

Course
*Superconducting Quantum Nanosystems,
Circuits and Devices*

Part 2: Circuits, Devices and Experiments

Federica Mantegazzini

f.mantegazzini@unitn.it fmantegazzini@fbk.eu

University of Trento, Physics Department

Academic Year 2025-2026

Note: The course "Superconducting Quantum Nanosystems, Circuits and Devices" is organised in two parts. The first part (teacher: Gianluca Rastelli) is dedicated to theory. The second part (teacher: Federica Mantegazzini) is dedicated to circuits, devices and experiments.

These notes are meant as a summary of the lectures and a helpful tool for the students. They are a living document. Suggestions to improve the content and corrections of typos or possible errors are more than welcome.

Contents

1	Introduction	3
2	Microwave Coplanar Waveguides	4
3	Superconducting Microwave Resonators	6
3.1	Coplanar waveguide resonators	6
3.2	Lumped element resonators	9
3.3	Quality factors	10
3.4	3D superconducting resonators	11
3.5	Representations of resonances	12
4	The Quantum Harmonic LC Oscillator	13
5	Josephson Junctions and Josephson inductance	15
6	Kinetic inductance	18
7	Superconducting devices: an overview	20
8	SQUIDs: Superconducting Quantum Interference Devices	21
8.1	Introduction	21
8.2	The zero-voltage state	21
8.3	The voltage state	25
8.4	SQUID operation	25
8.5	SQUID designs	28
8.6	Applications	28
9	Cryogenic detectors	29
9.1	Superconducting Nanowire Single-Photon Detectors (SNPSDs)	29
9.2	Cryogenic microcalorimeters	29
9.3	Transition Edge Sensors (TESs)	32
9.4	Magnetic Microcalorimeters (MMCs)	34
9.5	Kinetic Inductance Detectors (KIDs)	35
10	Superconducting parametric amplifiers	36
10.1	Introduction	36
10.2	Parametric amplification	36
10.3	Phase-preserving and phase-sensitive amplification	38
10.4	Resonant parametric amplifiers	40
10.5	Travelling Wave Parametric Amplifiers (TWPAs)	41
11	Superconducting qubits	42

1 Introduction

Superconducting circuits have attracted growing interest in the scientific community as a tool to implement quantum optics experiments in the microwave regime as well as a platform for quantum technologies.

Examples of applications of superconducting circuits include:

- *Quantum optics and circuit QED (cQED) experiments*, exploring the fundamental interactions between microwave light and matter;
- *High-resolution cryogenic detectors*, widely used for particle physics experiments (e.g. dark matter and neutrino physics) and astroparticle physics applications (e.g. for telescopes).
- *Quantum sensing*, exploiting the quantum nature of superconductivity to build sensors based on quantum tunnelling, superposition and entanglement phenomena.
- *Quantum computing*, realising superconducting artificial atoms that work as qubits, where the first two energy levels $|0\rangle$ and $|1\rangle$ are used as computational basis.

The key advantage of superconducting circuits as experimental platform for quantum applications is the fact that superconductivity is a *macroscopic quantum phenomenon*. This implies that relatively large, i.e. "macroscopic" superconducting structures, with sizes in the order of tens or hundreds of micrometers, preserve their quantum nature. This feature allows to design and produce circuits, with incredibly high flexibility, to control and exploits quantum phenomena. One example is the transmon qubit, a superconducting circuit widely used as basic building block for superconducting quantum processors, which we will discuss in section ??.

Quantum states are typically very fragile and the environmental disturbances, such as thermal noise, can be detrimental, easily destroying quantum coherence. Superconducting devices are operated at very low temperatures, typically in the order of millikelvin, and in the microwave regime, at frequencies of few GHz. Therefore, it results that the thermal noise $k_B T$ is much lower than the energy scale $\hbar\omega$, with $\omega = 2\pi f$ and $f \approx 5GHz$. This condition:

$$k_B T \ll \hbar\omega \tag{1}$$

is in fact crucial to allow the exploitation of quantum circuits to build qubits and other quantum devices.

Other advantages of the superconducting platform include:

- **Strong light-matter coupling:** In cQED, the interaction between photons and atoms is given by the ratio between the coupling strength and the bare energy of the excitations; with superconducting circuits, very high level of interactions can be achieved, enabling experiments in the strong and ultra-strong regimes;
- **The (quasi) absence of dissipation in superconducting circuits:** Superconductivity allows zero-resistance current flow in the circuit, which brings the dissipation close to zero;
- **Scalability:** Superconducting circuits are typically microfabricated on silicon chips, exploiting cleanroom techniques readily developed by the semiconductor industry, allowing for high control and high scalability.

2 Microwave Coplanar Waveguides

In order to build superconducting circuits, we need to be able to transmit microwaves, i.e. electromagnetic waves with frequencies in the order of few GHz, in a controlled way.

In the 3D world, we use coaxial cables to transmit microwaves. A coaxial cable comprises an inner conductor and an outer ground separated by dielectric material, as shown in figure 1.

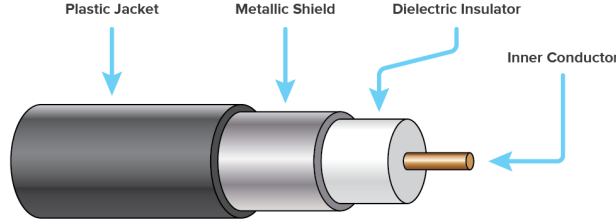


Figure 1: A coaxial cable, comprising the inner conductor, the dielectric insulator, the shield which serves as ground and a plastic jacket which serves as protection.

In a superconducting planar circuit, structures must feature a planar geometry and they are realised exploiting planar microfabrication techniques. Therefore, we need to implement a planar version of a coaxial cable. The simplest planar implementation of a transmission line is the *coplanar waveguide* shown in figure 2. A coplanar waveguide consists of a planar central conductor and a planar ground plane separated from the central conductor by a gap. In a typical superconducting circuit, a superconducting thin film (e.g. an aluminium film with a thickness in the order of 100 nm) is deposited on a silicon substrate and patterned by removing the film from the gap region.

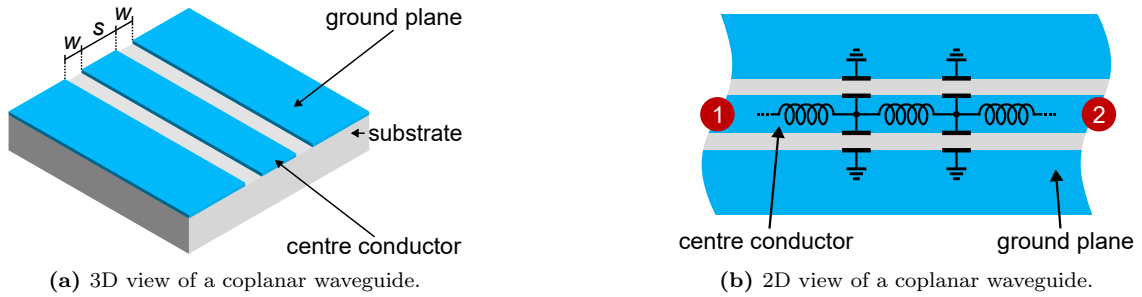


Figure 2: Coplanar waveguide comprising the central conductor and the ground plane.

A coplanar waveguide is characterised by an inductance, given by the geometrical inductance of the central conductor, and by a capacitance, given by the facing metallic planes of the central conductor and the ground.

Therefore, we can define two important parameters of a coplanar waveguide:

- The distributed inductance L' , i.e. the inductance per unit of length;
- The distributed capacitance C' , i.e. the capacitance per unit of length;

To calculate L' and C' , *conformal mapping* techniques are typically exploited. These are algebraic methods which are used to map a complex geometry, such as a coplanar waveguide, to a simpler and known geometry, such as a parallel plate capacitor. In this way it is possible to calculate in a rigorous way L' and C' for a given geometry of a coplanar waveguide. The free parameters for each geometry are the width of the central conductor and the width of the gap between the central conductor and the ground plane.

We can define four geometry dependent quantities:

$$\begin{aligned} k &= \frac{s}{s+2w}, \\ k' &= \sqrt{1-k^2}, \\ k_x &= u_1(x)/u_2(x), \\ k'(x) &= \sqrt{1-k_x^2} \end{aligned} \tag{2}$$

where

$$\begin{aligned} u_1(x) &= \frac{s}{2} + \frac{x}{\pi} \left[1 + \ln \left(\frac{2\pi s}{x} \frac{w}{s+w} \right) \right] \\ u_2(x) &= \frac{s}{2k} - \frac{x}{\pi} \left[1 + \ln \left(\frac{2\pi s}{xk} \frac{w}{s+w} \right) \right]. \end{aligned} \tag{3}$$

Subsequently, we can calculate the inductance per unit length L' and the capacitance per unit length C' as

$$\begin{aligned} L' &= \frac{\mu_0}{4} \frac{K(k'_{x=t/2})}{K(k_{x=t/2})} \\ C' &= 2\epsilon_0 \frac{K(k_{x=t})}{K(k'_{x=t})} + 2\epsilon_0 \epsilon_r \frac{K(k)}{K(k')} \end{aligned} \tag{4}$$

where ϵ_0 is the vacuum permittivity and $K(x)$ is the complete elliptic integral of the first kind.

From the distributed inductance L' and the capacitance C' , we can calculate the characteristic impedance Z_0 of the waveguide:

$$Z_0 = \sqrt{\frac{L'}{C'}} \tag{5}$$

and the phase velocity v_{ph} of the microwave travelling inside the waveguide:

$$v_{\text{ph}} = \frac{1}{\sqrt{L'C'}} \tag{6}$$

We can visualise the microwave travelling inside the waveguide looking at the electromagnetic field distribution, with the electric component $\vec{E}$ and the magnetic component the electric component $\vec{B}$, as shown in figure 3.

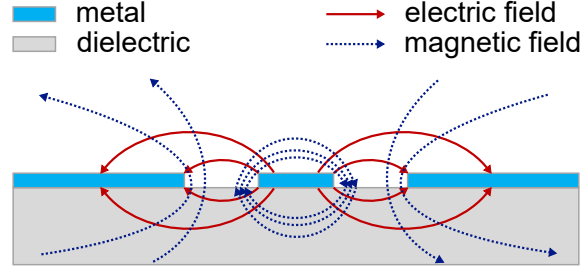


Figure 3: Visualisation of the field distribution of a quasi-TEM mode travelling in a coplanar waveguide.

3 Superconducting Microwave Resonators

Starting from coplanar waveguides, we can build planar microwave cavities. To do this, we can exploit interference and superposition effects of microwaves travelling in a waveguide with a certain boundary conditions.

In figure 4, we consider the case of a medium terminated with a closed end, equivalent to a coplanar waveguide terminated to ground. This configuration causes reflection of the electric component of the wave with a phase shift of π . A standing wave in a medium is a resonator: we have built a superconducting microwave resonator.

In the animation¹, we can see how an incoming wave (green) travelling from left to right in a medium is bouncing back when reflected at the end of the medium (right end of the figure). The reflected way (blue) is interfering with the incoming way (green) and the resulting superposition of the two gives a standing wave (red).

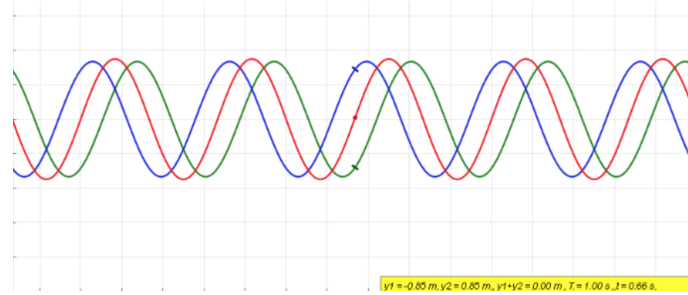


Figure 4: Animation here: <https://upload.wikimedia.org/wikipedia/commons/5/5d/Waventerference.gif>. The incoming wave (green) is reflected and the back-travelling wave (blue) is interfering resulting in the standing wave (red).

3.1 Coplanar waveguide resonators

Different boundary conditions in a transmission line result in different phase shifts and therefore different standing waves. In particular, we can think to two different configurations, namely the *quarter-wave resonator* and the *half-wave resonator*.

Quarter-wave resonator

A coplanar waveguide is interrupted at one end by a capacitor, which prevents current flowing, and is terminated to ground at the other end, allowing maximum current flow, as shown in figure 5a. We therefore have a current node (i.e. zero current, maximum voltage) at one end and a voltage node (i.e. maximum current, zero voltage) at the other end. The electrical components of the microwave is forced to be zero at one end and maximum at the other end. The resulting standing wave is characterised by a wavelength λ which corresponds to four times the length of the

¹<https://upload.wikimedia.org/wikipedia/commons/5/5d/Waventerference.gif>

resonator l : $\lambda = 4l$. In other words, a quarter of the wave fits in the resonator, hence generating a *quarter-wave* resonator.

The resonance frequency of a quarter-wave resonator is given by:

$$f_0 = \frac{\omega_0}{2\pi} = \frac{1}{4l_r\sqrt{L'C'}} \quad (7)$$

where the mode number is omitted for simplicity.

By characterising superconducting resonators with a Vector Network Analyser (VNA), their resonance frequencies can be extracted. A VNA is an RF instrument with two ports, which generates a microwave signal at port i and it compares in terms of amplitude and phase with the signal detected at port j , where $i, j \in \{1, 2\}$. The VNA returns the so called *scatter parameters* S_{ij} . The parameters S_{11} and S_{22} describe the forward and backwards reflections, while S_{21} and S_{12} describe the forward and backwards transmissions.

Using a VNA, we can probe a **quarter-wave resonator** with a reflection measurement, sending a signal to the port 1 of the device (see figure 5a) and sweeping the frequency of the signal. The plot in figure 6a shows the measured reflection as a function of the signal frequency:

- When the signal frequency does not match the resonance frequency, the reflection is total (i.e. at the maximum of 0 dB = 100% reflection). This happens because the sent microwave cannot fit in the quarter-wave resonator and therefore is entirely reflected back to port 1 of the device.
- When the signal frequency matches the resonance frequency, the reflection is at the minimum, generating a dip in the reflection graph. Under this condition, the microwave can enter the resonator and it is trapped there, so that the reflection is minimal.

The reflection is measured in logarithmic units of dB.

Half-wave resonator

A coplanar waveguide is interrupted at the two ends by two capacitors, which prevent flowing of current, as shown in figure 5b.. We therefore have two current nodes (i.e. zero current, maximum voltage) at the two ends, forcing the electrical components of the microwave to be zero. The resulting standing wave is characterised by a wavelength λ which corresponds to twice the length of the resonator l : $\lambda = 2l$. In other words, half wave fits in the resonator, hence generating a *half-wave* resonator.

The resonance frequency of a half-wave resonator is given by:

$$f_0 = \frac{1}{2l_r\sqrt{L'C'}} \quad (8)$$

where the mode number is omitted for simplicity.

We can probe a **half-wave resonator** with a transmission measurement, sending a signal to the port 1 of the device and reading the transmission at port 2 of the device (see figure 5b, while sweeping the frequency of the signal. The plot in figure 6b) shows the measured transmission as a function of the signal frequency:

- When the signal frequency does not match the resonance frequency, the transmission is very low (ideally $-\infty$). This happens because the sent microwave cannot fit in the half-wave resonator and therefore is entirely reflected back to port 1 of the device.
- When the signal frequency matches the resonance frequency, the signal can enter the resonator and be transmitted to the other side, generating a peak in the transmission graph. Under this condition, the microwave is trapped inside the resonator and subsequently transmitted both backwards (to port 1 of the device) and forward (to port 2 of the device).

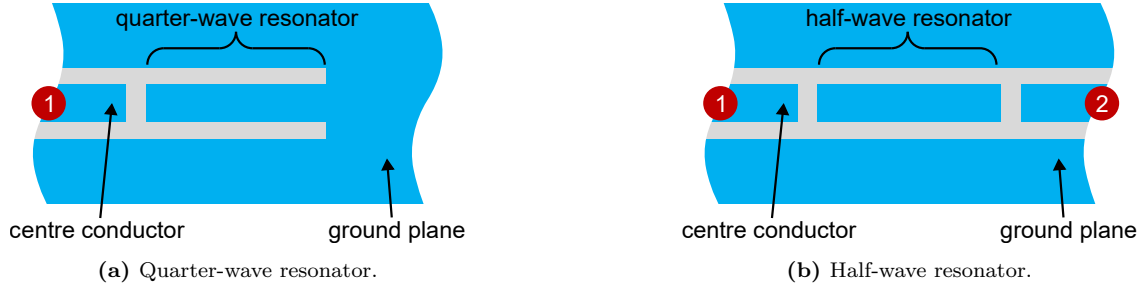


Figure 5: Planar resonators with different boundary conditions. The ports used to probe the resonators with a microwave tone are shown in red.

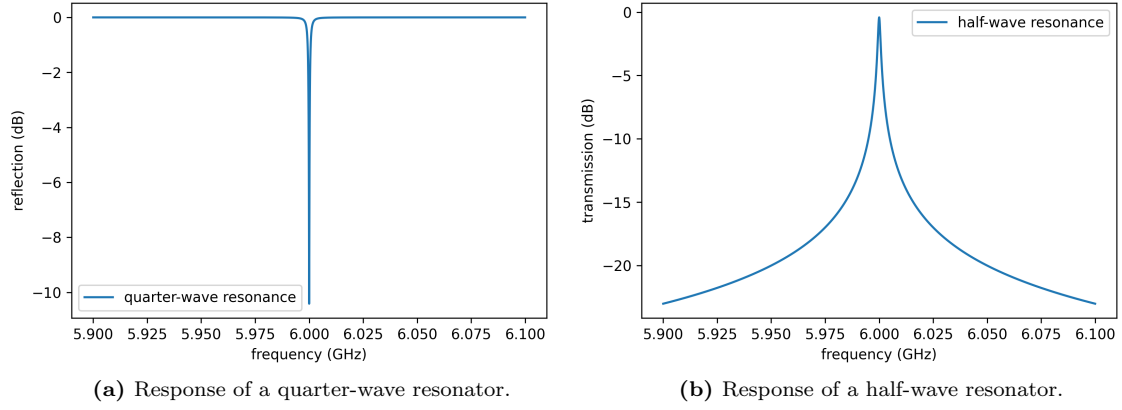


Figure 6: Microwave responses of superconducting resonators, showing the resonance frequency.

The transmission is measured in logarithmic units of dB.

In figure 7 the quarter-wave (left) half-wave (right) configurations are visualised, including the allowed harmonics. Note that for the quarter-wave resonator only the odd harmonics are allowed by the boundary conditions.

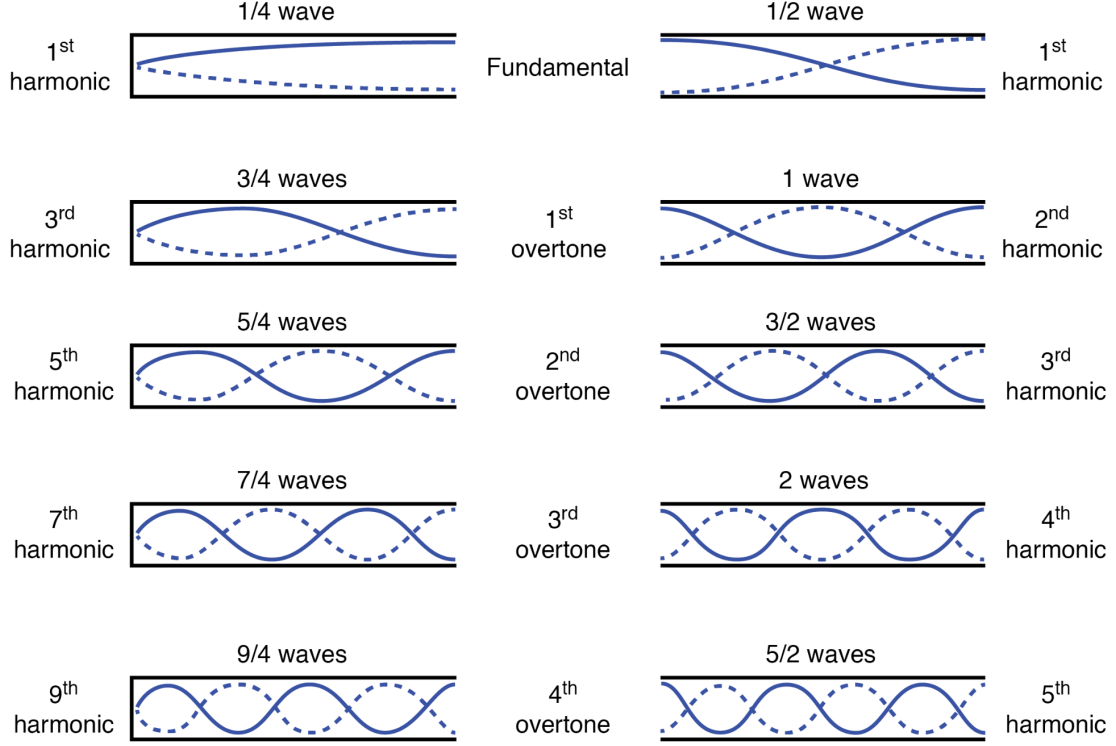


Figure 7: Schematics of the wave fitting in quarter-wave and half-wave resonators and their harmonics.

3.2 Lumped element resonators

The main disadvantage of resonator with a coplanar waveguide geometry is their size. An alternative and much more compact approach is to design and microfabricate superconducting microwave resonators with a *lumped element* geometry. Lumped elements are, by definition, elements which are much smaller than the resonant wavelengths, forming a tank circuit.

In a typical design, as shown in figure 8, the capacitor C is implemented with an interdigital structure, while the inductor L consists of a meander structure. In this example, the resonator is coupled to the transmission line via an interdigital coupling capacitor. The ground planes on either side of the transmission line are connected in order to suppress spurious modes caused by the asymmetry in the transmission line introduced by the coupling capacitance.

The resonance frequency of such a lumped element microwave resonator can be calculated [6]:

$$f_r = \frac{\omega_r}{2\pi} = \frac{1}{2\pi\sqrt{(L + L_T)(C + C_c^{\text{eff}})}} \quad (9)$$

where $C_c^{\text{eff}} = C_c C_p / (C_c + C_p)$ and L_T is the load inductor, as shown in figure 8.

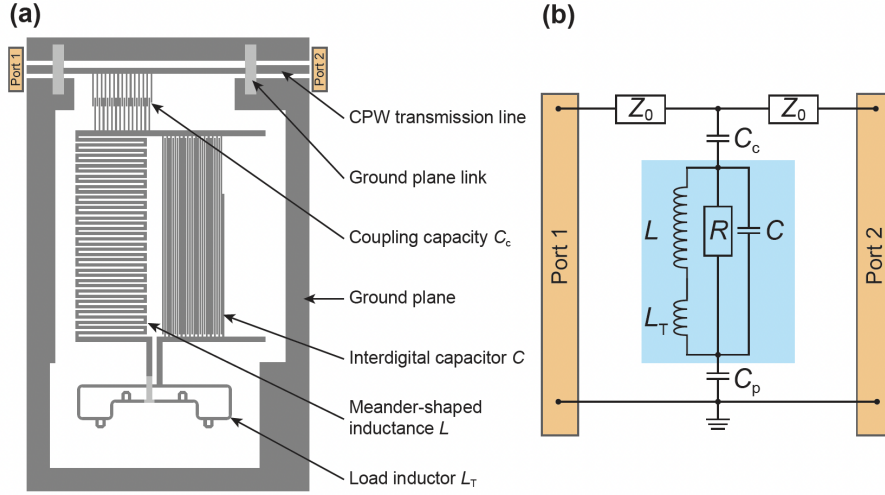


Figure 8: a) Layout of a planar lumped element microwave resonator coupled to a transmission line. b) Equivalent circuit of the lumped element resonator. Losses are accounted for by the resistance R , while the parasitic coupling to the ground planes is represented by the capacity C_p . Figure from [1]

3.3 Quality factors

The stored energy in a resonator in relation of the energy loss per cycle is measured by the *loaded quality factor* Q_l , which can be calculated as

$$\frac{1}{Q_l} = \frac{1}{Q_i} + \frac{1}{Q_c} \quad (10)$$

where Q_i is the *internal quality factor* (also called *intrinsic quality factor*) and Q_c is the *coupling quality factor*. The internal quality factor accounts for loss mechanism within the superconductor, such as:

- radiative losses, which depend on the geometry of the resonator and can be suppressed by reducing the width of the central conductor s ;
- dielectric losses, linked to the resonant energy absorption in the two-level systems (TLSs) present in the dielectric material;
- quasiparticle losses, which are due to the scattering of quasiparticles and therefore are expected to decrease drastically with falling temperature.

In a typical configuration, the superconducting resonator is capacitively coupled to a transmission line. The coupling quality factor Q_c depends on the coupling capacitance C_c . For example, for a quarter-wave resonator, the following relation holds:

$$Q_c = \frac{2\pi}{(2Z_0\omega_r C_r)^2} \quad (11)$$

where $\omega_r = 2\pi f_r$ and f_r is the resonance frequency of the resonator, taking into account the shift due to the coupling to the transmission line.

The resonator bandwidth Δf is linked to the loaded quality factor Q_l :

$$\Delta f = \frac{f_r}{Q_l}. \quad (12)$$

If we are able to minimise the internal losses, we can assume $Q_i \rightarrow \infty$, and therefore, according to equation 10

$$Q_l \approx Q_c. \quad (13)$$

In this way, we can control the loaded quality factor and consequently the resonator bandwidth tuning only the coupling to the transmission line.

3.4 3D superconducting resonators

Coplanar waveguide resonators are affected by internal losses, for example due to dielectric losses at interfaces and surfaces. In order to mitigate these effects, we can lower the ratio between the electric field energy stored at interfaces/surfaces and the energy stored in vacuum. In fact, planar resonators with larger sizes (i.e. larger geometrical parameters s and w) feature a weaker electric field near the interfaces/surfaces and they typically can achieve larger internal quality factors.

Following the same line, we can push this ratio even further and build 3D microwave cavities, where a large fraction of the field is stored inside the cavity and thus in vacuum. We can define the *surface participation ratio* as the ratio of the energy stored at surfaces versus in vacuum. In 3D metallic cavities, the surface participation ratio can reach 10^{-7} , while for planar structures typically values are around 10^{-5} . With 3D cavities it is possible to routinely achieve larger quality factors, up to 10^{10} , if compared to coplanar waveguide resonators or lumped-element resonators, allowing for longer average photon lifetimes.

An example of 3D rectangular cavity is shown in figure 9a.

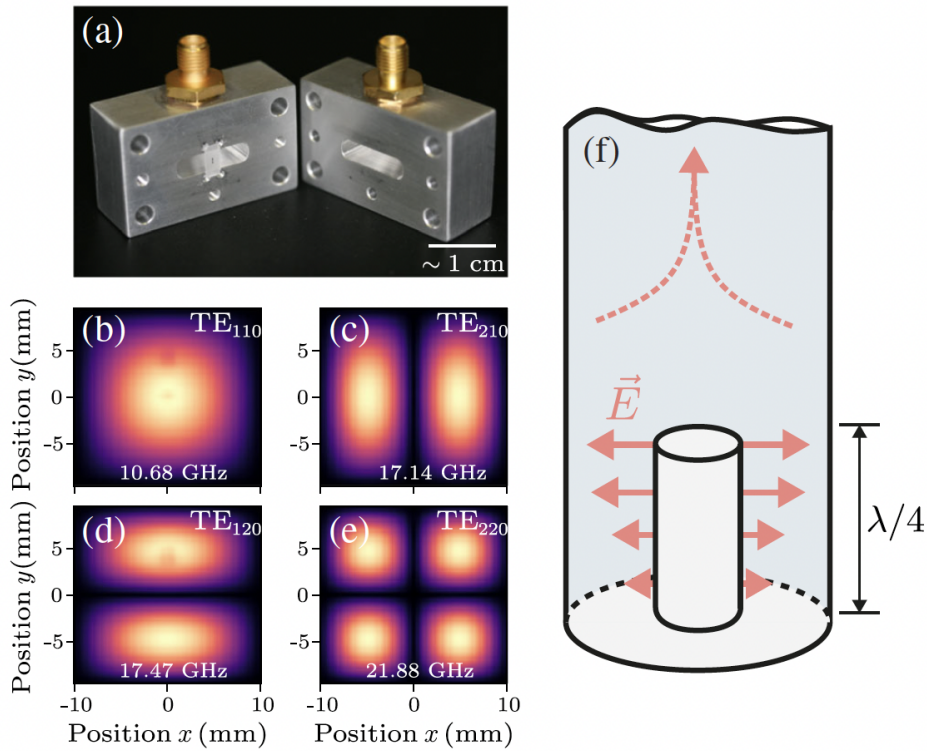


Figure 9: (a) Photograph of a 3D rectangular superconducting cavity. The inner volume houses a sapphire chip and transmon qubit, with two symmetric coaxial connectors for coupling signals in and out. (b)–(e) First four TE_{mnl} modes of a 3D rectangular superconducting cavity obtained from simulations. (f) Schematic representation of a coaxial quarter-wave cavity with electric field (solid line) pointing from the inner conductor to the sidewalls and evanescent field (dashed line) rapidly decaying from the top of the inner conductor. Figure from [2].

It consists of a finite section of a rectangular waveguide terminated by two metal walls acting as shorts, in analogy to the planar case previously discussed. The TE modes to which superconducting artificial atoms couple are illustrated in figure 9b-e.

Typical materials used to realise these types of cavities include Al and Cu, the first one reaching higher quality factors and the second one to allow usage of magnetic field (avoiding the limit of Meissner effect).

The walls of the cavity pose the boundary conditions which lead to the formation of TE and TM cavity independent modes at the following frequencies, labelled by the integers (m, n, l) :

$$\omega_{mnl} = c \sqrt{\left(\frac{m\pi}{a}\right)^2 + \left(\frac{n\pi}{b}\right)^2 + \left(\frac{l\pi}{d}\right)^2} \quad (14)$$

where a, b, d are the dimensions of the cavity. Dimensions in the order of a centimeter correspond to resonance frequencies in the order of gigahertz.

Another type of 3D microwave cavity is the coaxial cavity, shown in figure 9f. The main advantages of coaxial cavities is their robustness against fabrication imperfections and longer photon lifetimes, i.e. larger quality factors.

3.5 Representations of resonances

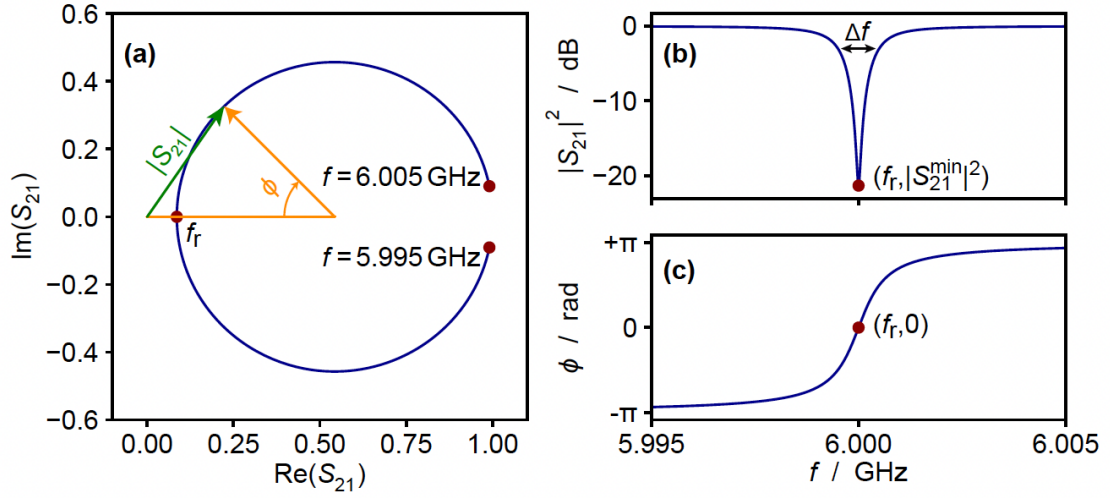


Figure 10: Different representations of the scatter parameter S_{21} and derived quantities in a frequency window of 10 MHz around a resonance at $f_r = 6 \text{ GHz}$ with $Q_1 = 6000$ and $Q_i = 70000$. (a) resonance circle in the complex plane, (b) transmission as a function of frequency and c) resonator phase as a function of frequency. Figure from [1].

4 The Quantum Harmonic LC Oscillator

A superconducting microwave LC resonator can be modelled as a quantum harmonic oscillator, where the energy of the system is oscillating between the magnetic energy in the inductor L and the electrical energy in the capacitor C . Figure 11a shows the circuit diagram.

The classical Hamiltonian of the LC resonator is given by:

$$H = \frac{Q^2}{2C} + \frac{\Phi^2}{2L} = \frac{1}{2}CV^2 + \frac{1}{2}LI^2 \quad (15)$$

where Q and Φ are the generalised circuit coordinates, charge and flux, respectively.

By definition, the flux Φ is the integral of the voltage V : $\Phi(t) = \int V(t')dt'$ and the charge Q is the integral of the current I : $Q(t) = \int I(t')dt'$. We also should remember that $V = LdI/dt$ and $I = CdV/dt$.

We can interpret the electrical energy as the kinetic energy and the magnetic energy as the potential energy of the oscillator, in analogy with a mechanical oscillator. In this analogy, the usual conjugate variables to describe a mechanical oscillators, position x and momentum p , are mapped to charge Q and flux Φ . The Hamiltonian in equation 15 is analogous to that of mechanical harmonic oscillator, with mass $m = C$ and momentum p . Expressed in position x and momentum p , the Hamiltonian reads $H = p^2/2m + m\omega^2x^2/2$.

To migrate from a classical to a quantum description, we introduce the quantum operators $\hat{\Phi}$ and $\hat{Q}$ which satisfy the commutation relation $[\hat{\Phi}, \hat{Q}] = i\hbar$. The operators hats are omitted in the following for simplicity.

We define the reduced flux $\phi := 2\pi\Phi/\Phi_0$ and the reduced charge $n := Q/2e$, which obey the commutation relation $[\phi, n] = i$. The quantum-mechanical Hamiltonian can then be derived:

$$H = 4E_C n^2 + \frac{1}{2}E_L \phi^2 \quad (16)$$

where $E_C = e^2/(2C)$ is the charging energy necessary to add a Cooper pair to the capacitor island and $E_L = (\phi_0/2\pi)^2/L$ is the inductive energy. The quantum operator n corresponds to the excess number of Cooper pairs on the island in the circuit shown in figure 11. The reduced flux ϕ corresponds to the phase across the inductor.

The resulting potential, shown in figure 11, exhibits equally spaced energy levels.

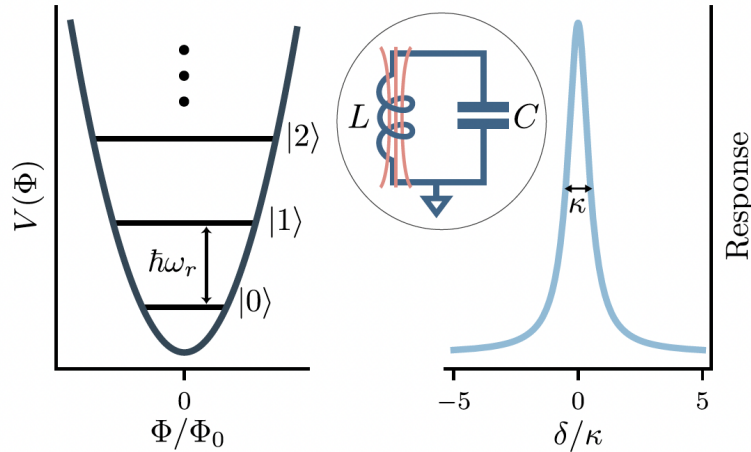


Figure 11: Left panel: harmonic potential vs flux of the LC circuit. Right panel: response of the oscillator to an external perturbation as a function of the detuning δ of the perturbation from the oscillator frequency. Figure adapted from [2]

The response of the oscillator to an external perturbation is shown in figure 11, where δ is the detuning of the perturbation from the oscillator frequency and $k = \omega_0/Q$ is the FWHM of the

oscillator response, where Q is the quality factor. The average lifetime of single-photon state $|1\rangle$ before decaying to $|0\rangle$ is given by $1/k$.

In a more compact form, we can write the Hamiltonian of the quantum harmonic oscillator as:

$$H = \hbar\omega \left(a^\dagger a + \frac{1}{2} \right) \quad (17)$$

where $a^\dagger$ and a are the creation and annihilation operators of a single excitation of the resonator.

The quantum LC oscillator and the superconducting resonators that we have reviewed so far are not tunable and their resonance frequency is fixed by the geometry of the circuit. Moreover, so far the circuit dynamics we have encountered is fully linear. In order to build interesting circuits for quantum optics and cQED experiments as well as to realise superconducting qubits, we need to introduce tunability and non-linearity in our superconducting circuits. In the superconducting toolbox, we have a special element that behaves as a controllable, dissipationless and tunable non-linear inductance, namely the Josephson junction.

5 Josephson Junctions and Josephson inductance

The Josephson junction is one of the most important circuit elements in the field of superconducting quantum devices. It consists of a sandwich of two superconducting electrodes separated by a thin insulating barrier which allows coherent tunnelling of Cooper pairs, the charge carriers in a superconductor.

Practically, Josephson junctions are realised overlapping two thin superconducting films, for example ~ 100 nm thick Al films, separated by a ~ 1 nm aluminium oxide barrier. A schematic of a Josephson junction is shown in figure 12.

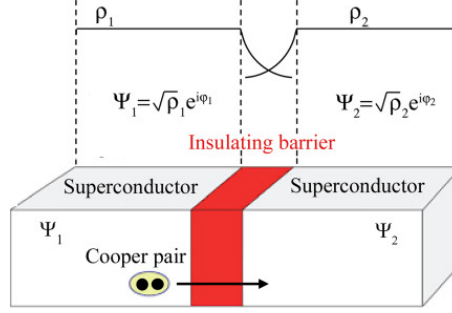


Figure 12: A simple schematic of a Josephson junction. The two superconducting electrodes are described by the macroscopic wavefunctions Ψ_1 and Ψ_2 .

Josephson equations

To describe the physics of a Josephson junction, we should recall the foundations of the Bardeen – Cooper – Schrieffer (BCS) theory, which describes the phenomenon of superconductivity as the result of the condensation of Cooper pairs. At low temperatures, i.e. when the energy of the system is highly reduced, an attractive potential between electrons is generated which causes binding of electrons into Cooper pairs. Thanks to their bosonic nature, Cooper pairs form a large Bose-Einstein condensate. Under these conditions, the quantum mechanical character of the ground state of the system appears on a macroscopic scale.

As a result, a single wave function $\Psi(x, t)$ can describe a macroscopic number of Cooper pairs that condensate in the same quantum state. We can therefore write a *macroscopic wave function* describing a macroscopic piece of superconductor, such as each superconductor forming a Josephson junction:

$$\Psi(x, t) = \sqrt{\rho} e^{i\phi(x, t)} \quad (18)$$

where ϕ is the common phase and ρ is the density in the macrostate, with $|\Psi|^2 = \rho$.

The time evolution of the macroscopic wave function in stationary conditions follows the Schrödinger equation:

$$i\hbar \frac{\partial \Psi}{\partial t} = H\Psi \quad (19)$$

Bearing this in mind, we can derive the dynamics of a Josephson junction by introducing two macroscopic wave functions Ψ_1 and Ψ_2 for the two superconductors (see figure 12) and using the Schrödinger equation with a coupling term describing the overlap of the wave functions in the barrier. The full derivation is beyond the scope of these lectures and it can be found in Feynman's lecture *The Schrödinger Equation in a Classical Context: A Seminar on Superconductivity* (see the section **Suggested Textbooks and Reviews**).

As a result, the dynamic of a Josephson junction is described by two simple and elegant equations called *Josephson equations*.

First Josephson equation The first Josephson equation describes the current I flowing in the Josephson junction:

$$I = I_c \sin(\phi) \quad (20)$$

where I_c is a parameter of the junction called *critical current* and $\phi = \phi_2 - \phi_1$ is the phase difference between the two superconductors.

From equation 20, it is evident that even in the absence of an applied voltage across the Josephson junction, a non-zero current can flow, up to the critical current value I_c . This phenomenon is called *dc Josephson effect* and it is a direct consequence of the finite phase difference between the macroscopic wave functions in the two superconductors.

Second Josephson equation The second Josephson equation described the evolution of the phase difference ϕ :

$$\frac{\partial \phi}{\partial t} = \frac{2eV}{\hbar} \quad (21)$$

where V is the voltage applied across the junction.

From equation 21, we observe that a non-zero voltage across the Josephson junction generates an oscillating current across the junction. This is easily shown by integrating equation 21: $\phi(t) = \frac{2eV}{\hbar}t + \phi_0$, and inserting the term $\phi(t)$ in equation 20:

$$I(t) = I_c \sin\left(\frac{2eV}{\hbar}t + \phi_0\right) \quad (22)$$

Equation 22 describes the so called *ac Josephson effect*, generating an ac current when a voltage is applied.

Josephson inductance

The Josephson equations not only describe the dc and ac Josephson effects previously reviewed, but they also determine that the behaviour of a Josephson junction can be modelled as a non-linear inductance.

Starting from equation 20, we can write the time derivative as:

$$\frac{\partial I}{\partial t} = I_c \cos \phi \frac{\partial \phi}{\partial t} = \frac{2\pi I_c}{\Phi_0} V \cos \phi \quad (23)$$

where $\Phi_0 = h/2e$ is the flux quantum.

We observe that the time derivative of the current is proportional to a voltage, representing an inductive behaviour, according to the general definition of inductance L : $V = L \frac{\partial I}{\partial t}$. We can therefore define the *Josephson inductance* L_J and derive it from equation 23:

$$L_J(V) = \frac{\Phi_0}{2\pi I_c \cos \phi(V)} = \frac{L_0}{\cos \phi(V)} \quad (24)$$

where $L_0 = \frac{\Phi_0}{2\pi I_c}$.

Equation 24 shows that the Josephson inductance L_J non-linearly depends on the phase difference ϕ , which in turns depends on the voltage V . For small values of V , equation 24 can be linearised, but for sufficiently large values of V , which corresponds to most of experimental cases, the inductance remains non-linear.

The crucial consequence is that Josephson junctions can be used as a non-linear and tunable inductances in superconducting circuits. For this reason, Josephson junctions are one of the most important building blocks for quantum circuits such as tunable resonators, parametric amplifiers and qubits.

6 Kinetic inductance

Another approach to introduce a nonlinear inductance in our circuit is to make use of the so called *kinetic inductance* of superconductors. To understand this concept, let's go back to Drude model and write the equation of motion for an electron gas with a scattering term:

$$m \frac{d\mathbf{v}}{dt} = e\mathbf{E}(t) - m \frac{\mathbf{v}}{\tau} \quad (25)$$

where τ is the scattering time (or life time) and $\mathbf{E}(t) = Ee^{-i\omega t}$ is the a.c. electric field.

Recalling the two fluid model, we can think to the charge carrier density n as the sum of two contributions, namely the superconducting contribution n_s and the dissipative contribution n_n : $n = n_s + n_n$. The first one corresponds to the density of Copper pairs, while the second one corresponds to the density of quasiparticles. The corresponding current densities can be written as:

$$\begin{aligned} \mathbf{J}_n(t) &= n_n e \mathbf{v}_n(t) \\ \mathbf{J}_s(t) &= n_s e \mathbf{v}_s(t) \end{aligned} \quad (26)$$

The current density $\mathbf{J}$ is related to the electric field $\mathbf{E}$ via the conductivity, according to the equation

$$\mathbf{J}(\omega) = \sigma(\omega) \mathbf{E} \omega \quad (27)$$

where $\sigma(\omega)$ is the complex conductivity. Its real part corresponds to currents in phase with the applied field (resistive component), while its imaginary part corresponds to out of phase currents (inductive component).

Deriving with respect to time equations 26, we obtain for the contribution from quasiparticles, i.e. from n_n :

$$\begin{aligned} \frac{d\mathbf{J}_n}{dt} &= n_n e \frac{d\mathbf{v}_n}{dt} = \frac{n_n e}{m} \left(e\mathbf{E} - \frac{m}{\tau_n} \mathbf{v}_n \right) \\ &= \frac{n_n e^2}{m} \mathbf{E} - \frac{\mathbf{J}_n}{\tau_n} = \frac{1}{\tau_n} \left[\left(\frac{n_n e^2 \tau_n}{m} \right) \mathbf{E} - \mathbf{J}_n \right] \\ &\rightarrow \tau_n \frac{d\mathbf{J}_n}{dt} = \sigma_n \mathbf{E} - \mathbf{J}_n \end{aligned} \quad (28)$$

The quasiparticle contribution to the complex conductivity can therefore be written as:

$$\sigma_n = \frac{n_n e^2 \tau_n}{m}. \quad (29)$$

For contribution from Cooper pairs, i.e. from n_s , we can assume infinite scattering time $\tau_s \rightarrow \infty$, and we can write:

$$\frac{d\mathbf{J}_s}{dt} = n_s e_s \frac{d\mathbf{v}_s}{dt} = n_s e_s \frac{e_s}{m_s} \mathbf{E} = \frac{n_s e_s^2}{m_s} \mathbf{E} \quad (30)$$

The Fourier transformation of a general electric field is:

$$\mathbf{E}(\mathbf{r}, t) = \int_{-\infty}^{+\infty} \mathbf{E}(\mathbf{r}, \omega) e^{-i\omega t} d\omega \quad (31)$$

with $\mathbf{E}(\mathbf{r}, -\omega) = \mathbf{E}^*(\mathbf{r}, \omega)$.

$$\mathbf{E}(\mathbf{r}, t) = \text{Re} [\mathbf{E}(\mathbf{r}, t) e^{-i\omega t}] = \frac{1}{2} \int_0^{+\infty} d\omega' (\mathbf{E}(\mathbf{r}, \omega') e^{-i\omega' t} + \text{c.c.}) \delta(\omega - \omega') \quad (32)$$

where the last term described a monochromatic field.

Considering that $\mathbf{J}_{n,s}(\mathbf{r}, t) = \text{Re} [\mathbf{J}_{n,s}(\mathbf{r}, t) e^{-i\omega t}]$, we can derive in the frequency domain:

$$\begin{aligned}
-i\omega\tau_n\mathbf{J}_n(\mathbf{r},\omega) &= \sigma_n\mathbf{E}(\mathbf{r},\omega) - \mathbf{J}_n(\mathbf{r},\omega) \\
-i\omega\mathbf{J}_s(\mathbf{r},\omega) &= \frac{n_s e_s^2}{m_s}\mathbf{E}_s(\mathbf{r},\omega)
\end{aligned} \tag{33}$$

and finally write:

$$\begin{aligned}
\mathbf{J}_n(\mathbf{r},\omega) &= \frac{\sigma_n}{1-i\omega\tau_n}\mathbf{E}(\mathbf{r},\omega) \\
\mathbf{J}_s(\mathbf{r},\omega) &= i\frac{n_s e_s^2}{m_s\omega}\mathbf{E}_s(\mathbf{r},\omega)
\end{aligned} \tag{34}$$

and

$$\sigma_s = \frac{n_s e_s^2}{m_s\omega}. \tag{35}$$

The total current density $\mathbf{J}$ is the sum of the two contributions $\mathbf{J}_n$ and $\mathbf{J}_s$:

$$\mathbf{J}(\mathbf{r},\omega) = \mathbf{J}_n(\mathbf{r},\omega) + \mathbf{J}_s(\mathbf{r},\omega) = \left(\frac{\sigma_n}{1-i\omega\tau_n} + i\frac{n_s e_s^2}{m_s\omega} \right) \mathbf{E}(\mathbf{r},\omega). \tag{36}$$

Assuming the limit of frequencies much lower than the superconducting gap, i.e. $\omega\tau_n \ll 1$, the final form is the following:

$$\mathbf{J}(\mathbf{r},\omega) = (\sigma_n + i\sigma_s(\omega))\mathbf{E}(\mathbf{r},\omega) \tag{37}$$

where the contribution to conductivity from the quasiparticles, σ_n , is given by equation 29, while the contribution from Cooper pairs, σ_s , is given by equation 35.

We notice that the conductivity due to Cooper pairs is a purely imaginary term, corresponding to out of phase currents, and therefore only inductive. On the other hand, the conductivity due to quasiparticles is a real term, corresponding to in phase currents, and therefore resistive and dissipative. Due to the real part, some amount of resistive response and therefore dissipation always remains, due to the presence of quasiparticles. The Cooper pairs, instead, give rise to an inductive behaviour and to a phase shift. The ratio between the densities n_s and n_n depends on the operational temperature T :

$$\frac{n_s}{n_s + n_n} = \frac{n_s}{n} = 1 - \left(\frac{T}{T_c} \right)^4 \tag{38}$$

where T_c is the critical temperature. Thus, the density of quasiparticles is suppressed when $T \ll T_c$.

We can interpret the phase shift due to the imaginary part of the conductivity as a phase lag between the oscillating electric field and the current density due to the inertia of the charge carriers. This phenomenon effectively behaves as an additional inductance in a circuit and it takes the name of *kinetic inductance*.

The conductivity of a superconductor depends on the density of Cooper pairs n_s , which in turns depends non-linearly on the current flowing through the superconductor. Therefore, the kinetic inductance L_k of a superconducting circuit is nonlinear with respect to the current I in the circuit and, in first approximation, it exhibits a quadratic behaviour:

$$L_k(I) \approx L_0 \left(1 + \frac{I^2}{I_*^2} \right) \tag{39}$$

with $L_0 = R_S \hbar / \pi \Delta$, where R_S is the sheet resistance of the film in normal state and Δ is the superconducting gap. I_* gives the nonlinearity scales and it is lower for higher resistances of the film: $I_* \propto 1/\sqrt{R_N}$, where R_N is the normal resistance of the film.

7 Superconducting devices: an overview

Superconducting devices range from 3D devices, such as superconducting magnets and 3D cavities, to planar circuits. Focusing on superconducting circuits, we can categorise them according to:

- Operation: DC operation (DC current or voltage signals, transport measurements) versus RF operation (microwave signals, transmission measurements);
- Platform: Josephson junctions, high kinetic inductance superconductors, ... ;
- Main application: magnetometry, metrology, radiation detection, thermometry, signal processing (e.g. microwave amplification), quantum computing,

The devices that will be discussed in this course are listed in the table below.

Superconducting devices			
Device	RF/DC operation	Building blocks	Main applications
DC SQUID (Superconducting Quantum Interference Device)	DC (transport)	Josephson junctions	Magnetometry, metrology, current sensing
TES (Transition Edge Sensor)	DC (transport)	Superconductors	Radiation detection (wide range), thermometry
MMC (Magnetic MicroCalorimeter)	DC (transport)	Superconductors + paramagnetic material	Radiation detection (wide range)
KID (Kinetic Inductance Detector)	RF (transmission)	High kinetic inductance superconductors	Radiation detection (wide range)
KICS (Kinetic Inductance Current Sensor)	RF (transmission)	High kinetic inductance superconductors	Current sensing, multiplexed detector readout
SNSPD (Superconducting Nanowire Single Photon Detector)	DC (transport)	High kinetic inductance superconductors	Photon counting
JPA, TWPA (Josephson Parametric Amplifier, Travelling-Wave Parametric Amplifier)	RF (transmission)	Josephson junctions / High kinetic inductance elements	Microwave amplification, microwave squeezing
Qubit (transmon, fluxonium, ...)	RF (transmission)	Josephson junctions	Quantum computing, quantum simulation

8 SQUIDS: Superconducting Quantum Interference Devices

8.1 Introduction

Superconducting Quantum Interference Devices (SQUIDS) are devices consisting of a superconducting loop interrupted by one or more Josephson junctions, as shown in figure 13. They are the most sensitive detectors for magnetic flux available. We know that also single Josephson junctions are sensitivity to magnetic fields, since the maximum Josephson current is strongly modulated by a magnetic field. However, intuitively, we understand that using a superconducting loop enhances the sensitivity to the magnetic flux with respect to a single junction.

The physics of SQUIDS is based on the combination of two physical phenomena, namely flux quantization in superconducting loops and the Josephson effect. In essence, a SQUID is a flux to voltage converter and it provides a flux dependent and periodic output voltage. The period corresponds to one flux quantum Φ_0 . Therefore, SQUIDS can measure all physical quantities that can be converted into magnetic flux. They are used as sensors for magnetic fields and magnetic field gradients, current, voltage, displacement, or magnetic susceptibility. For example, SQUIDS are used for medical diagnosis and bio-medical applications, where an extreme sensitivity to weak magnetic signals, such as the magnetic response of our brain, is required. Furthermore, SQUIDS are used for standards and metrology.

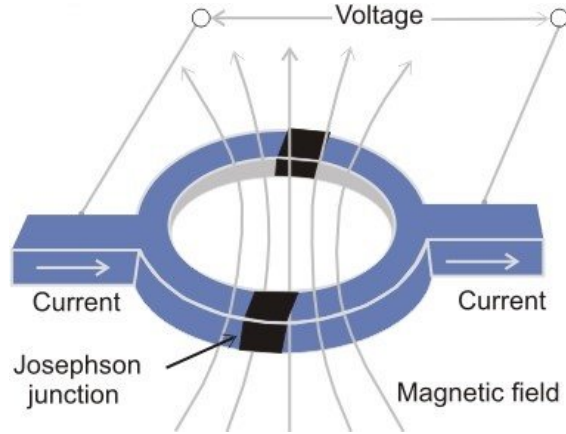


Figure 13: Simple schematic of a SQUID circuit, consisting of a superconducting loop interrupted by two Josephson junctions.

8.2 The zero-voltage state

The device depicted in figure 13 is known as *direct current superconducting quantum interference device*, or dc-SQUID.

In this circuit we have two Josephson junctions, that we treat as lumped elements, connected in parallel and joined by a superconducting loop. We assume that the two Josephson junctions are characterised by the same critical current I_c and thus by the same current-phase relation $I_{si} = I_c \sin \phi_i$, with $i = 1, 2$. Using Kirchhoff's law, we obtain the total current I_s :

$$I_s = I_{s1} + I_{s2} = I_c \sin \phi_1 + I_c \sin \phi_2 = 2I_c \cos \left(\frac{\phi_1 - \phi_2}{2} \right) \sin \left(\frac{\phi_1 + \phi_2}{2} \right). \quad (40)$$

Calculating the total phase change along the closed contour shown in figure ?? and considering that (i) the integration of the vector potential $\mathbf{A}$ gives the total flux Φ and (ii) if the integration path is deep inside the superconducting metal the current density is negligible, we obtain that the two phase differences are linked to each other by the flux quantisation condition²:

$$\phi_2 - \phi_1 = \frac{2\pi\Phi}{\Phi_0}. \quad (41)$$

²For the complete derivation, please refer to [4], chapter 4, section 4.1.1

Combining the last two expressions, we can derive the current I_s :

$$I_s = 2I_c \cos\left(\pi \frac{\Phi}{\Phi_0}\right) \sin\left(\phi_1 + \pi \frac{\Phi}{\Phi_0}\right) \quad (42)$$

If the flux Φ threading the loop was given just by the external flux Φ_{ext} , the maximum super-current would be given by:

$$I_s^{\text{max}} = 2I_c \cos\left|\pi \frac{\Phi}{\Phi_0}\right|. \quad (43)$$

However, typically we need to take into account the inductance L of the superconducting loop, which generates an extra flux Φ_L , with the total flux $\Phi = \Phi_{\text{ext}} + \Phi_L$. The currents flowing in the two arms of the SQUID loop are given by:

$$\begin{aligned} I_{s1} &= \tilde{I} + I_{\text{cir}} \\ I_{s2} &= \tilde{I} - I_{\text{cir}} \end{aligned} \quad (44)$$

where $\tilde{I} = (I_{s1} + I_{s2})/2$ and $I_{\text{cir}} = (I_{s1} - I_{s2})/2$ are the average current common in both arms and the screening current circulating in the loop, respectively.

The total flux Φ is given by:

$$\begin{aligned} \Phi &= \Phi_{\text{ext}} + LI_{\text{cir}} = \Phi_{\text{ext}} + \frac{LI_c}{2} (\sin \phi_1 - \sin \phi_2) \\ &= \Phi_{\text{ext}} + LI_c \sin\left(\frac{\phi_1 - \phi_2}{2}\right) \cos\left(\frac{\phi_1 + \phi_2}{2}\right). \end{aligned} \quad (45)$$

and combining this result with equation 41, we can write:

$$\Phi = \Phi_{\text{ext}} - LI_c \sin\left(\pi \frac{\Phi}{\Phi_0}\right) \cos\left(\phi_1 + \pi \frac{\Phi}{\Phi_0}\right) \quad (46)$$

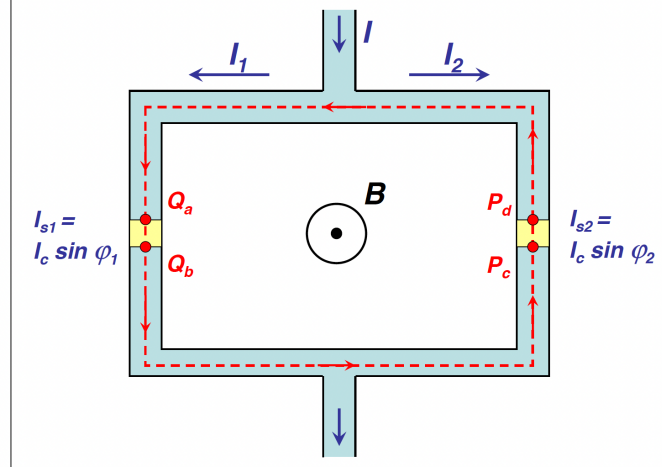


Figure 14: Schematic of the SQUID circuit, comprising two arms with two identical Josephson junctions. Figure from [4].

We define the so-called *screening parameter* the ratio between the magnetic flux generated by the maximum circulating current $I_{\text{cir}} = I_c$ and $\Phi_0/2$:

$$\beta_L := \frac{LI_c}{\Phi_0/2} = \frac{2LI_c}{\Phi_0}. \quad (47)$$

Following the derivation reported in [4] (chapter 4), we can consider the two following limits.

Negligible screening limit: $\beta_L \ll 1$

If $\beta_L \ll 1$, the screening circulating current is small when compared to one flux quantum and therefore can be neglected when compared to the external flux. Maximising equation 42 with respect to ϕ_1 , we impose $dI_s/d\phi_1 = 0$ and we obtain:

$$\cos\left(\phi_1 + \pi \frac{\Phi_{\text{ext}}}{\Phi_0}\right) = 0 \quad (48)$$

and therefore:

$$I_s^{\text{max}} \approx 2I_c \left| \cos\left(\pi \frac{\Phi_{\text{ext}}}{\Phi_0}\right) \right| \quad (49)$$

which corresponds to equation 43. The maximum supercurrent is a periodic function of the external flux, as shown in figure REF

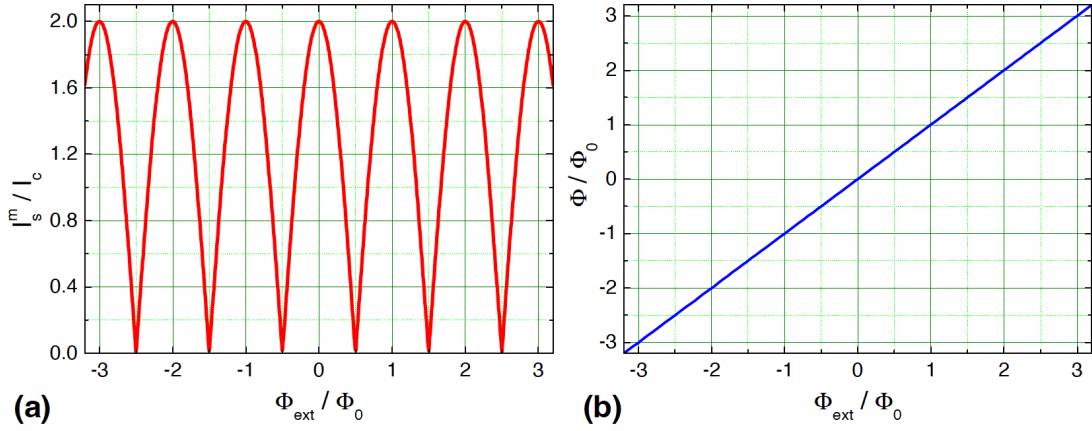


Figure 15: (a) The maximum supercurrent as a function of the applied magnetic flux in the limit of negligible screening. (b) The flux threading the SQUID loop as a function of the the applied flux in the same limit. Figure from [4].

Large screening limit: $\beta_L \gg 1$

For large loop inductances, the maximum screening supercurrent is bigger than one flux quantum: $LI_c \gg \Phi_0$ and the circulating current ensures flux quantisation by compensating the external flux. The total flux is quantised as:

$$\Phi = \Phi_{\text{ext}} + LI_c \approx n\Phi_0. \quad (50)$$

The current through the SQUID is

$$I_s = I_c \sin \phi_1 + I_c \sin \phi_2 \quad (51)$$

while the circulating current is

$$I_{\text{cir}} = I_c/2 \cdot (\sin \phi_1 - \sin \phi_2) \quad (52)$$

and these two relations are constraint by the condition:

$$\phi_2 - \phi_1 = \frac{2\pi\Phi}{\Phi_0} \quad (53)$$

where the flux is the sum of the external flux Φ_{ext} and the flux due to the screening current $\Phi_{\text{cir}} = LI_{\text{cir}}$. Given the current and the total flux, we can solve the two equations for the phase

differences and find Φ_{ext} and Φ_{cir} . For instance, if $I \approx 0$, we have that $\sin \phi_1 \approx -\sin \phi_2$ and we obtain:

$$\begin{aligned}\Phi_{\text{ext}} &= \Phi + LI_c \sin \left(\pi \frac{\Phi}{\Phi_0} \right) \\ \frac{\Phi_{\text{ext}}}{\Phi_0} &= \frac{\Phi}{\Phi_0} + \frac{\beta_L}{2} \sin \left(\pi \frac{\Phi}{\Phi_0} \right).\end{aligned}\tag{54}$$

Figure REF show the result of the inverted relation, i.e. the flux in the SQUID loop Φ as a function of the external flux Φ_{ext} , for different values of β_L . Only for small values of β_L , and in particular for $\beta_L \leq 2/\pi$, the relation is single-valued. Note that when $\Phi = n\Phi_0$, no screening current is flowing in the SQUID, $I_{\text{cir}} = 0$, and the flux is equal to the external flux, $\Phi = \Phi_{\text{ext}}$. In these points, the curve does not depend on β_L .

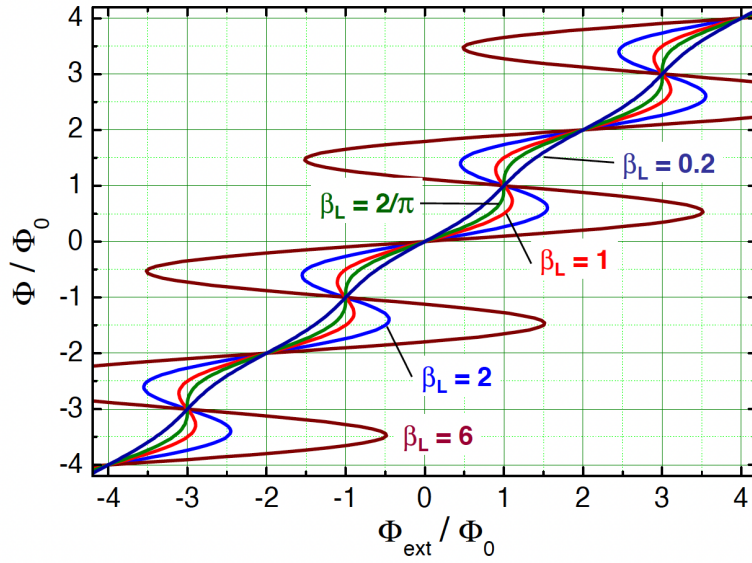


Figure 16: The total magnetic flux as a function of the applied magnetic flux, for different values of β_L . Figure from [4].

To obtain the dependence of the supercurrent on the applied magnetic flux, let us use equation 50 to write:

$$I_{\text{cir}} \approx -\frac{\Phi_{\text{ext}} - n\Phi_0}{L}\tag{55}$$

where $I_{\text{cir}} \rightarrow 0$ for $L \rightarrow \infty$. The current then divides equally in the two arms and the maximum current is $I_s^{\text{max}} \approx 2I_c$. Initially $n = 0$ and a small screening current $I_{\text{cir}} \approx -\Phi_{\text{ext}}/L$ is generated to screen the magnetic field. Therefore, with increasing Φ_{ext} , the current in the first arm, I_1 , will decrease, while the current in the second arm, I_2 , will increase until $I_2 \leq I_c$. Once I_2 reaches this limit, I_1 will decrease as $I_1 \approx I_c - 2\Phi_{\text{ext}}/L$. Using this expression for I_1 and assuming $I_2 \approx I_c \approx \text{const}$, we can write the maximum supercurrent as:

$$\begin{aligned}I_s^{\text{max}} &\approx 2I_c - \frac{2\Phi_{\text{ext}}}{L} \\ \frac{I_s^{\text{max}}}{2I_c} &\approx 1 - \frac{2\Phi_{\text{ext}}}{\Phi_0} \frac{1}{\beta_L}.\end{aligned}\tag{56}$$

The modulation of the maximum supercurrent depends on the parameter β_L and it decrease for increasing values of β_L .

8.3 The voltage state

Real SQUID sensors are typically operated at a constant bias current $I_b > I_s^{\max}(\Phi_{\text{ext}} = 0)$ and therefore in the voltage state. In this situation, the output voltage is related to the applied magnetic flux, making the SQUID a flux-to-voltage converter.

We analyse the voltage state assuming negligible screening, i.e. $\beta_L \ll 1$, and strong damping, i.e. $\beta_C \ll 1$. The total current is therefore given by the sum of the Josephson current and the resistive current in the two arms:

$$\begin{aligned} I &= I_c \sin \phi_1 + I_c \sin \phi_2 + 2 \frac{V}{R_N} \\ &= 2I_c \cos \left(\pi \frac{\Phi}{\Phi_0} \right) \sin \left(\phi_1 + \pi \frac{\Phi}{\Phi_0} \right) + 2 \frac{V}{R_N} \end{aligned} \quad (57)$$

using the two relations previously derived:

$$\begin{aligned} 1) \quad I_s &= 2I_c \cos \left(\frac{\phi_1 - \phi_2}{2} \right) \sin \left(\frac{\phi_1 + \phi_2}{2} \right) \\ 2) \quad \phi_2 - \phi_1 &= \frac{2\pi\Phi}{\Phi_0}. \end{aligned} \quad (58)$$

We define a new phase $\phi = \phi_1 + \pi\Phi/\Phi_0$ and due to the fact that $\Phi \approx \Phi_{\text{ext}} \approx \text{const}$, we have that:

$$\frac{d\phi}{dt} = \frac{d\phi_1}{dt} = \frac{2\pi}{\Phi_0} V(t) \quad (59)$$

using the second Josephson equation. Thus, we can rewrite equation 57 as:

$$I = I_s^{\max}(\Phi_{\text{ext}}) \sin \phi + \frac{2}{R_N} \frac{2\pi}{\Phi_0} \frac{d\phi}{dt} \quad (60)$$

with

$$I_s^{\max} = 2I_c \cos \left(\pi \frac{\Phi_{\text{ext}}}{\Phi_0} \right). \quad (61)$$

From the last expression, we can conclude that the SQUID can be modelled as a single Josephson junction with a maximum Josephson current that depends on the external flux. Thanks to this equivalence, we can reuse the RSCJ model discussed in the first part of the course and write the current-voltage relation:

$$\begin{aligned} \langle V(t) \rangle &= I_c R_N \sqrt{\left(\frac{I}{2I_c} \right)^2 - \left(\frac{I_s^{\max}(\Phi_{\text{ext}})}{2I_c} \right)^2} \\ &= I_c R_N \sqrt{\left(\frac{I}{2I_c} \right)^2 - \left[\cos \left(\pi \frac{\Phi_{\text{ext}}}{\Phi_0} \right) \right]^2} \end{aligned} \quad (62)$$

where we see that the voltage response of the SQUID depends periodically on the external magnetic flux. The corresponding 3D graph is shown in figure 17.

8.4 SQUID operation

The equivalent circuit model of a dc SQUID is shown in figure 18(a). The two Josephson junctions are described using the RCSJ model, with a resistance R_N and a capacitance C . To avoid hysteresis, the condition $\beta_C = 2e/\hbar \cdot I_c R_N^2 C \leq 1$ should be satisfied. As we will discuss later, this is typically achieved lowering the effective resistance by adding extra resistive elements in parallel, called *shunt resistors*.

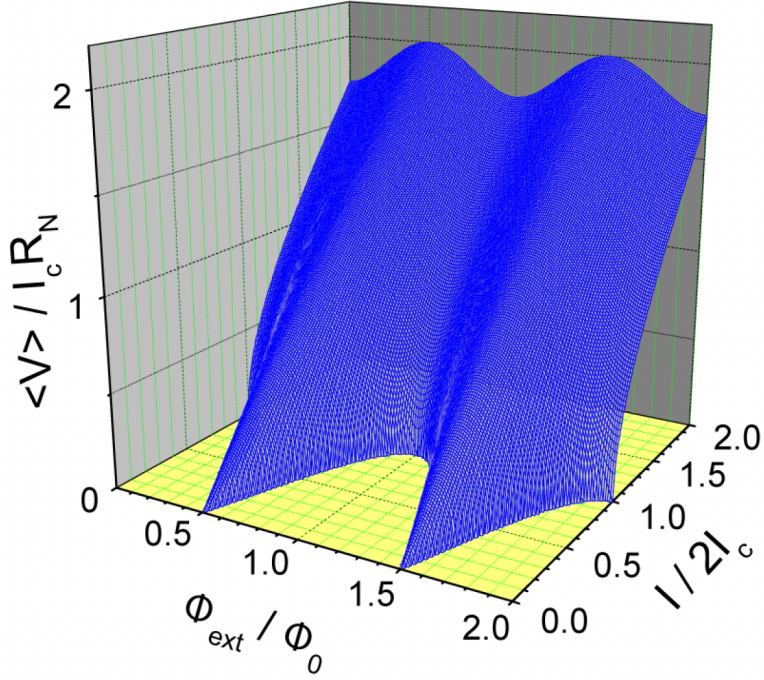


Figure 17: Current-voltage curve of a dc SQUID in the limit $\beta_L \ll 1$ and $\beta_C \ll 1$ as a function of the applied magnetic flux. Figure from [4].

For different values of external flux, we have different IV curves, as shown in figure 18(b). If we start with zero external magnetic field, we have the red curve. Increasing the magnetic field, we have intermediate curves, until we reach $\Phi_{\text{ext}} = \Phi_0/2$, corresponding to the blue curve. At this point, the periodicity previously discussed kicks in and we go back to intermediate curves, until we reach $\Phi_{\text{ext}} = \Phi_0$, corresponding again at the red curve. The two extreme curves correspond to the two values of external flux $\Phi_{\text{ext}} = (n + 1/2)\Phi_0$ and $\Phi_{\text{ext}} = n\Phi_0$, respectively.

The operation principle of a dc SQUID requires to choose a bias current value and a flux working point. Fixing a specific value of bias current, such as $I/2I_c = 1.8$, as indicated in figure 18(b) with the red dashed line, we can work out the expected voltage response as a function of external flux, as shown in figure 18(c), pink line. In order to maximise the modulation, the bias current should be just above $I/2I_c = 1$, i.e. $I \geq 2I_c$.

Once the bias current is fixed, and thus the voltage-flux characteristics is fixed, it is convenient to choose a flux working point on the steepest part of the slope. In this way, small variation of external flux will correspond to big voltage signals, maximising the sensitivity of our sensor. To do that, as we will see later, it is necessary to add a coil to add an extra flux to the SQUID.

To design and operate the dc SQUID, we have several nobs. We can decide how to build our Josephson junctions, setting their critical current and we can decide how to design the SQUID loop, setting its inductance. We can add shunt resistors to modify the effective resistance of the junctions and therefore to tune the β_C parameter. To understand how to design a SQUID, we can define three main figures of merit which represent the performance of the SQUID:

- The **flux-to-voltage transfer coefficient**, corresponding to the derivative of the voltage-flux characteristics for a given bias current:

$$\left| \left(\frac{\partial V}{\partial \Phi_{\text{ext}}} \right)_{I=\text{const}} \right| \quad (63)$$

- The **equivalent flux noise**, corresponding to the power spectral density of the voltage $S_V(f)$

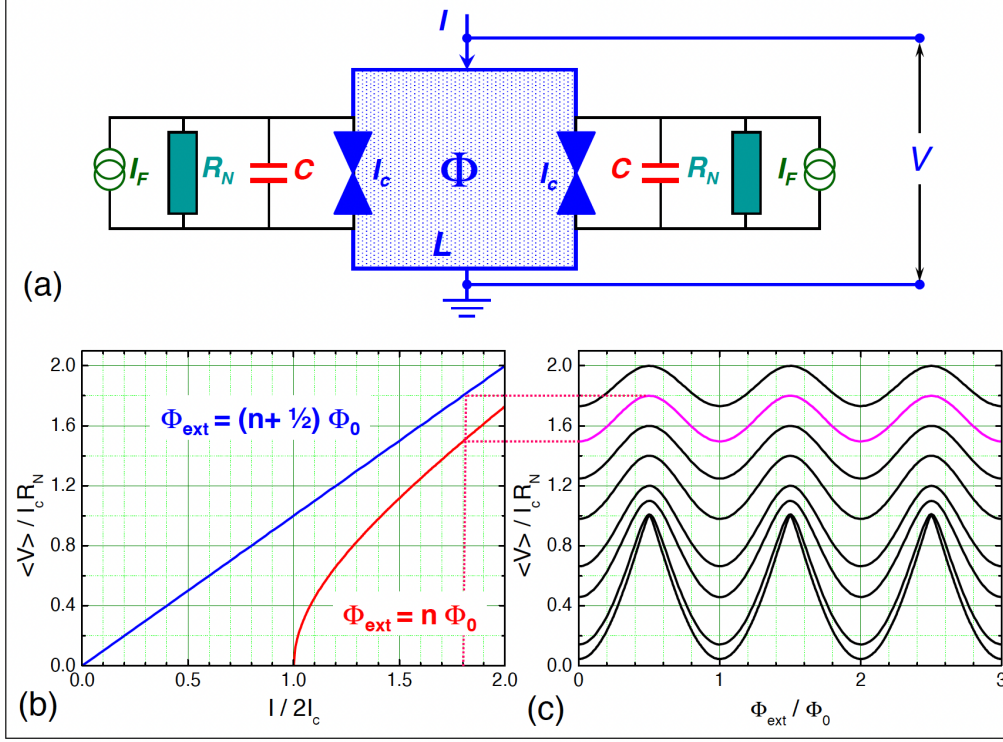


Figure 18: (a) Equivalent SQUID circuit using the RCSJ model. (b) Current-voltage characteristics for different flux values. (c) Flux-voltage characteristics for fixed bias currents. Figure from [4].

across the SQUID divided by the flux-to-voltage transfer coefficient square:

$$S_{\Phi}(f) = \frac{S_V(f)}{H^2} \quad (64)$$

- The **noise energy** associated to the equivalent flux noise, which is independent from the loop inductance L , being defined as:

$$\epsilon(f) = \frac{S_{\Phi}(f)}{2L} = \frac{S_V(f)}{2LH^2} \quad (65)$$

and it is therefore convenient to compare the performances of different SQUID with different geometries.

Using these figures of merit, we can deduce which are the best choices for our design values.

- **Bias current:** it should be just above $2I_c$ to maximise the modulation of the flux-voltage curve.
- **Flux bias:** for optimal bias current values, the flux bias should be close to $(2n + 1)\Phi_0/4$, to maximise H .
- **Critical current** of the Josephson junctions: it should be larger than the thermal noise current $I_{\text{th}} := 2\pi k_B T / \Phi_0$.
- **Loop inductance:** on the one hand, it should be large to maximise sensitivity; on the other hand, a large inductance L increases the thermal noise flux in the loop $\langle \Phi_{\text{th}}^2 \rangle^{1/2} = \sqrt{k_B T L}$.
- **Screening parameter** β_L : to avoid hysteresis in the flux-voltage curves, we need $\beta_L = 2I_c L / \Phi_0 \leq 1$. We can decrease L , but previously we stated that at the same time we need large L to increase sensitivity. Thus, we choose $\beta_L \approx 1$.

- **Screening parameter β_C :** the Stewart-McCumber parameter needs to be smaller than one to avoid hysteresis in the IV curves, i.e. $\beta_C = 2e/\hbar \cdot I_c R_N^2 C$. Tunnel junctions have intrinsically large capacitance C and hence $\beta_C \gg 1$. However, we can add shunt resistors $R_{\text{shunt}} \ll R_N$ to lower β_C . The limit is that too low values of R_{shunt} would lower the amplitude of the flux-voltage characteristics, replacing the pre-factor of equation 62 to $I_c R_{\text{shunt}} \ll I_c R_N$. Therefore, we choose $\beta_C \approx 1$.

8.5 SQUID designs

Practical SQUID system comprise different parts, as shown in figure 19. The blue circuit include the SQUID itself and the flux coil, used to set a flux working point. All these elements are typically on the same chip. The red circuit includes the antenna, or pick-up coil, i.e. the coil that convert the signal to flux and coupled it to the SQUID loop. The antenna, or pick-up coil, can be on the same chip of the SQUID or on a separate connected chip. Finally, the green part corresponds to the room temperature electronics, including two current generators to provide the bias current and the current to the flux coil, respectively, and an amplifier to amplify and read-out the voltage signal from the SQUID.

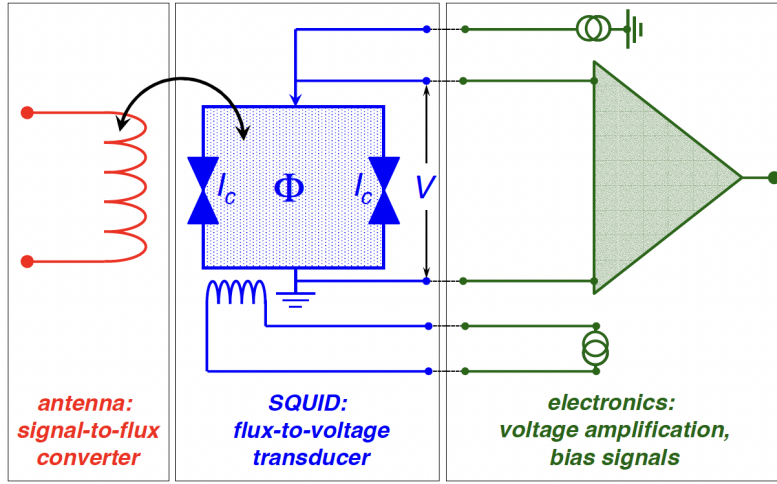


Figure 19: Figure from [4].

Examples of SQUID designs: see slides.

8.6 Applications

Dc SQUIDS are used for several applications, mainly as sensitive magnetometers and current sensors. Examples of magnetometry applications include medical diagnostics, such as magnetoencephalography (MEG) and magnetocardiography (MCG), to study the magnetic signals of our brains and hearts, respectively.

Another notable example of magnetometry application is geo-exploration and mineral surveys, where systems using SQUIDS are flying on helicopters to map the magnetic field on large portion of terrain.

(See slides.)

9 Cryogenic detectors

Building cryogenic detectors using superconducting materials and, more in general, exploiting low temperature physics brings several advantages:

- The **small energy gap** offers high sensitivity: Typical energy gap values for superconductors are in the order of meV. For instance, NbN features an energy gap of $\Delta = 1.76 \cdot k_B T_c = 2.5$ meV. Therefore, incoming particles with energy as small as ≥ 2.5 meV are able to break Cooper pairs in the material and give a measurable signal.
- The **small heat capacity** allows for sensitive calorimetric measurements. As we will discuss in section 9.2, cryogenic microcalorimeters measure the temperature increase ΔT due to an energy deposition E , with $\Delta T = E/C$, where C is the heat capacity of the detector. At cryogenic temperature, the heat capacity is suppressed, enhancing the temperature increase ΔT and, thus, the sensitivity of the calorimeter.
- The **sharp transition** of superconductors between normal state and superconducting state can be used to build precise thermometers, since a small change in temperature corresponds to a sizeable change in resistance.
- The dependence of the **kinetic inductance** on the Cooper pair density can be used to detect an incoming particle or an energy deposition that alters the Cooper pair density by breaking Cooper pairs.

These features are linked to the fundamental working principle of several cryogenic detectors: Superconducting Nanowire Single-Photon Detectors (SNPSDs), microcalorimeters, Transition Edge Sensors (TESs) and Kinetic Inductance Detectors (KIDs), respectively.

9.1 Superconducting Nanowire Single-Photon Detectors (SNPSDs)

(see slides)

9.2 Cryogenic microcalorimeters

By definition, a calorimeter is a detector that measures energy deposition. In a calorimeter, the particles to be measured are fully absorbed and their energy is transformed in a measurable quantity.

In high energy physics, the interaction of the incoming particle inside the calorimeter occurs via electromagnetic or strong processes. As a result, showers of secondary particles are generated and their energy is deposited in the active part of the calorimeter. The energy is detected in the form of ionisation (charges) or scintillation (light). A large fraction of the energy is deposited in form of heat and remains undetected. The energy resolution of high energy physics calorimeters is set by the statistical fluctuations of the number of charges or photoelectrons produced by the incoming particle. On the other hand, cryogenic microcalorimeter can measure the complete energy deposition, including the heat in form of phonons or quasi-particles in a superconductor. Therefore, cryogenic microcalorimeter can reach unprecedented energy resolutions, more than one order of magnitude better than semiconductor devices.

The development of cryogenic microcalorimeters has been driven mainly by the low energy particle physics community, for example for dark matter searches or neutrinoless double-beta decay experiments, where high sensitivity and extreme energy resolution are required.

Working principle and detector response

The microcalorimeter consists of an absorber, where the energy is deposited, which defines the interaction volume. The absorber is thermally well coupled to a thermometer, which measures the increase of temperature due to the energy deposition. The absorber is connected to a thermal bath via a weak thermal link with thermal conductance G . The thermal bath is kept at a base temperature, so that the temperature of the absorber is restored to the base temperature after an

event in the detector. Figure 20 shows a simple scheme of a microcalorimeter. The expected change in temperature ΔT is directly proportional to the deposited energy E and inversely proportional to the total heat capacity C : $\Delta T = E/C$. For an energy input $E = 3 \text{ keV}$, assuming a heat capacity $C = 1 \text{ pJ/K}$, the expected increase in temperature is $\Delta T = 0.5 \text{ mK}$.

The total heat capacity is given by $C = c \cdot V$, where c is the specific heat capacity and V is the volume. In order to maximise the temperature change ΔT , the capacity should be minimised. This is achieved by reducing the volume of the detector (thus "microcalorimeters") and by operating it at low temperature.

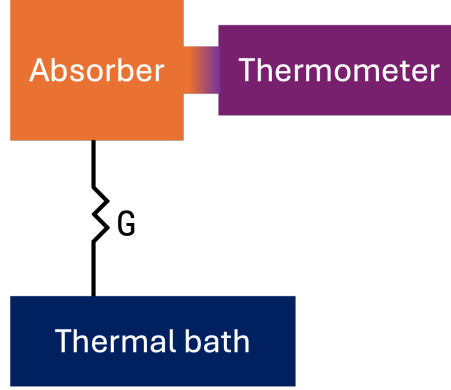


Figure 20: Scheme of a microcalorimeter consisting of absorber, thermometer and weak thermal link with thermal bath.

When an event occurs, i.e. a particle or a photon is depositing energy in the absorber, the interaction generates atomic and solid state excitations. As a consequence, electrons, photons (photoelectrons) and phonons are created. Their energy will degrade in time via electron-phonon interactions, phonon-phonon interactions and interactions with lattice irregularities. Eventually, thermal equilibrium will be reached.

Microcalorimeters can be operated

- in **equilibrium mode**: the detector is sensitive to thermal phonons and the best energy resolution can be achieved, thanks to the large number of thermal phonons with energies in the meV range;
- in **non-equilibrium mode**: the detector is sensitive only to quasi-ballistic phonons, i.e. high energy phonons, reaching a faster response, but featuring a poorer energy resolution.

Absorber

The role of the absorber is to absorb the input energy and therefore a sufficient stopping power is typically required, depending on the specific application. The material of the absorber can be dielectric, metal or superconducting.

Dielectric absorber According to the Debye model of a crystal, the specific heat capacity is given by the phonon specific heat:

$$c_{\text{diel}} = \beta \left(\frac{T}{\theta_D} \right)^3 = c_{\text{ph}} \quad (66)$$

where $\beta = 1944 \text{ Jmol}^{-1}\text{K}^{-1}$ and θ_D is the Debye temperature.

Metallic absorber The heat capacity of a metal comprises two contributions, namely the phononic contribution and the electronic contribution. The latter is due to thermally excited conduction electrons:

$$c_{\text{metal}} = \beta \left(\frac{T}{\theta_D} \right)^3 + \gamma T = c_{\text{ph}} + c_e \quad (67)$$

where γ is the Sommerfeld coefficient. At low temperatures, < 1 K, the electronic specific heat dominates, linearly decreasing with the temperature.

Superconducting absorber The specific heat includes the phononic contribution and a second term due to the thermally excited electrons across the energy gap Δ :

$$c_{\text{sc}} = \beta \left(\frac{T}{\theta_{\text{D}}} \right)^3 + a e^{-\frac{b\Delta}{k_{\text{B}}T}} \quad (68)$$

where a and b are material constants. The number of thermally excited electrons exponentially decrease with the temperature, because of the decrease of quasi-particles density. Thus, at very low temperatures, the phononic contribution dominates.

Detector response

The response of an ideal microcalorimeter upon an energy deposition E consists of a temperature signal with amplitude $\Delta T = E/C$, a fast rise time τ_r and a slow decay time $\tau_d = C/G$, which depends on the heat capacity C and the conductance G of the weak thermal link connecting the system to the thermal bath. Figure X shows a typical signal of a microcalorimeter. The different types of microcalorimeters mainly different from the mechanisms that translates the temperature increase into a measurable voltage (or current) signal.

Energy resolution

The intrinsic energy resolution of a microcalorimeter is limited by the thermal energy fluctuations due to the phonon exchange between the absorber and the thermal bath. The mean square energy fluctuations are given by $\langle \Delta E^2 \rangle = k_{\text{B}} T^2 C$. Assuming an effective number $N = C/k_{\text{B}}$ of phonon modes in the detector, a typical mean energy of one photon equal to $k_{\text{B}} T$ and the RMS fluctuation of one photon equal to 1, we can in fact intuitively write the mean square energy fluctuations as $N(k_{\text{B}} T)^2 = k_{\text{B}} T^2 C$.

For a practical microcalorimeter device, the energy resolution is given by:

$$\Delta E_{\text{FWHM}} = 2 \log 2 \xi \sqrt{k_{\text{B}} T^2 C} \quad (69)$$

where ξ depends on the sensitivity and noise characteristics of the device, and it typically assumes values between 1.0 and 1.2. Typical values of energy resolution can reach 1 to 2 eV for an energy input of 6 keV.

For a superconducting detector, the detection mechanism is different. In this case, the energy deposition E in the superconducting material with energy gap Δ produces $N = E/\Delta$ quasi-particles by breaking Cooper pairs as well as phonons. As long as the energy of these generated particles is larger than 2Δ (i.e. larger than the binding energy of one Cooper pair), they break more Cooper pairs in a cascade fashion, until their energy drops below the threshold ($< 2\Delta$). The usage of superconductors as detectors is motivated by the small energy gap and, thus, by their high sensitivity.

The energy resolution is given by:

$$\Delta E_{\text{FWHM}} = 2 \log 2 \sqrt{\Delta F E} \quad (70)$$

where the energy gap Δ is typically around 1 meV and the Fano factor F , taking into account the deviations from Poisson statistics in the generation of quasi-particles, is about 0.2. These values give an energy resolution of $\Delta E_{\text{FWHM}} = 2.6$ eV at 6 keV.

9.3 Transition Edge Sensors (TESs)

A Transition Edge Sensor (TES) is a phonon sensor which can be used as a cryogenic microcalorimeter. It consists of a superconducting film operated at a temperature within the transition region between the superconducting and the normal phase. Here, the TES's resistance changes between zero and its normal value, as shown in figure 21a. The strong dependence of the resistance change on temperature $R(T)$ is expressed using the dimensionless parameter α which is related to the sensitivity and can reach values up to 1000:

$$\alpha = \frac{d(\log R)}{d(\log T)} = \frac{T}{R} \frac{dR}{dT} \quad (71)$$

In a typical microcalorimeter configuration, a TES sensor is attached to an absorber. However, TESs can also be used as absorber and sensor at the same time. The input energy is changing the temperature of the superconducting film, thus strongly affecting its resistance.

TES detector can be operated in the *current bias* mode or in the *voltage bias* mode.

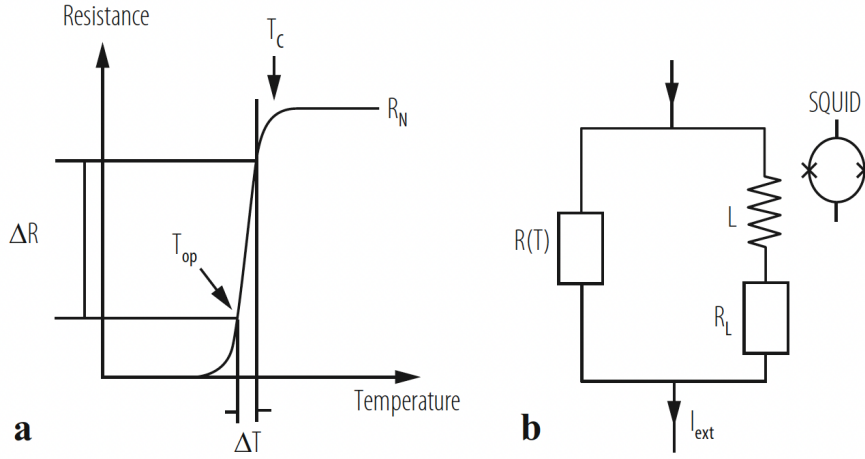


Figure 21: (a) Temperature as a function of resistance of a superconducting film close to the transition temperature. (b) dc-SQUID readout of a transition edge sensor in current bias in current bias. Figure from [7].

Current bias mode

In the current bias operation mode, a constant bias current is fed in the circuit, as shown in figure 21b. An energy input in the absorber connected to the TES sensor, or directly in the TES sensor, increases the temperature and thus the resistance. Because of this change of resistance, a larger current will flow in the parallel branch of the circuit. As a consequence, the magnetic flux will change in the inductance L , which is measured using a dc SQUID.

The main disadvantage of the current bias operation mode is Joule heating. It causes small fluctuations in the bath temperature, which limit the detector's performance. To solve this problem, the so called auto-biasing electro-thermal feedback (ETF) approach is introduced, operating the TES system in voltage mode.

Voltage bias mode

In the voltage bias mode, a constant voltage V_b is kept. A temperature rise in the TES causes an increase in its resistance and a corresponding decrease in the current ΔI , which results in a decrease of the Joule heating $V_b \cdot \Delta I$. In this way, the temperature of the superconducting film is brought back to the constant operating value, stabilising the operation of the detector and self-calibrating the system. The deposited energy E is reconstructed knowing the voltage bias and the integral of the current change, which is typically measured using a SQUID current amplifier:

$$E = V_{\text{b}} \int \Delta I(t) dt. \quad (72)$$

Energy resolution

The intrinsic energy resolution for an ideal TES calorimeter is given by the expression previously introduced:

$$\Delta E_{\text{FWHM}} = 2 \log 2 \xi \sqrt{k_{\text{B}} T^2 C} \quad (73)$$

with

$$\xi = 2 \sqrt{\frac{1}{\alpha}} \cdot \left(\frac{n}{2}\right)^{1/4} \quad (74)$$

with n depending on the thermal impedance between the phonons (absorber) and the electrons in the superconductor. A typical value is $n = 5$, for thin films, which gives an energy resolution of about 2 eV FWHM at 6 keV.

9.4 Magnetic Microcalorimeters (MMCs)

... Coming soon ...

9.5 Kinetic Inductance Detectors (KIDs)

... Coming soon ...

10 Superconducting parametric amplifiers

10.1 Introduction

Amplifiers act as bridges between the quantum circuit sitting at cryogenic temperatures and the classic apparatus in the laboratory at room temperature that we intend to use for the read-out. In 1982, C.M. Caves was writing: “The last essential quantum mechanical stage of a measuring apparatus is a high-gain amplifier; it produces an output that we can lay our grubby, classical hands on.”.

When measuring a superconducting quantum circuit, we need to route the weak microwave output signal of the device from the cryogenic stage to the room temperature electronics and we need to amplify it. Typically, we use several amplification stages in cascade, as shown in the sketch of figure 22, and each amplifier inevitably adds noise.

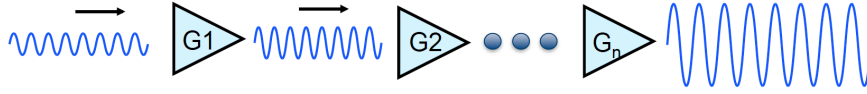


Figure 22: Simple sketch of a cascade of amplifiers in series, each one characterised by a gain G_i .

The total system noise temperature T_{sys} of a chain of M amplifiers in series is given by the Friis formula:

$$T_{\text{sys}} = \sum_{i=1}^M \frac{T_{N,i}}{\prod_{j=1}^{i-1} G_j} = T_{N,1} + \frac{T_{N,2}}{G_1} + \frac{T_{N,3}}{G_1 G_2} + \dots \quad (75)$$

where the G_i is the gain of the i -th amplifier and $T_{N,i}$ is its noise temperature.

If the first amplifier has a large gain, i.e. $G_1 \gg 1$, the total system noise temperature is dominated by the noise performance of the first amplifier. Therefore, in order to maximise the signal-to-noise ratio, we need a first-stage-amplifier with a sufficiently high gain and as low noise as possible. Moreover, the insertion loss between the superconducting circuit and the amplifiers needs to be minimised. Thus, the amplifier is typically co-located with the superconducting circuit under test at the lowest temperature stage of the cryogenic apparatus, reducing as much as possible the number of contacts and components between the signal source and the amplifier.

Standard low-noise cryogenic amplifiers (e.g. HEMTs - High Electron Mobility Transistors) feature noise levels of a few Kelvin. Moreover, they dissipate too high power, if compared to the available cooling power in the cryogenic system. Ideally, we would like to use an amplifier with minimal power dissipation and with the lowest possible added noise, i.e. at the quantum noise limit.

10.2 Parametric amplification

Any non-linear system where we can periodically modulate a parameter of the system (e.g. a circuit reactance), converting energy between conjugate field variables of the system (e.g. voltage and current), can produce a parametric behaviour. The nonlinear circuit mediates the transfer of energy from the pump to the signal to be amplified via wave-mixing processes. As a consequence of this nonlinear interaction, an additional mode known as the *idler* is generated. Figure 23 shows signal and pump entering a nonlinear medium and being amplified, alongside with the creation of the idler.

The simplest example of parametric amplification is a child on a swing, as shown in figure 24(left), where the potential is at the first order quadratic in displacement. The child is squatting at twice the frequency of the swing oscillations, therefore modulating the position of the centre of mass and achieving parametric gain in the oscillation amplitude. The parametric periodic

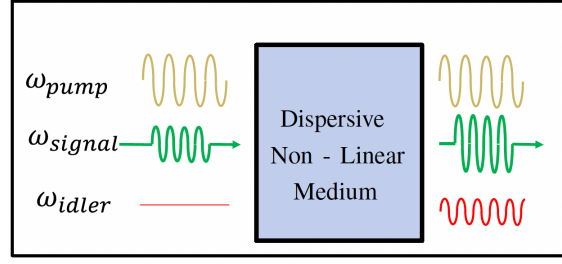


Figure 23: Wave mixing process leading to parametric amplification in a nonlinear medium.

modulation is often called *pumping* and the condition for parametric gain is $f_p \approx 2f_0$, where f_p is the pumping frequency and f_0 is the natural oscillator frequency.

Another curious example of parametric amplification is the famous Botafumeiro, a thurible used at the Santiago de Compostela Cathedral, in Spain. The Botafumeiro is suspended 20 m from a pulley mechanism under the dome on the roof of the church and it swings in a 65 m arc. The maximum amplitude of the oscillation is reached by parametric amplification, modulating the length of the suspension rope at twice the oscillation frequency. This modulation (*pumping*) effectively injects energy into the system, producing gain in the oscillation amplitude. Enjoy the video here: <https://www.youtube.com/watch?v=beP8N9X0nyw>.

We can implement parametric amplification in a non-linear LC resonator circuit, like the one sketched in figure 24(right), by periodically modulating the inductance or the capacitance of the system.

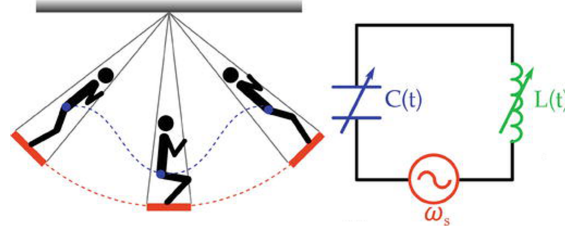


Figure 24: Two analogous examples of parametric amplification: the child on the swing, amplifying the oscillation amplitude by pumping the moment of inertia (left) and an LC resonator circuit, where the signal ω_s is amplified by pumping the magnetic flux in the inductor $L(t)$ or the electrical charge in the capacitor $C(t)$.

Depending on the order of the nonlinearity, parametric amplification can occur via three-wave mixing (3WM) or four-wave mixing (4WM). In 3WM, one pump photon at frequency ω_p is annihilated to create one signal photon at frequency ω_s and one idler photon at frequency ω_i . The following relation ensures energy conservation:

$$\omega_p = \omega_s + \omega_i . \quad (76)$$

In 4WM, instead, two pump photons at frequency ω_p are annihilated to create one signal photon at frequency ω_s and one idler photon at frequency ω_i , such that:

$$2\omega_p = \omega_s + \omega_i . \quad (77)$$

In the optical domain, parametric amplification is typically achieved by pumping with an intense laser a nonlinear crystal, whose the refractive index depends on the amplitude of the electromagnetic field. The pump modulates the the wave velocity in the medium and energy is transferred to the signal mode. In the microwave domain, superconducting circuits are typically used and the nonlinearity arises either from Josephson junctions or from the kinetic inductance of thin films. The microwave pump modulates the inductance of the circuit.

10.3 Phase-preserving and phase-sensitive amplification

In general, a (classical) microwave signal can be represented by two orthogonal quadratures I and Q :

$$V(t) = I(t) \cos(\omega t) + Q(t) \sin(\omega t) \quad (78)$$

commonly named *in-phase* (I) and *quadrature* (Q), respectively, as shown in figure 25.

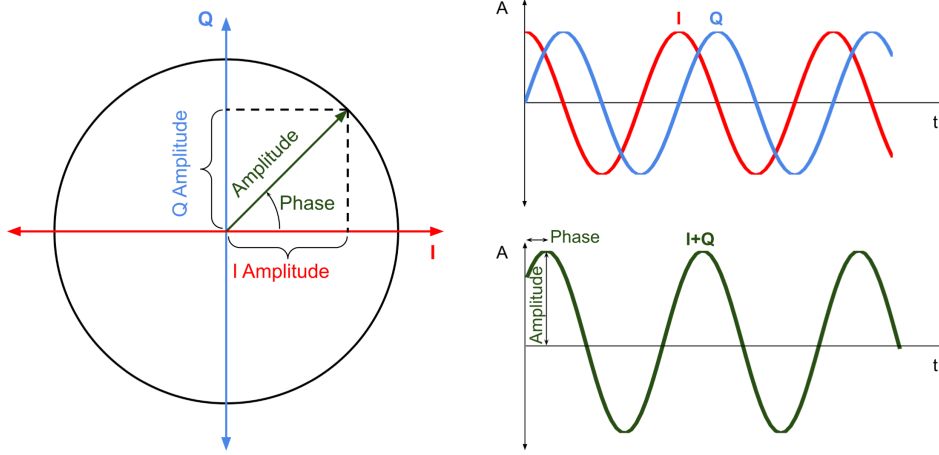


Figure 25: In-phase and quadrature components of a sinusoidal signal in the complex I-Q plane (left) and in time domain (right). Source: Wikipedia.

In superconducting electronics, in agreement with the quantum optics formalism, we can write the field annihilation and creation operators as:

$$\begin{aligned} \hat{a} &= \hat{X} + i\hat{P} \\ \hat{a}^\dagger &= \hat{X} - i\hat{P} \end{aligned} \quad (79)$$

where $\hat{X}$ and $\hat{P}$ are quadrature operators. Then, the microwave field operator $\hat{E}$ takes the form:

$$\hat{E}(t) = E_0 (\hat{a}e^{-i\omega t} + \hat{a}^\dagger e^{i\omega t}) . \quad (80)$$

The treatment of these quadratures in an amplification process determines whether amplification is phase-preserving or phase-sensitive. If we consider an amplifier with at its input a quantum field $\hat{a}_{\text{in}}$, the field at the amplifier output $\hat{a}_{\text{out}}$ is given by:

$$\hat{a}_{\text{out}} = \sqrt{G}\hat{a}_{\text{in}} \quad (81)$$

where G the amplifier power gain.

In order to fulfill the commutation relations $[\hat{a}_{\text{in}}, \hat{a}_{\text{in}}^\dagger] = [\hat{a}_{\text{out}}, \hat{a}_{\text{out}}^\dagger] = 1$, an extra degree of freedom must be added to equation 81. If this degree of freedom is given by the signal phase, the amplifier is phase-sensitive, otherwise the amplifier is phase-preserving.

Phase preserving amplifiers. In a phase preserving amplification process, both quadratures of an input signal are simultaneously and equally amplified by $\sqrt{G}$ and the amplifier intrinsic noise is also amplified [3], as shown in figure 26a.

The input-output relation requires an additional operator for the reason discussed above:

$$\hat{a}_{\text{out}} = \sqrt{G}\hat{a}_{\text{in}} + \hat{F}. \quad (82)$$

Here, the operator $\hat{F}$ represent the added noise of the amplifier. It is uncorrelated to the input mode $\hat{a}_{\text{in}}$ and it has zero mean:

$$\begin{aligned} [\hat{a}_{\text{in}}, \hat{F}] &= [\hat{a}_{\text{in}}^\dagger, \hat{F}] = 0, \\ \langle \hat{F} \rangle &= 0. \end{aligned} \quad (83)$$

Imposing the commutation relations on $\hat{a}$ and $\hat{a}^\dagger$, we have that $[\hat{F}, \hat{F}^\dagger] = 1 - G$.

A phase-preserving amplifier must add at least half a quantum of noise to the input signal, as we will show in the following. This is commonly called *standard quantum limit* (SQL).

We can quantify the output noise considering the variance of the output mode:

$$\begin{aligned} (\Delta \hat{a}_{\text{out}})^2 &= G(\Delta \hat{a}_{\text{in}})^2 + \frac{1}{2} \langle \{\hat{F}, \hat{F}^\dagger\} \rangle \\ &\geq G(\Delta \hat{a}_{\text{in}})^2 + \frac{1}{2} |\langle \hat{F}, \hat{F}^\dagger \rangle| \\ &= G(\Delta \hat{a}_{\text{in}})^2 + \frac{|G - 1|}{2}. \end{aligned} \quad (84)$$

Referring the output noise to the input, we have that:

$$\frac{(\Delta \hat{a}_{\text{out}})^2}{G} \geq (\Delta \hat{a}_{\text{in}})^2 + \frac{|1 - 1/G|}{2}. \quad (85)$$

Thus, in the limit of no gain, $G = 1$, the amplifier does not add noise, while in the large gain limit, $G \gg 1$, half a quantum of noise must be added:

$$\frac{(\Delta \hat{a}_{\text{out}})^2}{G} \geq (\Delta \hat{a}_{\text{in}})^2 + \frac{1}{2}. \quad (86)$$

In a phase-sensitive amplifier, the extra degree of freedom is given by the only other internal mode $\hat{b}$, that fulfills the commutation relation $[\hat{b}, \hat{b}^\dagger] = 1$. This is the **idler mode**. Therefore, the noise operator $\hat{F}$ can be written as:

$$\begin{aligned} \hat{F} &= \sqrt{G - 1}\hat{b}, \\ \hat{F}^\dagger &= \sqrt{G - 1}\hat{b}^\dagger. \end{aligned} \quad (87)$$

The input-output relation yields to:

$$\hat{a}_{\text{out}} = \sqrt{G}\hat{a}_{\text{in}} + \sqrt{G - 1}\hat{b}^\dagger. \quad (88)$$

The noise added by the amplifier directly arises from the idler mode. If the added noise is exactly half quantum, i.e. it is at the standard quantum limit, the amplifier is said to be a *phase-preserving, quantum-limited amplifier*.

Phase sensitive amplifiers. In a phase sensitive amplification process, the phase space volume is conserved. This happens when the amplifier is sensitive to the input signal phase. In a phase sensitive amplifier, the signal and idler modes are temporally and spatially degenerate, $\omega_s = \omega_i$. The only remaining degree of freedom to fulfill the commutation relation is the signal phase ϕ (with respect to the pump phase).

The input-output relation becomes:

$$\hat{a}_{\text{out}} = \sqrt{G}\hat{a}_{\text{in}} + e^{i\phi}\sqrt{G - 1}\hat{a}_{\text{in}}^\dagger. \quad (89)$$

Writing this relation in terms of signal quadratures $\hat{X}$ and $\hat{P}$, in the large-gain limit $G \gg 1$, we have:

$$\begin{aligned}\hat{X}_{\text{out}}^{\phi=0} &= \sqrt{G}\hat{X}_{\text{in}}, \\ \hat{P}_{\text{out}}^{\phi=0} &= \frac{1}{\sqrt{G}}\hat{P}_{\text{in}}.\end{aligned}\tag{90}$$

A phase-sensitive amplifier amplifies only one quadrature of the input field, while the orthogonal quadrature is simultaneously attenuated, as shown in figure 26b. Same happens for the noise, where one quadrature is *squeezed* below the input noise while the other is "stretched" above it. If the input noise is originally dictated by quantum fluctuations, the output noise of one quadrature is reduced below the SQL, but it does not contradict Heisenberg principle, since the other quadrature noise is amplified.

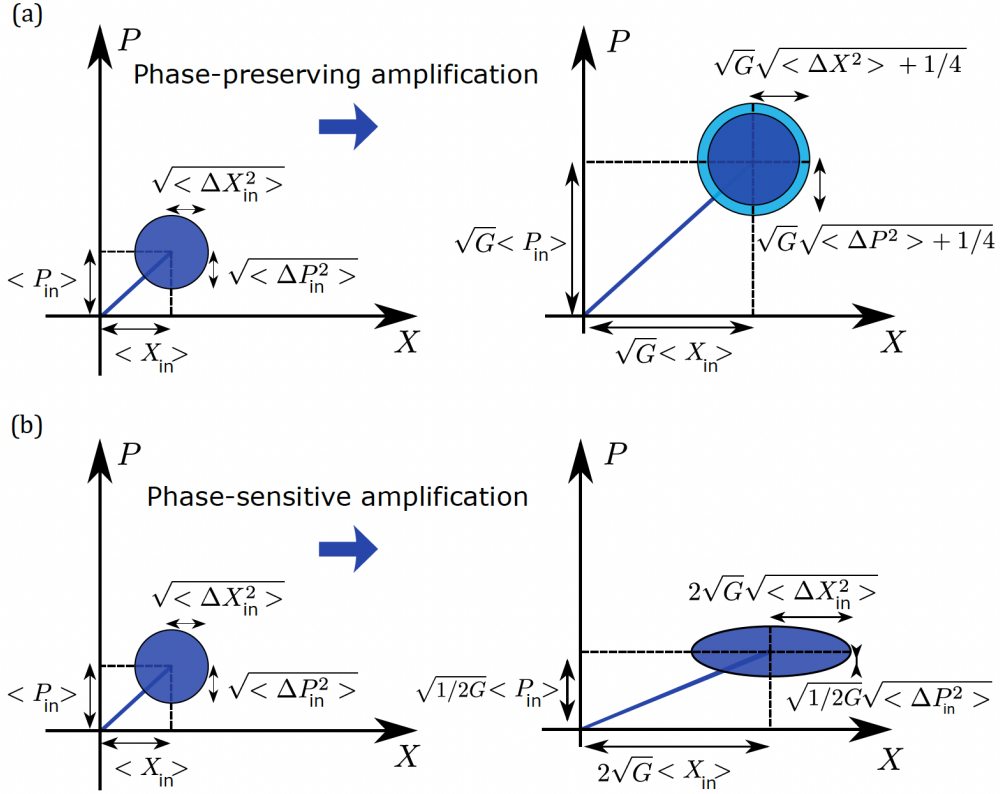


Figure 26: Contour of Wigner distribution of input and output states. The light blue contour shows the amplification of the amplifier intrinsic noise. (a) Phase-preserving amplification. (b) Phase sensitive amplification. Figure from [5].

10.4 Resonant parametric amplifiers

... coming soon ...

Josephson Parametric Amplifier

... coming soon ...

Kinetic Inductance Parametric Amplifier

... coming soon ...

10.5 Travelling Wave Parametric Amplifiers (TWPAs)

... coming soon ...

Josephson TWPAs

... coming soon ...

Kinetic Inductance TWPAs

... coming soon ...

11 Superconducting qubits

... Coming soon ...

Suggested Textbooks and Reviews

A. Blais, A. L. Grimsmo, S. M. Girvin and A. Wallraff, *Circuit quantum electrodynamics*, Rev. Mod. Phys. 93, 025005, 2021

Nathan K. Langford, *Circuit QED - Lecture Notes*, arXiv:1310.1897, 2013

Pozar, D. M. Wiley, Hoboken, NJ, *Microwave engineering*, 4th ed edition, 2012

Yvonne Y. Gao, M. Adriaan Rol, Steven Touzard, and Chen Wang, *Practical Guide for Building Superconducting Quantum Devices*, PRX Quantum 2, 040202, 2021

R. Gross and A. Marx, "Applied Superconductivity: Josephson Effect and Superconducting Electronics", 2005

Stolz, R., Schiffler, M., Becken, M. et al., "SQUIDS for magnetic and electromagnetic methods in mineral exploration", Miner Econ 35, 467–494(2022)

R.H. Hadfield, G. Johansson, Superconducting Devices in Quantum Optics, 2016. Part of Quantum Science and Technology, p. 3-30.

K. Pretzl, Cryogenic Detectors, (Bern U., LHEP), 2020. Part of Particle Physics Reference Library. Volume 2: Detectors for Particles and Radiation, p. 871-912

References

- [1] Felix Ahrens. "Cryogenic read-out system and resonator optimisation for the microwave SQUID multiplexer within the ECHO experiment". PhD thesis. Heidelberg University, 2022.
- [2] Alexandre Blais et al. "Circuit quantum electrodynamics". In: *Rev. Mod. Phys.* 93 (2 May 2021), p. 025005. DOI: 10.1103/RevModPhys.93.025005. URL: <https://link.aps.org/doi/10.1103/RevModPhys.93.025005>.
- [3] Carlton M. Caves. "Quantum limits on noise in linear amplifiers". In: *Phys. Rev. D* 26 (8 Oct. 1982), pp. 1817–1839. DOI: 10.1103/PhysRevD.26.1817. URL: <https://link.aps.org/doi/10.1103/PhysRevD.26.1817>.
- [4] R. Gross and A Marx. *Applied Superconductivity: Josephson Effect and Superconducting Electronics*. chapter 4. 2005.
- [5] Luca Planat. "Resonant and traveling-wave parametric amplification near the quantum limit". PhD thesis. Université Grenoble Alpes, 2020.
- [6] David M. Pozar. *Microwave engineering*. 4th ed. Includes index. New Delhi : Wiley, 2012.
- [7] Klaus Pretzl. "Cryogenic Detectors". In: *Particle Physics Reference Library. Volume 2: Detectors for Particles and Radiation*. Ed. by Christian Wolfgang Fabjan and Herwig Schopper. 2020, pp. 871–912. DOI: 10.1007/978-3-030-35318-6_19.