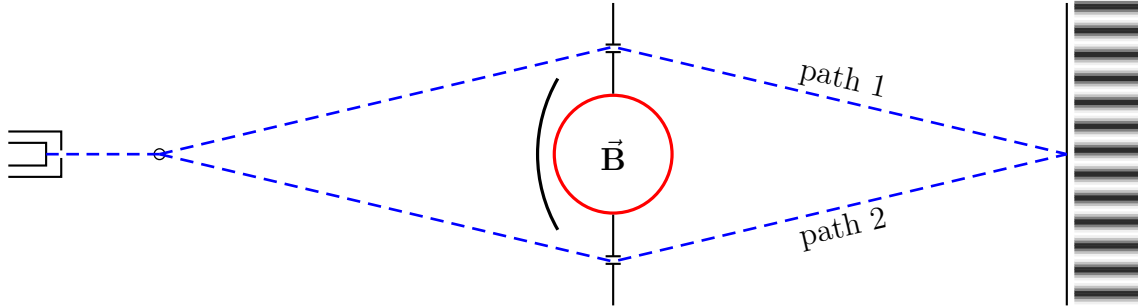


Aharonov–Bohm Effect and SQUIDS

AHARONOV–BOHM EFFECT

In classical mechanics, the motion of a charged particle depends only on the electric and magnetic tension fields $\mathbf{E}$ and $\mathbf{B}$; the potentials A^0 and $\mathbf{A}$ do not have any direct effect. Also, the motion depends only on the $\mathbf{E}$ and $\mathbf{B}$ fields along the particle's trajectory — the EM fields in some volume of space the particle never goes through do not affect it at all. But in *quantum mechanics*, the *interference between two trajectories a charged particle might take depends on the magnetic field between the trajectories, even if along the trajectories themselves $\mathbf{B} = 0$* . This effect was first predicted by Werner Ehrenberg and Raymond E. Siday in 1949, but their paper was not noticed until the effect was re-discovered theoretically by David Bohm and Yakir Aharonov in 1959 and then confirmed experimentally by R. G. Chambers in 1960.

Consider the following idealized experiment: Take a two-slit electron interference setup, and put a solenoid between the two slits as shown below:



The solenoid is thin, densely wound, and very long, so the magnetic field outside the solenoid is negligible. Inside the solenoid there is a strong $\mathbf{B}$ field, but the electrons do not go there; instead, they fly outside the solenoid along paths 1 and 2. But despite $\mathbf{B} = 0$ along both paths, the magnetic flux Φ inside the solenoid affects the interference pattern between the two paths.

The key to the Aharonov–Bohm effect is the vector potential $\mathbf{A}$. Outside the solenoid $\mathbf{B} = \nabla \times \mathbf{A} = 0$ but $\mathbf{A} \neq 0$ because for any closed loop surrounding the solenoid we have a

non-zero integral

$$\oint_{\text{loop}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} = \oint_{\substack{\text{inside the loop} \\ \text{including the solenoid}}} \mathbf{B}(\mathbf{x}) \cdot d^2\mathbf{Area} = F, \quad (1)$$

the magnetic flux through the solenoid. (Technically, F is the magnetic flux through the whole loop surrounding the solenoid, but since the $\mathbf{B}$ field outside the solenoid is negligible, the flux F comes from the solenoid itself.)

Locally, a curl-less vector potential is a gradient of some function, so it (the vector potential $\mathbf{A}(\mathbf{x})$) can be removed by a gauge transform,

$$\mathbf{A}(\mathbf{x}) \rightarrow \mathbf{A}'(\mathbf{x}) = \mathbf{A}(\mathbf{x}) + \nabla\Lambda(\mathbf{x}) = 0 \quad \text{for some } \Lambda(\mathbf{x}), \quad (2)$$

but *globally* no single-valued $\Lambda(\mathbf{x})$ can gauge away the vector potential along both paths around the solenoid. Instead, we have two separate gauge transforms — the $\Lambda_1(\mathbf{x})$ that gauges away $\mathbf{A}(\mathbf{x})$ along the path #1, and the $\Lambda_2(\mathbf{x})$ that gauges away $\mathbf{A}(\mathbf{x})$ along the path #2 — but they are different transforms, $\Lambda_1 \neq \Lambda_2$. To see how this works, let $\mathbf{x}_g$ be the electron gun's location while $\mathbf{x}_s$ is some point on the screen. *Along path #1* from $\mathbf{x}_g$ to $\mathbf{x}_s$,

$$d\Lambda_1(\mathbf{x}) = -\mathbf{A}(\mathbf{x}) \cdot d\mathbf{x}, \quad (3)$$

hence

$$\Lambda_1(\mathbf{x}_s) - \Lambda_1(\mathbf{x}_g) = \int_{\text{path\#1}} d\Lambda_1 = - \int_{\text{path\#1}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x}. \quad (4)$$

Likewise, *along path #2* from the same $\mathbf{x}_g$ to the same $\mathbf{x}_s$,

$$d\Lambda_2(\mathbf{x}) = -\mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \quad (5)$$

and hence

$$\Lambda_2(\mathbf{x}_s) - \Lambda_2(\mathbf{x}_g) = \int_{\text{path\#2}} d\Lambda_2 = - \int_{\text{path\#2}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x}. \quad (6)$$

However, the integrals in eq. (4) and (6) are not equal to each other; instead

$$\int_{\text{path\#1}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} - \int_{\text{path\#2}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} = \oint_{\mathcal{L}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \quad (7)$$

where $\mathcal{L}$ is the closed loop made from path#1 from the electron gun $\mathbf{x}_g$ to the point $\mathbf{x}_s$ on the screen and then path#2 in reverse from $\mathbf{A}(\mathbf{x}_s)$ back to the electron gun $\mathbf{x}_g$. By the Stokes theorem, the loop integral (7) is the magnetic flux through the loop $\mathcal{L}$, and since $\mathcal{L}$ surrounds the solenoid

$$\begin{aligned} \int_{\text{path\#1}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} - \int_{\text{path\#2}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} &= \oint_{\mathcal{L}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \\ &= F[\text{through } \mathcal{L}] \\ &= F[\text{through the solenoid}]. \end{aligned} \quad (8)$$

Consequently,

$$(\Lambda_1(\mathbf{x}_s) - \Lambda_1(\mathbf{x}_g)) - (\Lambda_2(\mathbf{x}_s) - \Lambda_2(\mathbf{x}_g)) = -F \neq 0, \quad (9)$$

which means we cannot possibly have the same single-valued $\Lambda_1(\mathbf{x}) \equiv \Lambda_2(\mathbf{x})$ gauge parameter for both paths.

The other key to the Aharonov–Bohm effect is the the local phase transform of the charged particle’s wave function which must accompany the gauge transform of the vector potential,

$$\left. \begin{aligned} \Psi'(\mathbf{x}) &= \exp\left(\frac{iq}{\hbar}\Lambda(\mathbf{x})\right) \Psi(\mathbf{x}) \\ \mathbf{A}'(\mathbf{x}) &= \mathbf{A}(\mathbf{x}) - \nabla\Lambda(\mathbf{x}) \end{aligned} \right\} \text{ for the same } \Lambda(\mathbf{x}). \quad (10)$$

Let’s translate this local phase transform of the wave function to the language of the propagation amplitude (AKA the evolution kernel) $U(\mathbf{x}_2, \mathbf{x}_1)$ from one point $\mathbf{x}_1$ to another point $\mathbf{x}_2$. For example from the electron gun $\mathbf{x}_1 = \mathbf{x}_g$ to some particular point $\mathbf{x}_2 = \mathbf{x}_s$ on the

screen. By definition, the propagation amplitude during flight time t is

$$U(\mathbf{x}_2, \mathbf{x}_1) \stackrel{\text{def}}{=} \langle \mathbf{x}_2 | \exp(-it\hat{H}/\hbar) | \mathbf{x}_1 \rangle, \quad (11)$$

$$\begin{aligned} & \Downarrow \\ \Psi(\mathbf{x}_2, t_2 = t) &= \int U(\mathbf{x}_2, \mathbf{x}_1) \Psi(\mathbf{x}_1, t_1 = 0) d^3\mathbf{x}_1. \end{aligned} \quad (12)$$

When a gauge transform is accompanied by a local phase transform of the wave function as in eq. (10), the propagation amplitude also changes its phase. Indeed, in order to keep eq. (12) working in a new gauge, we need

$$U'(\mathbf{x}_2, \mathbf{x}_1) = \exp\left(+i\frac{q}{\hbar}\Lambda(\mathbf{x}_2)\right) \times U(\mathbf{x}_2, \mathbf{x}_1) \times \exp\left(-i\frac{q}{\hbar}\Lambda(\mathbf{x}_1)\right). \quad (13)$$

where the first phase factor changes the phase of the $\Psi(\mathbf{x}_2, t_2 = t)$ while the second phase factor compensates for the changed phase of the $\Psi(\mathbf{x}_1, t_1 = 0)$, thus

$$\begin{aligned} \Psi'(\mathbf{x}_2, t) &= \int U'(\mathbf{x}_2, \mathbf{x}_1) \times \Psi'(\mathbf{x}_1, 0) d^3\mathbf{x}_1 \\ &= \int \exp\left(+i\frac{q}{\hbar}\Lambda(\mathbf{x}_2)\right) U(\mathbf{x}_2, \mathbf{x}_1) \exp\left(-i\frac{q}{\hbar}\Lambda(\mathbf{x}_1)\right) \times \exp\left(+i\frac{q}{\hbar}\Lambda(\mathbf{x}_1)\right) \Psi(\mathbf{x}_1, 0) d^3\mathbf{x}_1 \\ &= \exp\left(+i\frac{q}{\hbar}\Lambda(\mathbf{x}_2)\right) \times \int U(\mathbf{x}_2, \mathbf{x}_1) \Psi(\mathbf{x}_1, 0) d^3\mathbf{x}_1 \\ &= \exp\left(+i\frac{q}{\hbar}\Lambda(\mathbf{x}_2)\right) \times \Psi(\mathbf{x}_2, t). \end{aligned} \quad (14)$$

In particular, suppose $\mathbf{B} \equiv 0$ along the electron's path from $\mathbf{x}_1$ to $\mathbf{x}_2$ but the vector potential does not vanish, $\mathbf{A} \neq 0$. Then *locally* the vector potential is gauge-equivalent to zero, meaning there exist some $\Lambda(\mathbf{x})$ such that

$$\mathbf{A}_0(\mathbf{x}) = \mathbf{A}(\mathbf{x}) + \nabla\Lambda(\mathbf{x}) = 0, \quad (15)$$

if not everywhere then at least throughout the neighborhood of the electron's path. Then comparing the propagation amplitude $U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1)$ in presence of the vector potential with

the similar amplitude $U_0(\mathbf{x}_2, \mathbf{x}_1)$ for $\mathbf{A}_0 \equiv 0$, we find

$$\begin{aligned}
U_0(\mathbf{x}_2, \mathbf{x}_1) &= U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1) \times \exp \left(\frac{iq}{\hbar} (\Lambda(\mathbf{x}_2) - \Lambda(\mathbf{x}_1)) \right) \\
&= U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1) \times \exp \left(\frac{iq}{\hbar} \int_{\mathbf{x}_1}^{\mathbf{x}_2} \nabla \Lambda \cdot d\mathbf{x} \right) \\
&= U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1) \times \exp \left(\frac{-iq}{\hbar} \int_{\mathbf{x}_1}^{\mathbf{x}_2} \mathbf{A} \cdot d\mathbf{x} \right),
\end{aligned} \tag{16}$$

and therefore

$$U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1) = U_0(\mathbf{x}_2, \mathbf{x}_1) \times \exp \left(\frac{iq}{\hbar} \int_{\mathbf{x}_1}^{\mathbf{x}_2} \mathbf{A} \cdot d\mathbf{x} \right). \tag{17}$$

Thus, even when the vector potential $\mathbf{A}$ does not lead to a magnetic field in the region the electron travels through, it still manages to change the phase of its propagation amplitude.

Note: if the $\mathbf{B}$ field vanishes along the electron's path but does not vanish somewhere else, then we can make the gauge-transformed potential $\mathbf{A}' = \mathbf{A} + \nabla \Lambda$ vanish along the path, but it would not vanish somewhere else. Consequently, the relation

$$\Lambda(\mathbf{x}_2) - \Lambda(\mathbf{x}_1) = \int_{\mathbf{x}_1}^{\mathbf{x}_2} \nabla \Lambda \cdot d\mathbf{x} = - \int_{\mathbf{x}_1}^{\mathbf{x}_2} \mathbf{A} \cdot d\mathbf{x}$$

works only if we integrate $\mathbf{A} \cdot d\mathbf{x}$ along the electron path rather than some other line. In the context of eq. (17), this means that

$$U_{\mathbf{A}}(\mathbf{x}_2, \mathbf{x}_1) = U_0(\mathbf{x}_2, \mathbf{x}_1) \times \left(\frac{iq}{\hbar} \int_{\text{electron's path}} \mathbf{A} \cdot d\mathbf{x} \right). \tag{18}$$

In the Aharonov–Bohm experiment, the electron can take two different paths from the same point $\mathbf{x}_g$ (the electron gun) to the same point $\mathbf{x}_s$ on the screen. The interference pattern

on the screen follows from the net amplitude

$$U^{\text{net}}(\mathbf{x}_s, \mathbf{x}_g) = U^{\text{path 1}}(\mathbf{x}_s, \mathbf{x}_g) + U^{\text{path 2}}(\mathbf{x}_s, \mathbf{x}_g), \quad (19)$$

which depends on the phase difference between the amplitudes for each path,

$$\Delta\varphi(\mathbf{x}_s) = \text{phase}(U^{\text{path 1}}(\mathbf{x}_s, \mathbf{x}_g)) - \text{phase}(U^{\text{path 2}}(\mathbf{x}_s, \mathbf{x}_g)). \quad (20)$$

Note that along both paths $\mathbf{B} = 0$ but $\mathbf{A} \neq 0$, which affects the phases of the each amplitude according to eq. (18), specifically

$$\begin{aligned} \text{phase}(U_{\mathbf{A}}^{\text{path 1}}(\mathbf{x}_s, \mathbf{x}_g)) &= \text{phase}(U_0^{\text{path 1}}(\mathbf{x}_s, \mathbf{x}_g)) + \frac{q}{\hbar} \int_{\text{path 1}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x}, \\ \text{phase}(U_{\mathbf{A}}^{\text{path 2}}(\mathbf{x}_s, \mathbf{x}_g)) &= \text{phase}(U_0^{\text{path 2}}(\mathbf{x}_s, \mathbf{x}_g)) + \frac{q}{\hbar} \int_{\text{path 2}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x}. \end{aligned} \quad (21)$$

Consequently, the phase difference (20) is affected by the vector potential according to

$$\begin{aligned} \Delta\varphi_{\mathbf{A}} &= \Delta\varphi_0 + \frac{q}{\hbar} \int_{\text{path 1}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} - \frac{q}{\hbar} \int_{\text{path 2}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \\ &= \Delta\varphi_0 + \frac{q}{\hbar} \times F, \end{aligned} \quad (22)$$

where F is the magnetic flux through the solenoid, and the second equality follows from eq. (8).

For different points $\mathbf{x}_s$ on the screen we have different $\Delta\varphi_0(\mathbf{x}_s)$, that's why we see the interference pattern on the screen! The magnetic flux term in eq. (22) is the same for all points on the screen,

$$\Delta_{\mathbf{A}}\varphi(\mathbf{x}_s) = \Delta_0\varphi(\mathbf{x}_s) + \frac{q}{\hbar} \times F, \quad (23)$$

so it *shifts the whole interference patten along the screen!* Thus, **even though $\mathbf{B} = 0$ along both paths an electron might take from the gun to the screen, the quantum interference between the paths depends on the magnetic flux in the solenoid!**

In the mathematical language, the Aharonov–Bohm effect feels the *cohomology* of the vector potential $\mathbf{A}(\mathbf{x})$. In a topologically trivial space — like the flat 3D space without any holes — specifying $\mathbf{A}(\mathbf{x})$ *modulo* gauge transforms $\mathbf{A}(\mathbf{x}) \rightarrow \mathbf{A}(\mathbf{x}) + \nabla\Lambda(\mathbf{x})$ is equivalent to specifying the magnetic field $\mathbf{B}(\mathbf{x}) = \nabla \times \mathbf{A}$. However, in spaces with holes the vector potential *modulo* $\nabla\Lambda(\mathbf{x})$ for *single-valued* $\Lambda(\mathbf{x})$ contains more information than the magnetic field: In addition to $\mathbf{B}(\mathbf{x})$ for $\mathbf{x}$ outside the holes, the vector potential also knows the magnetic fluxes through the holes! Indeed, the integrals along closed loops

$$\oint_{\text{loop}} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} = F(\text{loop}) \quad (24)$$

are gauge-invariant *for single-valued* $\Lambda(\mathbf{x})$, and when $\nabla \times \mathbf{A} \equiv 0$ everywhere outside the holes, then the fluxes (24) depend only on the topologies of the loops in question — which hole(s) they surround and how many times. In math, such integrals are called *cohomologies* of the one-form $\mathbf{A}(\mathbf{x})$.

In classical mechanics, the motion of a charged particle depends on the magnetic field $\mathbf{B}$ in the region of space through which the particle travels, and it does not care about any cohomologies of the vector potential $\mathbf{A}$. But in quantum mechanics, the Aharonov–Bohm effect makes quantum interference sensitive to the cohomologies that the classical mechanics does not see. Specifically, when the space has some holes through which the particle does not get to travel — like the solenoid (and a bit of space around it) in the AB experiment — the interference between alternative paths on different sides of a hole depends on the cohomology of $\mathbf{A}$ for that hole — *i.e.*, the magnetic flux through the hole.

To be precise, the interference between two paths depends on the phase difference (23) only modulo 2π — changing the phase by $2\pi n$ for some integer n would not affect the interference at all. Consequently, the Aharonov–Bohm effect is un-detectable for

$$F = \frac{2\pi\hbar}{q} \times \text{an integer}, \quad (25)$$

or in other words, the AB effect measures only the fractional part of the magnetic flux through the solenoid in units of

$$F_1 = \frac{2\pi\hbar}{|q|} \quad (26)$$

where q is the electric charge of the particles used in the experiment. In particular, the

Chambers's experiment using electron beams was sensitive to the magnetic flux in units of

$$F_1^e = \frac{2\pi\hbar}{e} = 4.135\,667\,697 \times 10^{-15} \text{ Wb} \quad (\text{Weber} = \text{Tesla} \times \text{m}^2). \quad (27)$$

A more practical version of the Aharonov–Bohm experiment is a SQUID (Superconducting Quantum Interferometry Device) magnetometer, which is explained in the next section of these notes. A SQUID uses Cooper pairs instead of single electrons; the electric charge of such a pair is $-2e$, so a SQUID measures the fractional part of the magnetic flux in units of

$$F_1^{\text{Cp}} = \frac{2\pi\hbar}{2e} = 2.067\,833\,848 \times 10^{-15} \text{ Wb}. \quad (28)$$

Note that particles of different charges would measure the fractional part of the magnetic flux F in different units! Thus, were Nature kind enough to provide us with two particle species with an irrational charge ratio q_1/q_2 , then measuring the fractional part of the same flux F in two different units F_1 and F_2 with irrational F_1/F_2 , we would be able to reconstruct the whole flux F and not just its fractional part. However, in reality all the electric charges are integral multiplets of the fundamental charge units e . Consequently, the AB effect using any existing particle species can measure only the fractional part of the magnetic flux in universal units

$$F_1^u = F_1^e = \frac{2\pi\hbar}{e} = 4.135\,667\,697 \times 10^{-15} \text{ Wb}. \quad (29)$$

Superconducting Quantum Interferometry Devices

INTRODUCTION

The superconducting quantum interferometry devices — commonly called the SQUIDs — are extremely sensitive magnetometers, capable of measuring magnetic field variations by as little as 10^{-14} Tesla, or even smaller. The SQUIDs operate on a principle very similar to the Aharonov-Bohm effect, that's why these notes include a section on the SQUIDs.

Warning: This section of the notes was written for an extra lecture to the *graduate students* who have already been through 3 lectures on superconductivity and Josephson junctions. For the undergraduate students who read these notes, here is a very quick summary of the background material necessary to understand the SQUIDs.

In a superconducting metal, electron-phonon interactions create attractive forces between electrons with near-opposite momenta near the Fermi surface. At very low temperatures, such electron pairs form bound states called Cooper pairs, and then these Cooper pairs — which act as slow bosonic particles — form a Bose–Einstein condensate. This condensate acts as a superfluid similar to the liquid helium II, but because the Cooper pairs are electrically charged, this superfluid conducts electric current (called the supercurrent to distinguish it from the current carried by the un-paired electrons) and is sensitive to the magnetic field.

In a Bose–Einstein condensate, all Cooper pairs are in the same quantum state with some wavefunction $\psi_{\text{pair}}(\mathbf{x})$. The Landau–Ginzburg complex scalar field

$$\Psi(\mathbf{x}) = \sqrt{N_{\text{pairs}}} \times \psi_{\text{pair}}(\mathbf{x}) \quad (30)$$

governs the density and the flow velocity of the condensate. In particular, the condensate density is simply $n_s = |\Psi|^2$, while the velocity leads to the supercurrent

$$\mathbf{J}_s = -\frac{2e}{M_{\text{pair}}} n_e \left(\hbar \nabla \text{phase}(\Psi) + 2e\mathbf{A} \right). \quad (31)$$

In a bulk superconductor with uniform n_s , this equation leads to

$$\nabla \times \mathbf{J} = -\frac{4e^2 n_s}{M} \mathbf{B} \quad (32)$$

and hence

$$(\nabla^2 + \kappa^2)\mathbf{B}(\mathbf{x}) = 0 \quad \text{for} \quad \kappa^2 = \frac{4e^2 \mu_0 n_s}{M}. \quad (33)$$

Thanks to this equation, the magnetic field cannot penetrate a superconductor beyond the London’s penetration depth $(1/\kappa) \sim 10$ to 100 nm. From a thick superconductor, the magnetic field is completely expelled beyond a thin surface level; this is the Meissner effect.

In a thick superconducting wire, the current flows only on its surface, so in the middle of the wire

$$\nabla \text{phase}(\Psi) = -\frac{2e}{\hbar} \mathbf{A}. \quad (34)$$

In particular, in the absence of the magnetic field — and in the gauge where $\mathbf{A} = 0$, — the $\text{phase}(\Psi)$ is constant along the wire.

Now consider a Josephson junction: Two pieces of superconducting wire separated by a very thin barrier. The electrons — and hence the Cooper pairs — cannot classically flow through the barrier, but they can jump through it by quantum tunneling. Consequently, a weak supercurrent can flow through the barrier, but it requires different phases $\phi_1 \neq \phi_2$ of the condensate on the two sides of the Josephson junction. Specifically,

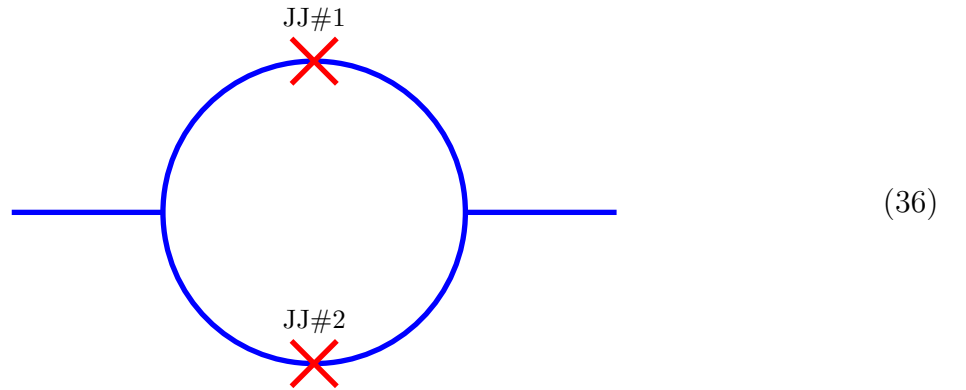
$$I = I_0 \times \sin(\phi_1 - \phi_2) \quad (35)$$

where I_0 depends on the junction's geometry. Experimentally, it's measured as the strongest current that can flow through the Josephson junction without resistance, *i.e.* without causing a voltage drop across the junction.

Josephson junctions are very interesting devices, but explaining them goes beyond the scope of these notes. Even explaining the origin of the current formula (35) is beyond the scope of these notes. Instead, let me explain how a pair of Josephson junctions makes a SQUID magnetometer.

SQUIDS

The SQUIDS come in many shapes, but the simplest version consists of two similar Josephson junctions in a single loop of superconducting wire,



For simplicity, let both Josephson junctions have the same maximal supercurrent I_0 . Then, in complete absence of the magnetic field, the maximal supercurrent which can flow through

the SQUID without generating a voltage is obviously $I_{\max} = 2I_0$. However, in presence of the weak magnetic field, the maximal zero-voltage current through the SQUID is reduced to

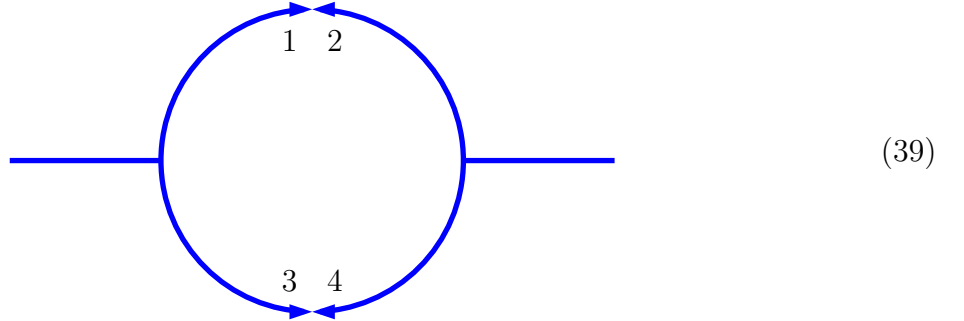
$$I_{\max}(B) = 2I_0 \times \left| \cos \frac{\pi F}{F_0} \right| \quad (37)$$

where F is the flux of the magnetic field through the squid's loop and F_0 is the magnetic flux quantum in the superconductor,

$$F_0 = F_1^{\text{Cp}} = \frac{2\pi\hbar}{2e} = 2.067\,833\,848 \times 10^{-15} \text{ Wb}. \quad (38)$$

Thanks to the very small value of this magnetic flux quantum, the maximal zero-voltage current through the SQUID — which can be easily measured — is extremely sensitive to the tiniest changes of the magnetic field through the SQUID's loop. And when even higher sensitivity is needed, one may combine a SQUID with a magnetic amplifier, or with a cascade arrangement of amplifiers; the engineering of magnetic couplings between SQUIDs and amplifier loops is tricky, but the physics is quite straightforward.

Eq. (37) follows from the Aharonov–Bohm-like interference between the Cooper pairs — or rather between the Cooper pair condensates — flowing through the two Josephson junctions. To see how this interference works, consider the phases $\phi_1, \dots, \phi_4$ of the Cooper pair condensate at 4 points of the SQUID: The two ends (1) and (2) of the top Josephson junction, and the two ends (3) and (4) of the bottom junction:



Eq. (35) for the current through each junction tells us

$$\begin{aligned} I^{\text{top}} &= I_0^{\text{top}} \times \sin(\phi_1 - \phi_2), \\ I^{\text{bot}} &= I_0^{\text{bot}} \times \sin(\phi_3 - \phi_4), \end{aligned} \quad (40)$$

so assuming similar makes (and hence similar I_0) of the two junctions, the net current

through the SQUID is

$$I^{\text{net}} = I^{\text{top}} + I^{\text{bot}} = I_0 \times \left(\sin(\phi_1 - \phi_2) + \sin(\phi_3 - \phi_4) \right). \quad (41)$$

Now consider the left half of the SQUID, specifically the SC wire going from the left end (3) of the bottom junction to the left end (1) of the top junction. Assuming this wire is thick enough, the magnetic field and the supercurrent are expelled by the Meissner effect from the middle of the wire; instead, the supercurrent flows in a thin layer (of thickness ℓ = London's penetration depth) along the wire's surface. Thus, in the middle of the wire

$$\mathbf{J}_s = \frac{-2e\hbar n_s}{M_{\text{eff}}} \left(\nabla_{\text{phase}} + \frac{2e}{\hbar} \mathbf{A} \right) = 0, \quad (42)$$

hence

$$\text{along the wire's axis : } d(\text{phase}) = -\frac{2e}{\hbar} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \quad (43)$$

and therefore

$$\phi_1 - \phi_3 = \frac{-2e}{\hbar} \int_3^1 \mathbf{A} \cdot d\mathbf{x} \quad \text{along the left wire.} \quad (44)$$

The right half of the SQUID — *i.e.*, the SC wire going from the right end (4) of the bottom junction to the right end (2) of the top junction — also has the supercurrent flowing only along the wire's surface, so in the middle of the wire

$$d(\text{phase}) = -\frac{2e}{\hbar} \mathbf{A}(\mathbf{x}) \cdot d\mathbf{x} \quad (45)$$

and therefore

$$\phi_2 - \phi_4 = \frac{-2e}{\hbar} \int_4^2 \mathbf{A} \cdot d\mathbf{x} \quad \text{along the right wire.} \quad (46)$$

Now let's take a difference between eqs. (44) and (46) for the two halves of the SQUID:

$$(\phi_1 - \phi_3) - (\phi_2 - \phi_4) = \frac{-2e}{\hbar} \left[\int_{\substack{\text{left wire} \\ \text{from 3 to 1}}} \mathbf{A} \cdot d\mathbf{x} - \int_{\substack{\text{right wire} \\ \text{from 4 to 2}}} \mathbf{A} \cdot d\mathbf{x} \right]. \quad (47)$$

Geometrically, the SQUID has a much larger size than each of the two Josephson junctions,

so the distances between the two superconductors within each junction — *i.e.*, between (1) and (2), or between (3) and (4), — are much shorter than the distance between the two junctions along either side of the SQUID. So assuming a non-singular vector potential near either junction, we may approximate the integrals in eq. (47) as integrals from the bottom junction to the top junction along 2 different paths,

$$\begin{aligned} \int_3^1 \mathbf{A} \cdot d\mathbf{x} &\approx \int_{\text{bottom JJ}}^{\text{top JJ}} \mathbf{A} \cdot d\mathbf{x} \quad \text{along the left wire,} \\ \int_4^2 \mathbf{A} \cdot d\mathbf{x} &\approx \int_{\text{bottom JJ}}^{\text{top JJ}} \mathbf{A} \cdot d\mathbf{x} \quad \text{along the right wire.} \end{aligned} \tag{48}$$

Hence, the difference between these two integrals is the integral along a closes path which runs up the left wire from the bottom JJ to the top JJ and then down the right wire from the top JJ down to the bottom JJ,

$$\begin{aligned} \int_3^1 \mathbf{A} \cdot d\mathbf{x} - \int_4^2 \mathbf{A} \cdot d\mathbf{x} &= \oint_{\text{whole SQUID}} \mathbf{A} \cdot d\mathbf{x} \\ \langle\langle \text{by the Stokes' theorem} \rangle\rangle & \\ &= F, \quad \text{the magnetic flux through the SQUID.} \end{aligned} \tag{49}$$

In the context of eq. (47), this formula means

$$(\phi_1 - \phi_3) - (\phi_2 - \phi_4) = -\frac{2e}{\hbar} F = -2\pi \frac{F}{F_0}. \tag{50}$$

Next, let's re-organize the LHS here in terms of differences $\phi_1 - \phi_2$ and $\phi_3 - \phi_4$ instead of $\phi_1 - \phi_3$ and $\phi_2 - \phi_4$, thus

$$(\phi_1 - \phi_2) - (\phi_3 - \phi_4) = (\phi_1 - \phi_3) - (\phi_2 - \phi_4) = \frac{-2\pi F}{F_0}. \tag{51}$$

Also, let θ denote the average between $\phi_1 - \phi_2$ and $\phi_3 - \phi_4$,

$$\theta = \frac{1}{2}(\phi_1 - \phi_2) + \frac{1}{2}(\phi_3 - \phi_4); \tag{52}$$

then in terms of this θ and the magnetic flux F through the SQUID,

$$(\phi_1 - \phi_2) = \theta - \frac{\pi F}{F_0}, \quad (\phi_3 - \phi_4) = \theta + \frac{\pi F}{F_0}. \quad (53)$$

Finally, let's plug these phase difference across each Josephson junctions into eq. (41) for the net current through the SQUID:

$$\begin{aligned} \frac{I^{\text{net}}}{I_0} &= \sin(\phi_1 - \phi_2) + \sin(\phi_3 - \phi_4) \\ &= \sin\left(\theta - \frac{\pi F}{F_0}\right) + \sin\left(\theta + \frac{\pi F}{F_0}\right) \\ &= 2 \sin \theta \times \cos \frac{\pi F}{F_0}, \end{aligned} \quad (54)$$

or equivalently

$$I^{\text{net}} = \left(2I_0 \cos \frac{\pi F}{F_0}\right) \times \sin \theta. \quad (55)$$

Experimentally, we control the magnetic flux F through the SQUID but we have no direct control over the averaged phase difference θ . Instead, we control the net current through the SQUID while the θ adjusts itself to whatever it takes to carry the desired current. However, for any possible value of θ , the $\sin \theta$ ranges between -1 and $+1$ and cannot exceed these limits. Consequently, the net supercurrent through the SQUID varies in the range

$$-2I_0 \left| \cos \frac{\pi F}{F_0} \right| < I^{\text{net}} < +2I_0 \left| \cos \frac{\pi F}{F_0} \right| \quad (56)$$

but cannot get any stronger than this in either direction. In other words, *the maximal supercurrent through the SQUID* is

$$I_{\text{max}} = 2I_0 \times \left| \cos \frac{\pi F}{F_0} \right|. \quad (37)$$

Quod erat demonstrandum.

PS: As written, eq. (37) is valid for the relatively weak magnetic fields — strong enough to affect the interference between the two Josephson junctions, but not so strong to affect the Josephson junctions themselves. In stronger fields, eq. (37) becomes

$$I_{\max} = 2I_0(B) \times \left| \cos \frac{\pi F}{F_0} \right| \quad (57)$$

where $I_0(B)$ decreases with the magnetic field. The specific analysis of the $I_0(B)$ is quite beyond the scope of this notes, so let me simply say two things: (1) the details depend on the precise geometry of the Josephson junctions, and (2)

$$I_0(B) \approx I_0(0) \quad \text{as long as} \quad B \times (\text{Junction's area}) \ll F_0. \quad (58)$$

Note that the junction's area is much smaller than the area of the SQUID's loop, so this condition can hold while at the same time

$$F[\text{squid}] \gg F_0. \quad (59)$$