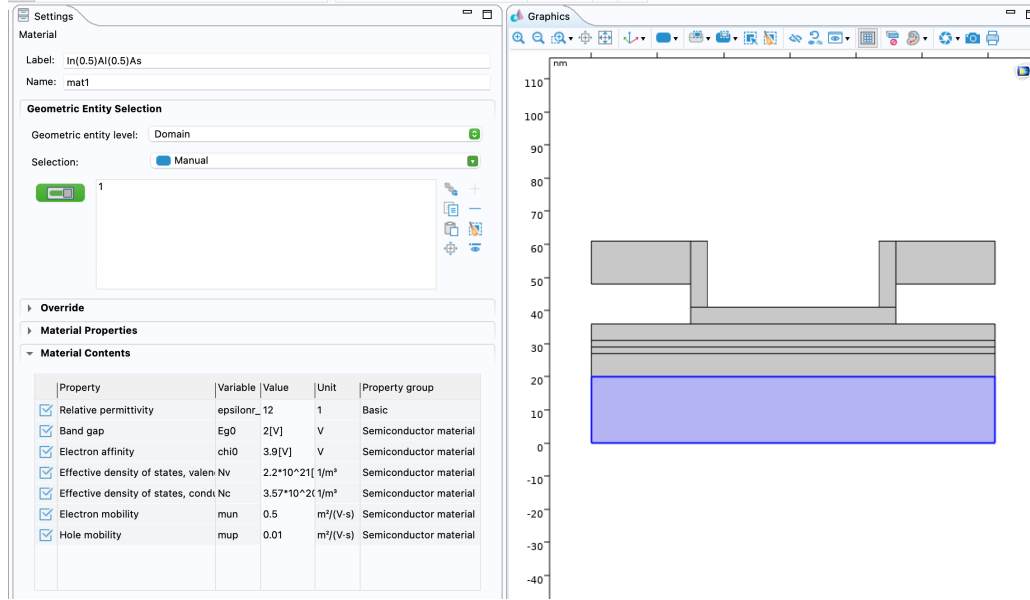


Figure A.17: The bottom layer consisted of $In_{0.5}Al_{0.5}As$.

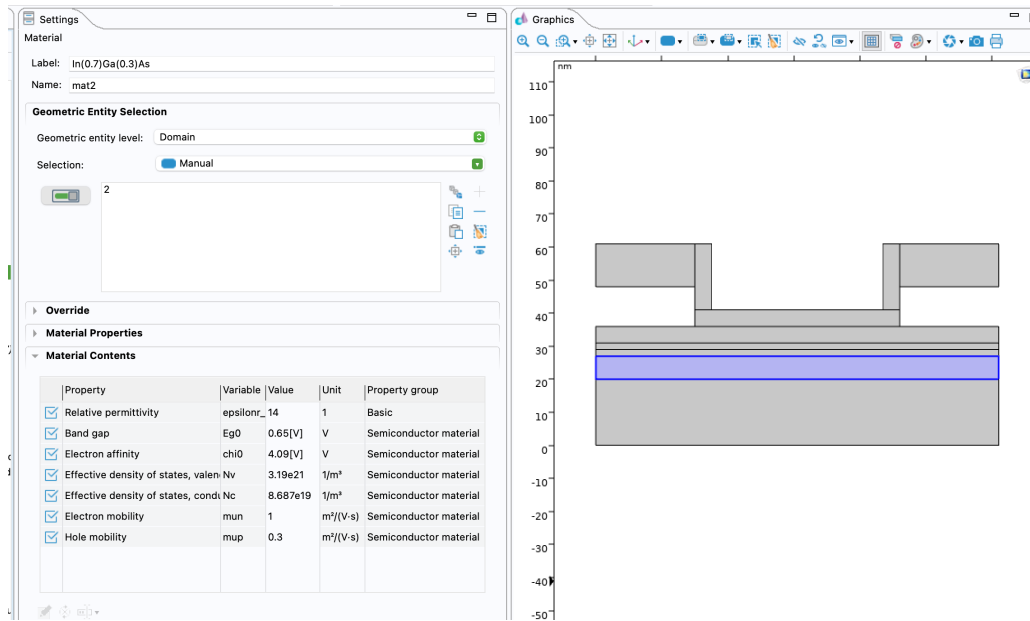


Figure A.18: The second layer consisted of $In_{0.7}Ga_{0.3}As$.

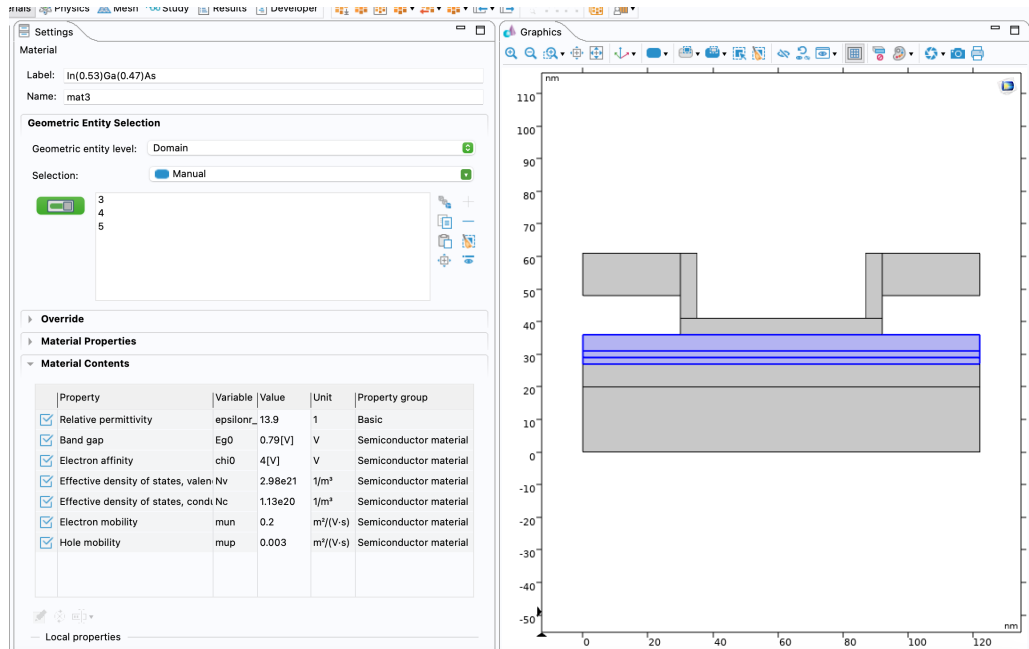


Figure A.19: The third, fourth and fifth layer consisted of $\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$.

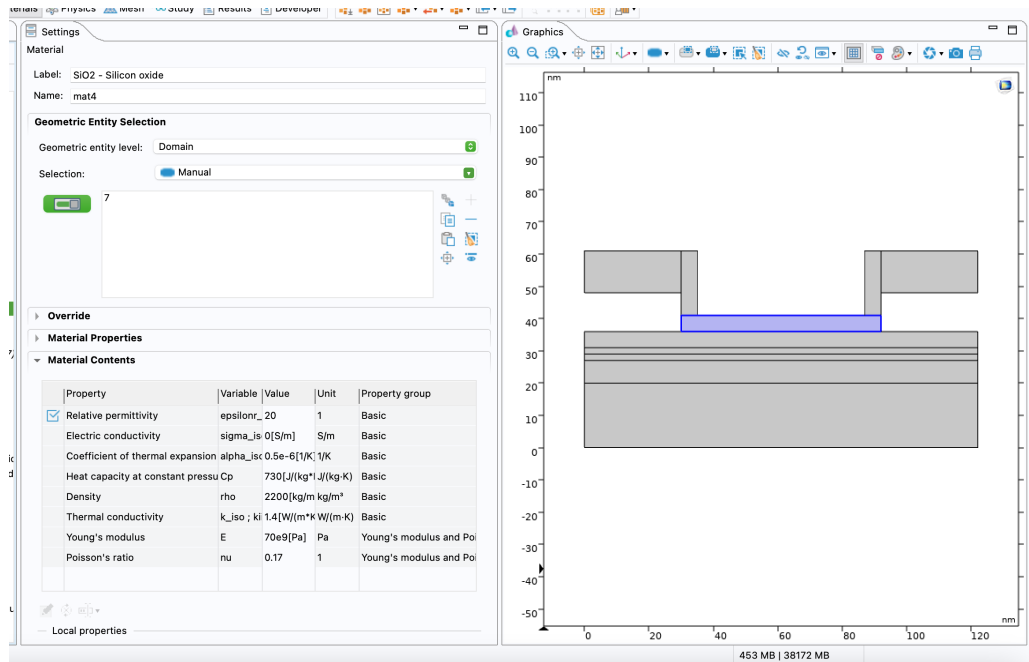


Figure A.20: The oxide layer consisting of SiO_2 .

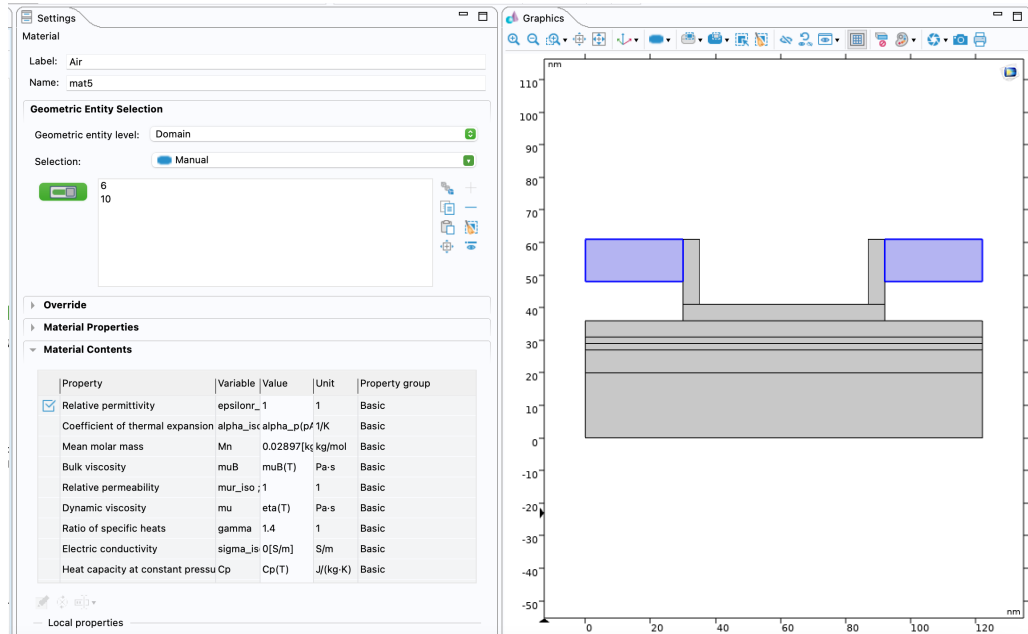


Figure A.21: Two rectangles consisting of air.

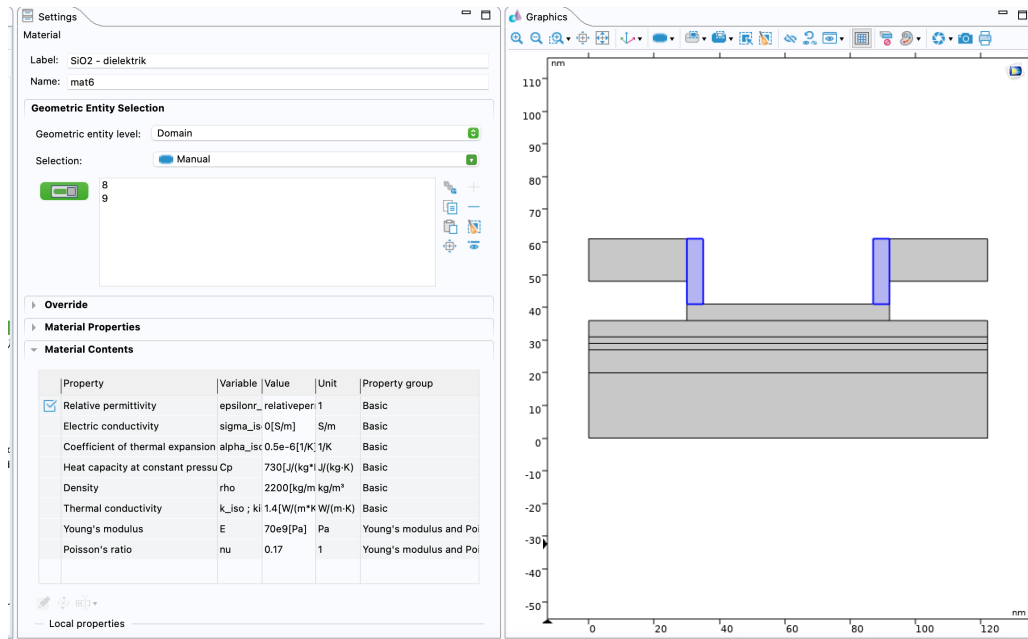


Figure A.22: SiO₂ dielectric domains.

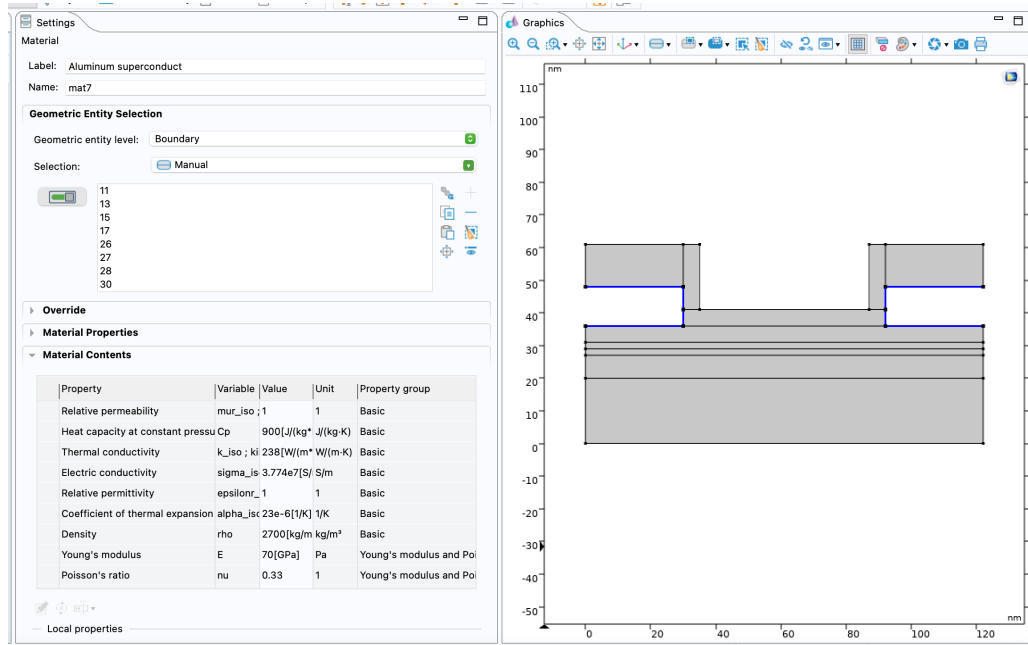


Figure A.23: Aluminum assigned to boundaries to create superconducting contacts.

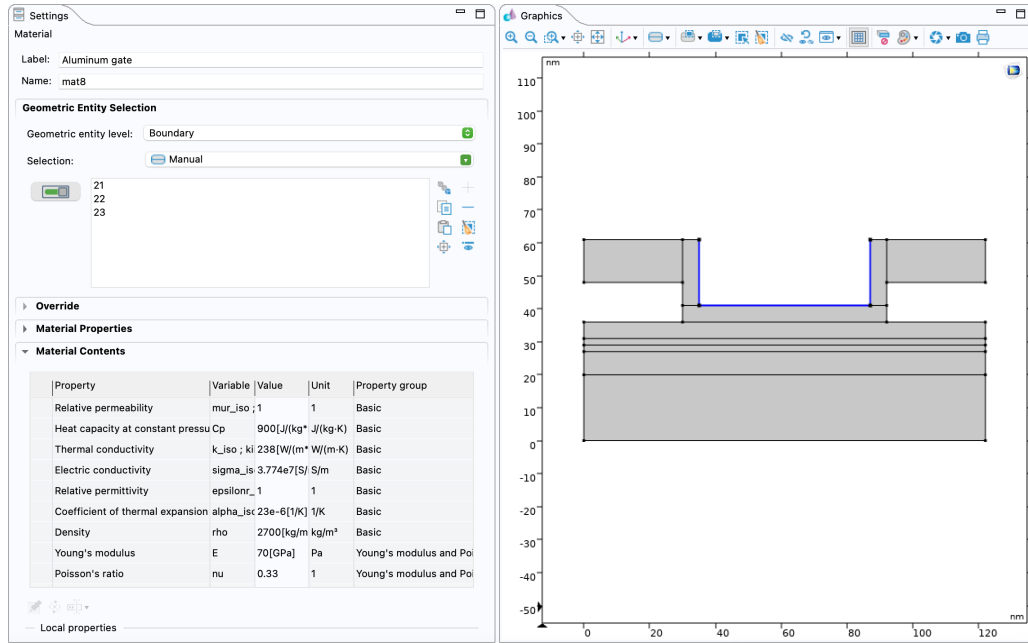


Figure A.24: Aluminum added to boundaries to create the gate.

A.4 Boundary conditions

The electrical behavior of the structure was defined using the boundary conditions shown in Fig. A.25–A.28.

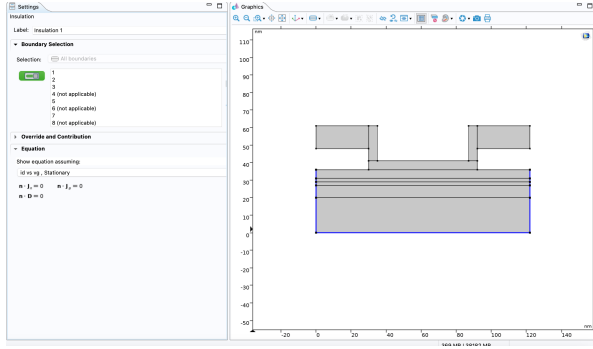


Figure A.25: Insulation applied to the outer air interfaces.

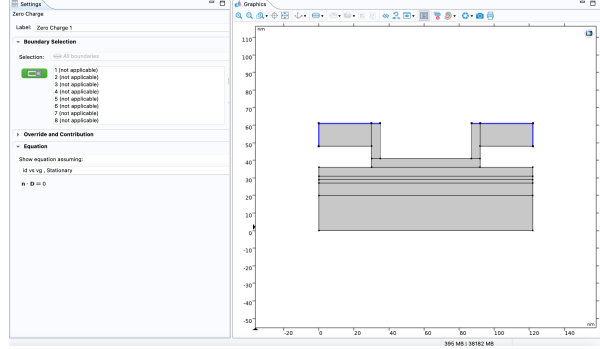


Figure A.26: Zero-charge condition applied to the outer air interfaces.

Fig. A.25 and A.26 show the insulation and zero-charge boundaries applied to the outer air interfaces. The insulation condition prevents current from leaving the model, while the zero-charge condition enforces zero normal electric displacement.

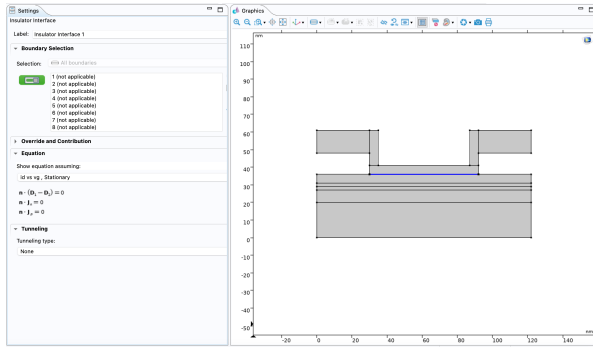


Figure A.27: Insulator interface applied to oxide-semiconductor boundaries.

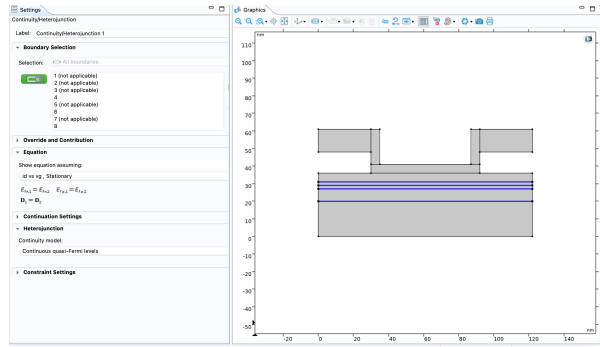


Figure A.28: Continuity/hetero junction interface applied to interior semiconductor boundaries.

The insulator interface condition in Fig. A.27 was applied to the oxide-semiconductor boundaries to prevent carrier flow across the interface. Fig. A.28 shows the boundaries where the continuity/hetero junction condition was applied.

A.5 The physics of the model

The main settings for the semiconductor interface are shown in Fig. A.29. Maxwell-Boltzmann carrier statistics were used, and both electrons and holes were included.

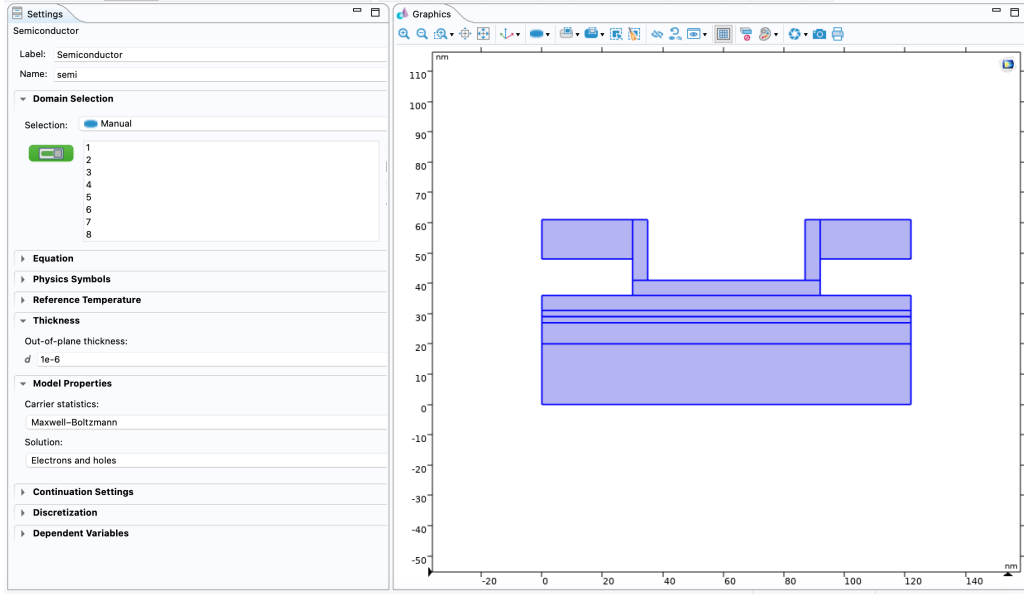


Figure A.29: Overview of the main settings under the semiconductor node.

The semiconductor material model, shown in Fig. A.30, was applied to all semiconductor domains.

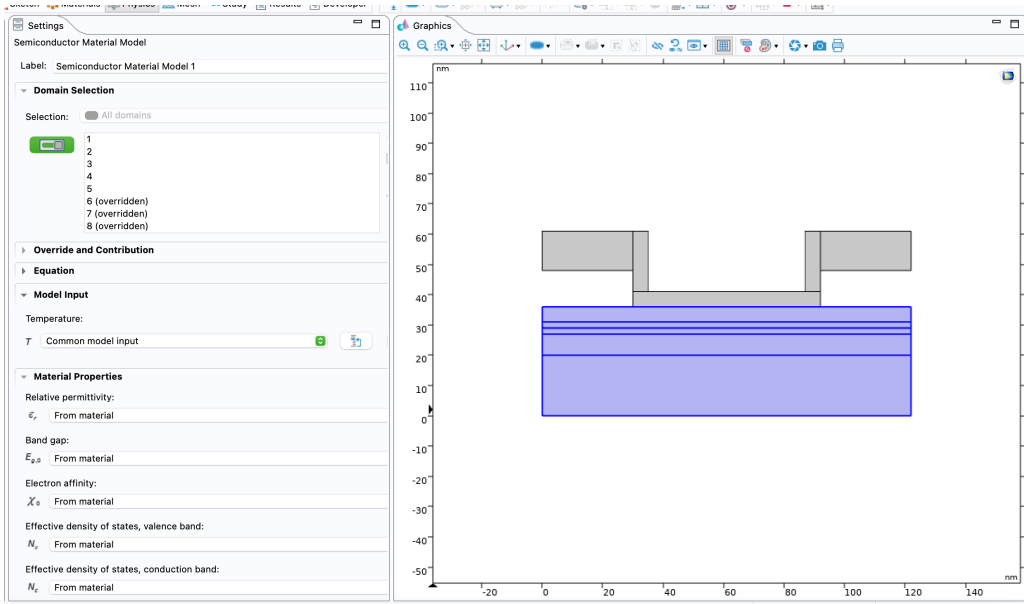


Figure A.30: The semiconductor material model applied to all semiconductor domains.

Initial values for the electric potential and carrier concentrations were assigned as shown in Fig. A.31. These values were chosen to help the solver converge to the stationary solution and reflect the doping polarity.

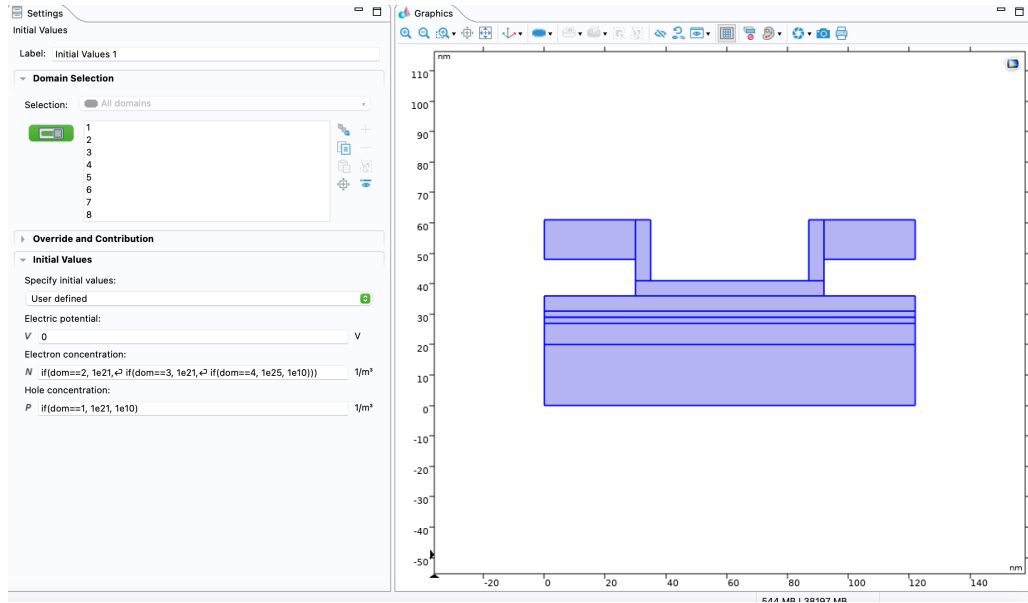


Figure A.31: Initial values for the electric potential and carrier concentrations.

The charge conservation settings can be seen in Fig. A.32, and these were applied to the oxide, dielectric and air domains. This node specifies how the electric displacement field is related to the electric field through the material's constitutive relation. In this model, the temperature was manually set to 50 K to achieve a more realistic material response while still ensuring stable convergence of the stationary solver.

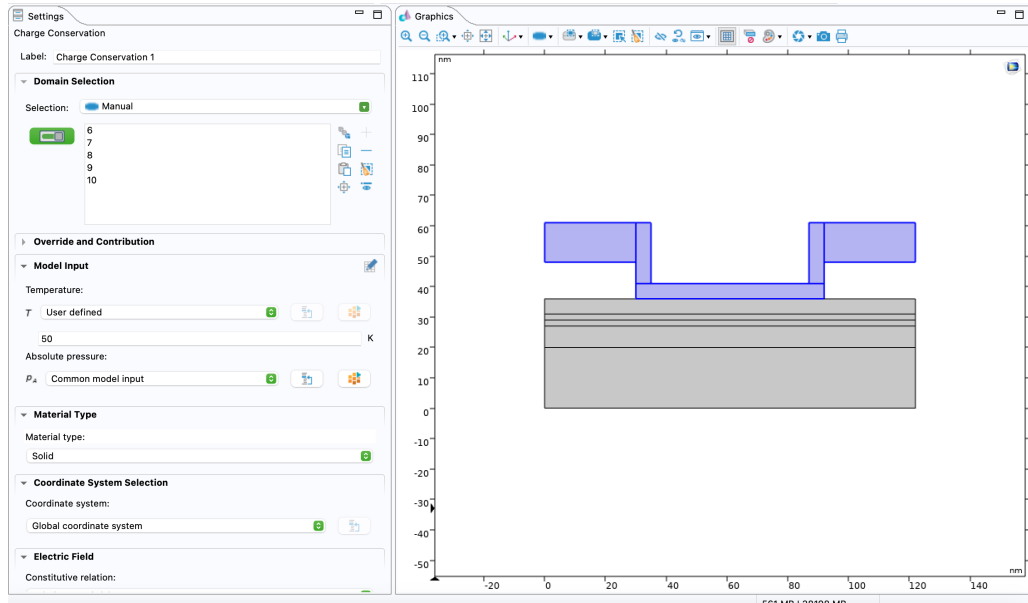


Figure A.32: Charge conservation applied to the oxide, dielectric and air domains.

Doping profiles were specified under the semiconductor nodes. Fig. A.33–A.36 show the acceptor and donor concentrations assigned to each doped semiconductor domain.

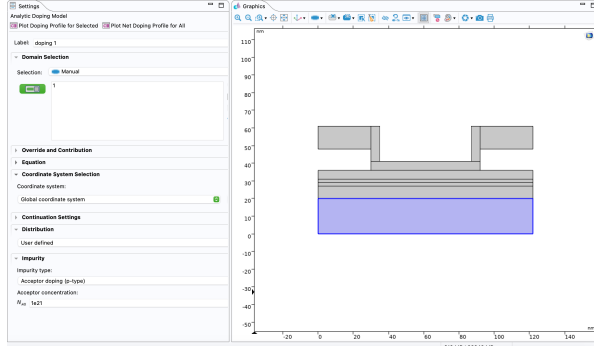


Figure A.33: The bottom semiconductor layer with p -doping of 10^{21} m^{-3} .

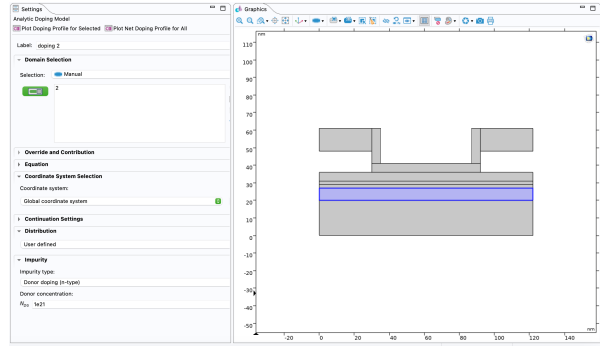


Figure A.34: The second semiconductor layer with n -doping of 10^{21} m^{-3} .

The first domain was p -doped with an acceptor concentration of 10^{21} m^{-3} , and the second domain was n -doped with a donor concentration of 10^{21} m^{-3} .

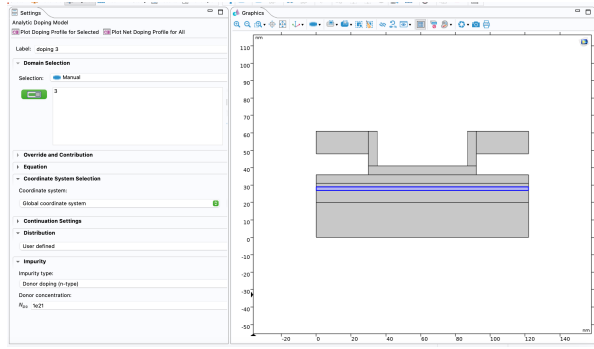


Figure A.35: The third semiconductor layer with n -doping of 10^{21} m^{-3} .

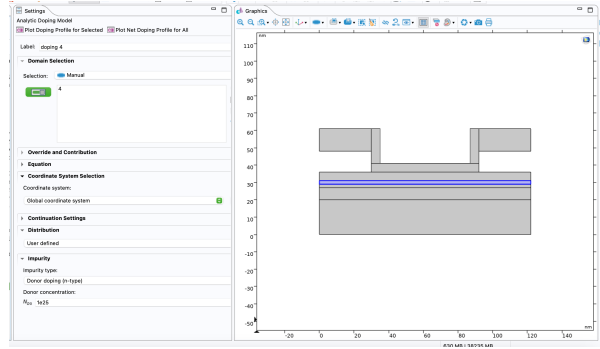


Figure A.36: The fourth semiconductor layer with n -doping of 10^{25} m^{-3} .

Following the second layer, the next layer also had a donor concentration of 10^{21} m^{-3} , and the fourth layer had a donor concentration of 10^{25} m^{-3} . The uppermost semiconductor layer was the spacer and therefore undoped [27].

A.6 Meshing

With the physics settings defined, an additional step before creating a study and computation is required: the meshing. The mesh divides the structure into many small elements and is used to represent the geometry with a finite number of elements. During this project, the mesh was adjusted for certain parameter sweeps when needed for convergence. An overview of a standard mesh used for this model can be seen in Fig. A.37, and the finer mesh applied to the dielectric and oxide domains to improve convergence can be seen in Fig. A.38. A coarser mesh was used for most of the model, while a finer mesh was applied to sensitive regions.

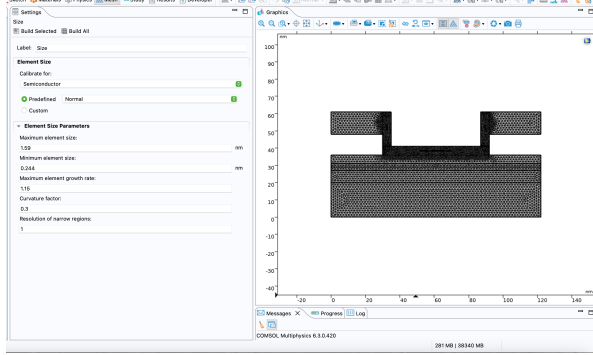


Figure A.37: An overview of the meshing of the model.

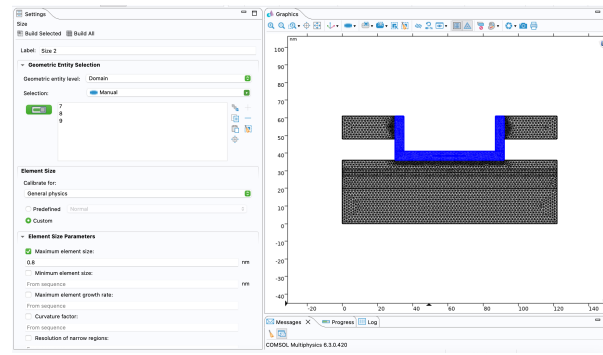


Figure A.38: Finer mesh applied to the dielectric and oxide domains.

A.7 An example of a study used

To finish this learning appendix, an example of how to set up a study with the correct settings is presented. This study consisted of three components: a stationary step, a frequency-domain perturbation step, and a parametric sweep. The stationary step, seen in Fig. A.39, had no modifications from the default settings. It is good practice to run the stationary step first to ensure convergence.

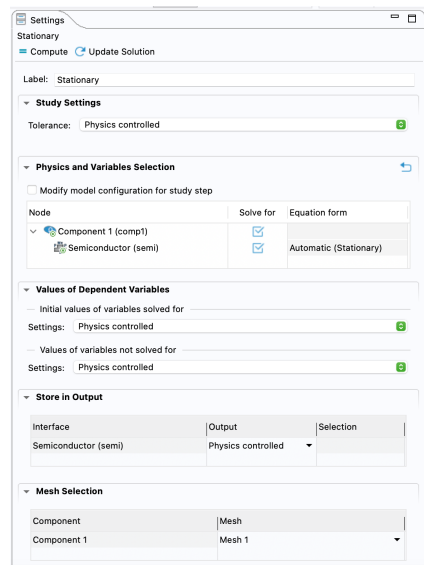


Figure A.39: The stationary step settings in a study.

Following this, the frequency-domain perturbation step seen in Fig. A.40 was added. Two important settings were changed from the default ones:

- “Reuse solution from previous step” was set to “Yes” to improve convergence.
- “Initial values of variables solved for” was changed from “Physics controlled” to “User controlled”, with the stationary step set as the initial solution.

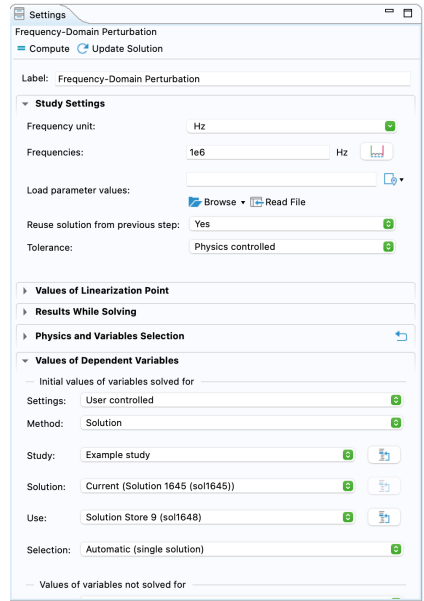


Figure A.40: Frequency-domain perturbation step settings.

After running the study with the two previous steps and checking convergence, a parametric sweep was added. Here, parameters can be swept over ranges or specific values. To ensure convergence, it is best to start with realistic values and sweep one parameter at a time. An example is shown in Fig. A.41.

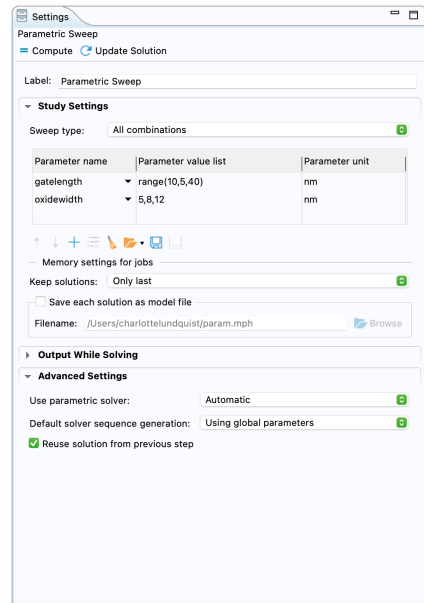


Figure A.41: An example of a parametric sweep where one parameter is swept from 10-40 nm in steps of 5 nm, while another is swept for the specific values 5, 8 and 12 nm.

The settings changed from the default were enabling “Reuse solution from previous step” and changing “Keep solutions” from “All” to “Only last”. This improves convergence and reduces computation time.

Having learned the basics of creating a model in COMSOL, the next section explains how to use plots from COMSOL in MATLAB to calculate mobility and transport times.

B MATLAB scripts

The following scripts were generated with help from Microsoft Copilot.

The MATLAB scripts in this appendix were used to process the quantum-mechanical output from COMSOL and to compute all transport-related quantities presented in the main text. The scripts take the exported wavefunctions, eigenenergies and position grids as input, and perform the following tasks:

- compute the form factor $F(q)$ using a double integral over the wavefunction,
- evaluate the screened Coulomb potential $\tilde{U}(q)$ using the Thomas–Fermi expression,
- calculate the transport relaxation time and mobility for each gate voltage,
- extract the gate-dependent sheet density $n_{2D}(V_g)$ directly from the wavefunction,
- compute the differential capacitance $C(V_g) = e dn_{2D}/dV_g$,
- generate the mobility–capacitance relationship used in the analysis.

All numerical parameters, physical constants and integration routines match those described in the main text. The full MATLAB code is included here for reproducibility and to allow further development of the transport model.

B.1 compute_Fq_double.m

```
function [Fq, q] = compute_Fq_double(z, psi2, z_dope_min, z_dope_max, qmax, Nq)
% F(q) = \int_{z_dope_min}^{z_dope_max} \int_{z_dope_min}^{z_dope_max} |psi(z)|^2 |psi(z')|^2 e^{-q|z-z'|} dz dz'
% Input:
%   z           - position in meters
%   psi2        - |psi(z)|^2, normalized so that \int psi2 dz = 1
%   z_dope_min  - lower limit for doping profile (m)
%   z_dope_max  - upper limit for doping profile (m)
%   qmax        - maximum value of q (m^-1)
%   Nq          - number of q-points
%
% Output:
%   Fq          - dimensionless form factor F(q)
%   q           - q-values (m^-1)

if nargin < 6
    Nq = 400;
end

z = z(:);
psi2 = psi2(:);

if length(z) ~= length(psi2)
    error('z and psi2 has to have the same length.');
```

```

dz = z(2) - z(1);

q = linspace(0, qmax, Nq).';
q(1) = 1e-6 * qmax;

z_prime = z(z >= z_dope_min & z <= z_dope_max);
dzp = z_prime(2) - z_prime(1);

[Z, Zp] = meshgrid(z, z_prime); % Z: z, Zp: z'
abs_diff = abs(Z - Zp);        % |z - z'|

Fq = zeros(size(q));

for i = 1:length(q)
    qi = q(i);
    kernel = exp(-qi * abs_diff); % e^{-q/|z - z'|}
    inner = sum(kernel, 1) * dzp; % dz' e^{-q/|z - z'|}
    integrand = psi2' .* inner; % |psi(z)|^2 * inner integral
    Fq(i) = sum(integrand) * dz; % dz ...
end
end

```

B.2 main_mobility.m

```

% main_mobility.m
% =====

clear; clc;

%% 0. Universal and material constants
hbar = 1.054e-34; % [J*s]
e = 1.602e-19; % [C]
eps0 = 8.854e-12; % [F/m]
eps_r = 13.5; % effective relative permittivity
m0 = 9.109e-31; % [kg]

%% 1. Load wavefunction from COMSOL
z = load('psix_correct.txt'); % nm
psi = load('psinorm_correct.txt');

z = z * 1e-9; % convert to meter

% Normalize |psi|^2 correctly
dz = z(2) - z(1);
psi2 = abs(psi).^2;
psi2 = psi2 ./ (sum(psi2) * dz);

%% 2. Electron density n2D from COMSOL
n2D = 7.2e12; % [1/m^2]

%% 3. Effective mass m_eff from m*(z)
m_well = 0.06 * m0;
m_barrier = 0.09 * m0;

```

```

% New well domains
z_oxide_max = 10e-9;      % undoped/oxide up to 10 nm
z_well_min  = 20e-9;      % well: 20-28 nm
z_well_max  = 28e-9;
z_well_center = (z_well_min + z_well_max)/2;

m_z = zeros(size(z));
m_z(z >= 0 & z <= z_oxide_max) = m_barrier;
m_z(z > z_oxide_max & z < z_well_min) = m_barrier;
m_z(z >= z_well_min & z <= z_well_max) = m_well;
m_z(z > z_well_max) = m_barrier;
m_z(m_z == 0) = m_barrier;

inv_m_eff = trapz(z, psi2 ./ m_z);
m_eff      = 1 / inv_m_eff;

fprintf('Effective mass m_eff/m0 = %.4f\n', m_eff/m0);

%% 4. Fermi-wave number and q-grid
kF  = sqrt(2*pi*n2D);
qmax = 3 * kF;
Nq   = 400;

%% 5. Doping: barriers and well
% Left barrier: 10-20 nm, ND = 1e25 m^-3
z_top_min = 10e-9;
z_top_max = 20e-9;
d_dope_top = z_top_max - z_top_min;
z_top_center = (z_top_min + z_top_max)/2;
d_top        = abs(z_well_center - z_top_center);
Nimp_top     = 1e25 * d_dope_top;      % [1/m^2]

% Well: 20-28 nm, ND = 1e21 m^-3
d_well      = z_well_max - z_well_min;
Nimp_well   = 1e21 * d_well;          % [1/m^2]

% Right barrier: 78-128 nm, NA = 1e21 m^-3
z_bot_min = 78e-9;
z_bot_max = 128e-9;
d_dope_bot = z_bot_max - z_bot_min;
z_bot_center = (z_bot_min + z_bot_max)/2;
d_bot        = abs(z_bot_center - z_well_center);
Nimp_bot     = 1e21 * d_dope_bot;     % [1/m^2]

fprintf('Nimp_top  = %.2e 1/m^2, d_top  = %.2f nm\n', Nimp_top, d_top*1e9);
fprintf('Nimp_bot   = %.2e 1/m^2, d_bot  = %.2f nm\n', Nimp_bot, d_bot*1e9);
fprintf('Nimp_well = %.2e 1/m^2 (in-well)\n', Nimp_well);

%% 6. Screening
qTF = 2 * m_eff * e^2 / (eps_r * eps0 * hbar^2);
fprintf('qTF = %.3e 1/m\n', qTF);

%% 7. Form factor with double integral (dimensionless)

```

```

[Fq_top, q] = compute_Fq_double(z, psi2, z_top_min, z_top_max, qmax, Nq);
[Fq_bot, ~] = compute_Fq_double(z, psi2, z_bot_min, z_bot_max, qmax, Nq);
%% 7b. Plot form factors (Left/Right barrier)
figure;
plot(q, Fq_top, 'LineWidth', 2); hold on;
plot(q, Fq_bot, 'LineWidth', 2);
set(gca, 'YScale', 'log'); % logarithmic y-axis
grid on;
xlabel('q (1/m)');
ylabel('F(q)');
title('Form factor F(q)');
legend('Left barrier', 'Right barrier', 'Location', 'northeast');

%% 8. Coulombpotentialer (without Nimp)
Uq_top = (e^2 ./ (2 * eps0 * eps_r)) .* Fq_top ./ (q + qTF);
Uq_bot = (e^2 ./ (2 * eps0 * eps_r)) .* Fq_bot ./ (q + qTF);
Uq_well = (e^2 ./ (2 * eps0 * eps_r)) ./ (q + qTF); % F(q) 1

U_top = @(qq) interp1(q, Uq_top, qq, 'pchip', 0);
U_bot = @(qq) interp1(q, Uq_bot, qq, 'pchip', 0);
U_well = @(qq) interp1(q, Uq_well, qq, 'pchip', 0);

%% 9. Angular integral (Nimp is used here)
Ntheta = 1000;
theta = linspace(0, 2*pi, Ntheta);
q_theta = 2 * kF .* sin(theta/2);

U_theta_top = U_top(q_theta);
U_theta_bot = U_bot(q_theta);
U_theta_well = U_well(q_theta);

integrand = ...
    Nimp_top * abs(U_theta_top).^2 .* (1 - cos(theta)) + ...
    Nimp_bot * abs(U_theta_bot).^2 .* (1 - cos(theta)) + ...
    Nimp_well * abs(U_theta_well).^2 .* (1 - cos(theta));

I_theta = trapz(theta, integrand);
tau_inv = (m_eff / (2 * pi * hbar^3)) * I_theta;
tau_tr = 1 / tau_inv;
mu = e * tau_tr / m_eff;

%% 10. Results
fprintf('\nTransport time tau_tr = %.3e s\n', tau_tr);
fprintf('Mobility: %.3e m^2/Vs (%.1f cm^2/Vs)\n', mu, mu * 1e4);

```

B.3 mobility_run_energies.m

```

%% === Vg and eigenvalues (MODE 1) ===

Vg_values = [-0.25 -0.20 -0.15 -0.10 -0.05 0 0.05 0.10 0.15 0.20 0.25];

E = [
    7.9993; % -0.25
    7.9662; % -0.20

```

```

7.9334;    % -0.15
7.9015;    % -0.10
7.8712;    % -0.05
7.8375;    %  0.00
7.7951;    % +0.05
7.7533;    % +0.10
7.7121;    % +0.15
7.6711;    % +0.20
7.6299     % +0.25
];

%% === Find all psi-files ===
files = dir('1psi*.txt');

data = struct('z',{}, 'psi',{}, 'E',{}, 'Vg',{}, 'n2D',{});

Vg_list  = zeros(length(files),1);
mu_list  = zeros(length(files),1);
n2D_list = zeros(length(files),1);

%% === Loop over all files ===

for i = 1:length(files)

    fname = files(i).name;

    % Match Vg with file name
    tokens = regexp(fname, '1psi(minus)?(\d+)\.txt', 'tokens');

    if strcmp(tokens{1}{1}, 'minus')
        rawVg = -str2double(tokens{1}{2});
    else
        rawVg = str2double(tokens{1}{2});
    end

    Vg = rawVg / 100;

    idx = find(abs(Vg_values - Vg) < 1e-6, 1);

    % MODE 1 energy
    E_mode = E(idx);

    % Read psi-file
    raw = readmatrix(fname);

    if isempty(raw)
        fprintf(" Tom fil: %s\n", fname);
        continue
    end

    z_nm = raw(:,1);
    psi  = raw(:,2);

```

```

valid = ~isnan(psi);
z = z_nm(valid) * 1e-9;
psi = psi(valid);

if isempty(psi) || all(psi == 0)
    fprintf(" Broken psi-file: %s\n", fname);
    continue
end

% === n2D from the wavefunction ===

n2D = trapz(z, psi.^2);

% Spara
Vg_list(i) = Vg;

data(i).z      = z;
data(i).psi     = psi;
data(i).E      = E_mode;
data(i).Vg     = Vg;
data(i).n2D    = n2D;

% Mobilitet
mu_list(i) = compute_mobility_from_main(data(i));
n2D_list(i) = n2D;

end

%% === Sort ===

[Vg_sorted, idx] = sort(Vg_list);
mu_sorted = mu_list(idx);
n2D_sorted = n2D_list(idx);

%% === Capacitance ===

e = 1.602e-19;
dn_dV = gradient(n2D_sorted, Vg_sorted);
C_extracted = e * dn_dV;

%% === Plots ===

figure; plot(Vg_sorted, mu_sorted, '-o'); grid on;
xlabel('Vg'); ylabel('Mobility');

figure; plot(Vg_sorted, n2D_sorted, '-o'); grid on;
xlabel('Vg'); ylabel('n2D');

figure; plot(Vg_sorted, C_extracted, '-o'); grid on;
xlabel('Vg'); ylabel('C');

figure; plot(C_extracted, mu_sorted, '-o'); grid on;
xlabel('C'); ylabel('Mobility');

```