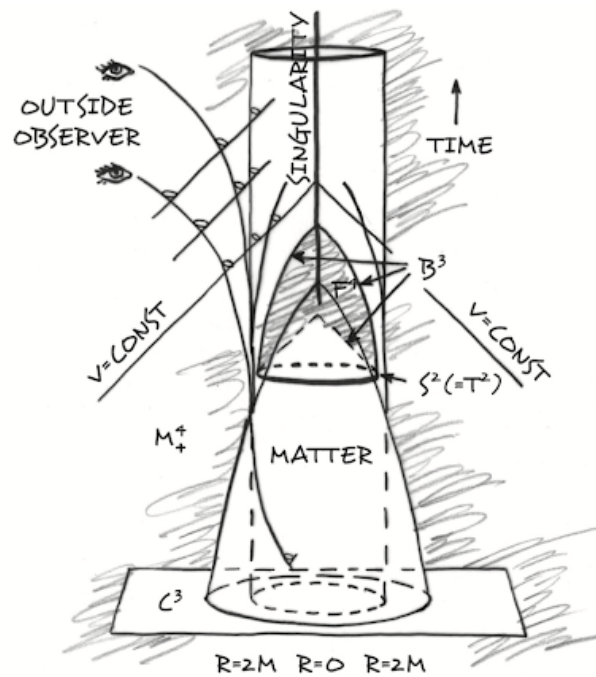


FROM EINSTEIN TO SCHRÖDINGER

MATHEMATICAL FOUNDATIONS OF THEORETICAL PHYSICS, VOLUME I

WRITTEN BY
TIANYI WANG



PUBLISHED
IN THE WILD

Preface

There are two perspectives one can take to appreciate this intricate connection between mathematics and physics: The first perspective is what I think most physicists take. That is, to be amazed by the fact that using only simple and elegant mathematical language can one model how our universe works. But I prefer to think of this the other way around: Why does the universe work in a way that can be described by profound and elegant mathematics? I believe there must be some deeper reasons for this question, since the interplay of mathematics and physics is too beautiful to be just a coincidence.

This book starts with an introduction to classical physics: Newtonian mechanics, classical field theory, special and general theory of relativity. Then a considerable portion of the book is dedicated to quantum mechanics. I should be absolutely clear that this book is by no means comprehensive: We only cover the basic knowledge of theoretical physics needed by mathematicians interested to work in mathematical physics. So certain topics (despite their importance) are ignored. For example, there is no discussion of black holes in the general relativity section, nor is there a discussion on perturbation theory in the quantum mechanics portion.

Of course, as suggested by the title of this book, a significant portion of this book will be dedicated to the mathematics that arise from these theories. There are several major topics that I choose to cover: Riemannian geometry is carefully developed in the general relativity portion, since this is the foundation of the theory. For quantum mechanics, people have different tastes: Many people choose to cover a lot of analysis and operator theory, such as [5]. But my taste is more geometric, so I made a decision to include Lie groups and representation theory before the analysis portion as this is such a beautiful theory with profound applications to quantum mechanics.

A few words about the prerequisites to read this book: I assume the reader is comfortable with basic analysis, linear algebra, and group theory. A basic understanding of topology and smooth manifolds are recommended, but not required. A good reference for such material will be [1], but this is too much. For practical purposes, the first part of my notes [16] would suffice. Although it could be very helpful, no previous knowledge of physics is assumed.

The idea to write this book came from a year long course on mathematical physics that the author took in his first year graduate study taught by Prof. Ron Donagi at the University of Pennsylvania. But there are many more people to thank: The chapter on general relativity was based on the course taught by Prof. Gary Horowitz. The chapter on quantum mechanics was mostly based on a reading course with Prof. David Morrison, but I also incorporated some notes from my quantum mechanics course with Prof. Gary

Horowitz. Thank you all for being such amazing teachers and for constantly inspiring me in my study. This book is dedicated to all of you.

Tianyi Wang¹

Sep 2024, Philadelphia.

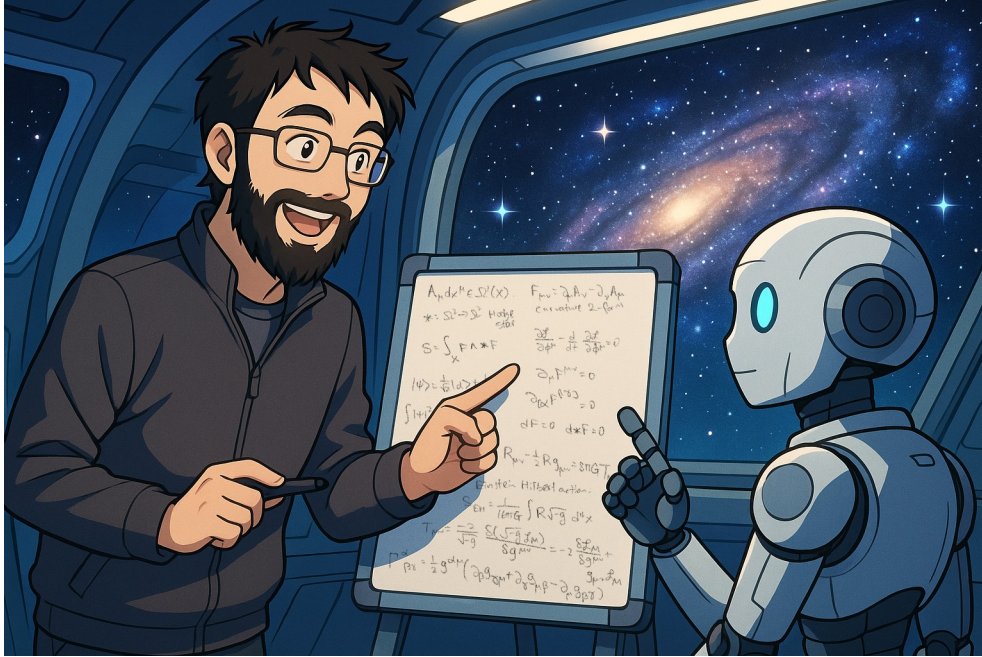


Figure 1: A drawing of me on a spaceship, explaining Yang-Mills theory and Einstein's general relativity to a robot. It is a good approximation of me generated by ChatGPT as I used my actual photo as input.

¹tywang9@sas.upenn.edu

Contents

Preface	1
I Relativity	7
1 A Quick Review of Classical Mechanics	8
1.1 Hamiltonian Formalism	8
1.2 Lagrangian Formalism	11
2 Special Relativity	15
2.1 Newtonian Gravity	15
2.2 Frame of Reference	16
2.3 Newton's First Law and Inertial Frame	20
2.4 Galilean Transformation	20
2.5 Special Relativity	24
2.6 Vector, Dual Vectors, and Tensors	35
2.7 Relativistic Dynamics	42
2.8 Energy Momentum Tensor	43
3 Geometry of Spacetime	48
3.1 Principle of Equivalence	49
3.2 Smooth Manifolds	54
3.3 Tangent Vectors	56
3.4 Tensors	59
3.5 Differential Forms and Integration on Manifolds	61
3.6 Covariant Derivatives	65
3.7 Parallel Transport and Geodesics	68
3.8 The Riemann Curvature Tensor	71
3.9 Geodesic Deviation	77
4 Einstein's Field Equation	79
4.1 Physics in Curved Spacetime	79
4.2 Einstein's Equation	82
4.3 *Einstein-Hilbert Action	84
4.4 Linearized Gravity	84

II	Quantum Mechanics	89
5	Quantum Theory and Representations	90
5.1	Introduction and Overview	90
5.1.1	Axioms of Quantum Mechanics	90
5.1.2	Group Action and Representation	90
5.1.3	Representations and Quantum Mechanics	91
5.2	$U(1)$ and its Representations	91
5.2.1	Irreducible Representations	91
5.2.2	$U(1)$ and Its Representation	92
5.2.3	The Charge Operator	92
5.2.4	Conservation of Charge and $U(1)$ Symmetry	93
5.3	Lie Groups and Lie Algebras	93
5.3.1	Lie Groups	93
5.3.2	Lie Algebras	94
5.3.3	Lie Group and Lie Algebra Representations	95
5.4	Irreducible Representations of $SU(2)$	96
5.4.1	Structure of $\mathfrak{su}(2)$	97
5.4.2	The Highest Weight Theorem	98
5.4.3	Irreducible Representations of $SU(2)$	100
5.4.4	Unitarity of the Representation	102
5.5	Irreducible Representations of $SO(3)$	102
5.5.1	Properties of $SO(3)$ and $\mathfrak{so}(3)$	102
5.5.2	Irreducible Representation of $\mathfrak{so}(3)$	104
5.5.3	Irreducible Representations of $SO(3)$	106
5.5.4	Spherical Harmonics	107
5.5.5	The Casimir operator	110
5.6	Tensor Products	111
5.6.1	Tensor Products in Quantum Mechanics: Pauli Exclusion and Entanglement	111
5.6.2	Tensor Products of Representations	112
5.6.3	Clebsch-Gordan Decomposition	113
5.6.4	Complexification of Vector Space	115
5.7	Representation of the Group $(\mathbf{R}, +)$	116
5.7.1	Hamiltonian and Momentum Operator	117
5.8	Free Particles	119
5.8.1	Schwartz Space and Tempered Distributions	119
5.8.2	Fourier Transform	122
5.8.3	Dirac Notation, Position and Momentum Space	123
5.9	Representations of Euclidean Group	126
5.9.1	Solution to the Free Particle SDE in Momentum Space	126
5.9.2	Semidirect Products and The Euclidean Group	126
5.9.3	Irreducible Representations of $E(2)$	128
6	The Schrödinger Equations	132
6.1	Schrödinger Equation in Position Space	132
6.2	Infinite Square Well	134
6.3	The Harmonic Oscillator	136
6.4	The Uncertainty Principle	140
6.5	Quantum Mechanics in Three Dimensions	142

6.5.1	The Angular Equation	144
6.5.2	The Radial Equation	146
6.6	Hydrogen Atom	146
7	Operator Theory in Quantum Mechanics	153
7.1	Basic Theory of Bounded Operators	153
7.2	Spectral Theorem for Bounded Operators	162
	Index	178
	Bibliography	180

Part I

Relativity

Chapter 1

A Quick Review of Classical Mechanics

1.1 Hamiltonian Formalism

In classical mechanics, the space that we use to describe physical systems that we can observe is called the *configuration space*, which we denote by B . Say one has a single particle moving in 3d space, then its position is completely characterized by 3 coordinates, $(q_1(t), q_2(t), q_3(t))$. The configuration space, i.e., the set of all possible positions that the particle can go to, will then be $\mathbf{R}^3$. If we have n particles, then their trajectory can be described by specifying $3n$ coordinates, hence the configuration space will be $\mathbf{R}^{3n}$.

The more interesting situation arises when the systems have some constraints: e.g. if some particles are attached to each other by rigid rods that has fixed length. But since these constraints are usually given by algebraic equations, we have that in general, the configuration space B will be some submanifold of $\mathbf{R}^{3n}$. For now, let us assume that the configuration space is a manifold B , with local coordinates $q_1, \dots, q_n$.

The *phase space*, which describes all possible position and velocity of the collection of particles, is defined to be $M = T^*B$, the cotangent bundle of the configuration space B . One might be confused why we did not use the tangent bundle TB to define phase space, since we are talking about velocities. As it turns out, on T^*B , there is a natural symplectic structure, which leads us to the following

Definition 1.1. A *symplectic manifold* is a pair (M, ω) , where M is a smooth manifold, and ω is a closed, nondegenerate 2-form on M , where:

- closed means $d\omega = 0$.
- non-degenerates means that the map $\omega : T_p M \rightarrow T_p^* M$ given by $X_p \mapsto \omega(X_p, -)$ is an isomorphism at every $p \in M$.

Proposition 1.2. Given any manifold B , its cotangent space T^*B is always a symplectic manifold with a natural symplectic form.

Proof. We construct the natural symplectic form ω in local coordinates first: Let $q_1, \dots, q_n$ be local coordinates on B , and let $p_1, \dots, p_n$ be fiber coordinates

on T^*B . Then consider the 1-form

$$\alpha = \sum_{i=1}^n p_i dq_i.$$

We then define a 2-form by

$$\omega = d\alpha = \sum_{i=1}^n dp_i \wedge dq_i.$$

Then one can check that α, ω can be globally defined, because we can define them locally as above, and we may check that in the intersection of two charts, their values agree. We claim that ω is symplectic. Closedness is easy since $d\omega = d^2\alpha = 0$. Nondegeneracy is also easy to check. $\square$

In classical mechanics, the **observables**, i.e., quantities we may measure through experiments, are smooth functions on the phase space, i.e., $C^\infty(T^*B)$. Think of them as functions of positions and momentums. There is an observable $H \in C^\infty(T^*B)$ that is of particular importance, called the **Hamiltonian**, which corresponds to the energy of the system. The Hamiltonian is important because it governs the evolution of all other physical observables. First, there is an associated vector field X_H , called **Hamiltonian flow**, defined as follows:

Definition 1.3. Let $H \in C^\infty(T^*B)$ be the Hamiltonian of the system. The **Hamiltonian flow** X_H is defined as follows: Let $\omega_p^{-1} : T_p^*M \rightarrow T_pM$ be the inverse isomorphism defined by the symplectic form ω , and let dH denote the differential of H . Then

$$(X_H)_p = \omega_p^{-1}(dH_p).$$

Proposition 1.4. In local coordinates $p_1, \dots, p_n, q_1, \dots, q_n$, the Hamiltonian flow is given by

$$X_H = \sum_{i=1}^n \left(\frac{\partial H}{\partial p_i} \frac{\partial}{\partial q_i} - \frac{\partial H}{\partial q_i} \frac{\partial}{\partial p_i} \right).$$

The basic axiom of Hamiltonian mechanics is that the trajectories of a Hamiltonian system (M, ω, H) are the integral curves of X_H . That is, if

$$(q_1(t), \dots, q_n(t), p_1(t), \dots, p_n(t))$$

denote the said trajectory in local coordinates, then they are governed by the Hamilton equation

$$\dot{p}_i = -\frac{\partial H}{\partial q_i} \quad \dot{q}_i = \frac{\partial H}{\partial p_i}.$$

Lemma 1.5. The Hamiltonian flow preserves ω . In other words, let g_t denote the flow of X_H , for some real t . Then $g_t^*\omega = \omega$ for all t .

Proof. Note that it suffices to prove $g_t^*\omega = \omega$ for small t , since for large t we can compose flows with small t many times. Hence it suffices to prove that $\frac{d}{dt} \big|_{t=0} (g_t^*\omega) = 0$. But note that $\frac{d}{dt} \big|_{t=0} (g_t^*\omega) = L_{X_H}\omega$ is the Lie derivative of ω with respect to X_H .

Now, using the fact that $L_X(df) = dX(f)$ and $L_{X_H}(dp \wedge dq) = L_X(dp) \wedge dq - dp \wedge L_X(dq)$, and using the local coordinate expression of X_H , we have that

$$\begin{aligned} L_{X_H}\omega &= L_{X_H}(dp \wedge dq) \\ &= L_{X_H}(dp) \wedge dq - dp \wedge L_{X_H}(dq) \\ &= -d(\partial H/\partial q) \wedge dq + dp \wedge d(\partial H/\partial p) \\ &= -d^2H = 0. \end{aligned}$$

This proves the claim. $\square$

Definition 1.6. The form

$$\frac{1}{n!}\omega^n \in \Omega^{2n}(T^*B)$$

is called the **Liouville volume form** of (M, ω) .

Theorem 1.7 (Liouville's Theorem). *Hamiltonian flow preserves the Liouville volume form.*

Proof. The Hamiltonian flow preserves ω , and pullback commutes with wedge product. $\square$

Definition 1.8. Let $S \subset (M, \omega)$ be a submanifold. We say that S is **isotropic** if $\omega|_S = 0$. A maximal isotropic submanifold in M is called a **Lagrangian submanifold**.

Corollary 1.9. *The Hamiltonian flow sends Lagrangian submanifolds to Lagrangian submanifolds.*

Definition 1.10. A smooth manifold M , together with a map $\{\cdot, \cdot\} : C^\infty(M) \times C^\infty(M) \rightarrow C^\infty(M)$ is called a **Poisson manifold**, if $\{\cdot, \cdot\}$

- gives a Lie algebra structure to $C^\infty(M)$;
- is a derivation: $\{fg, h\} = f\{g, h\} + g\{f, h\}$;
- satisfies the Libniz condition: $\{f, \cdot\} : C^\infty(M) \rightarrow C^\infty(M)$ is a vector field on M .

Proposition 1.11. Any symplectic manifold is Poisson.

Proof. Given any $f \in C^\infty(M)$, we again define $(X_f)_p = \omega_p^{-1}(df_p)$. Then we may define $\{f, g\} = \omega(X_f, X_g)$. $\square$

Definition 1.12. Let $f : (M, \omega) \rightarrow (M', \omega')$ be a map between symplectic manifolds. f is said to be **symplectic** if $f^*\omega' = \omega$. If f is symplectic and a diffeomorphism, then f is said to be a **symplectomorphism**.

Let G be the group of symplectomorphisms on M . Then there is a natural action $\rho : G \times M \rightarrow M$ given by $(f, p) \mapsto f(p)$. For each $p \in M, \xi \in \mathfrak{g}$, the differential of the map $\rho(-, p) : G \rightarrow M$ evaluated at $\mathfrak{g}$ gives an element $\rho(\xi) \in T_pM$. Do this for all p we get a vector field on M . Thus we get a 1-form $\omega(\rho(\xi)) \in \Omega^1(M)$.

Definition 1.13. A **momentum map** is a map $\mu : M \rightarrow \mathfrak{g}^*$ such that $d\langle \mu, \xi \rangle = \omega(\rho(\xi))$. A G -action is called **Hamiltonian** if a momentum map exists.

1.2 Lagrangian Formalism

Let M be the configuration space. This time we consider the tangent bundle TM . Let $q^1, \dots, q^n$ be coordinates on M , the base space. A path $\gamma : [t_0, t_1] \rightarrow M$ has a natural lift to $\tilde{\gamma} : [t_0, t_1] \rightarrow TM$ via $t \mapsto (\gamma(t), \gamma'(t))$. In this case the lifted path has coordinate representation

$$\tilde{\gamma}(t) = (q^1(t), \dots, q^n(t), \dot{q}^1(t), \dots, \dot{q}^n(t)).$$

The physics is specified by a **Lagrangian** $L : TM \times \mathbf{R} \rightarrow \mathbf{R}$, or $L(q, \dot{q}, t)$.

Consider the space

$$\mathcal{P}(M) = \{\gamma : [t_0, t_1] \rightarrow M : \gamma(t_0) = q_0, \gamma(t_1) = q_1\}.$$

Generally, $\mathcal{P}(M)$ is a Fréchet manifold. For $\gamma \in \mathcal{P}(M)$, it is not hard to see that $T_\gamma \mathcal{P}(M)$ is the set of vector fields on M along γ that vanishes at the endpoints q_0, q_1 . We formalize this idea below:

Definition 1.14. A **variation** of γ is a map $\Gamma : [t_0, t_1] \times [-\epsilon, \epsilon] \rightarrow M$ such that $\Gamma(t_0, a) = q_0, \Gamma(t_1, a) = q_1$, for all $a \in [-\epsilon, \epsilon]$, and $\Gamma(t, 0) = \gamma(t)$.

Definition 1.15. An **action** is a function $S : \mathcal{P}(M) \rightarrow \mathbf{R}$, given by

$$S(\gamma) = \int_{t_0}^{t_1} L(\gamma(t), \gamma'(t), t) dt.$$

The fundamental axiom of Lagrangian mechanics is the **principle of least action**, which says that the physical path is a minimum of S . In other words, the physical path γ satisfies the following minimality condition: Let γ_ϵ be any variation of γ , then

$$\left. \frac{d}{d\epsilon} \right|_{\epsilon=0} S(\gamma_\epsilon) = 0.$$

Theorem 1.16. *The equation of motion is given by the following **Euler-Lagrange equation**:*

$$\frac{\partial L}{\partial q} - \frac{d}{dt} \frac{\partial L}{\partial \dot{q}} = 0.$$

Proof. By the condition $\left. \frac{d}{d\epsilon} \right|_{\epsilon=0} S(\gamma_\epsilon) = 0$, we have that

$$\begin{aligned} 0 &= \left. \frac{d}{d\epsilon} \right|_{\epsilon=0} \int_{t_0}^{t_1} L(q(t, \epsilon), \dot{q}(t, \epsilon), t) dt \\ &= \sum_{i=1}^n \int_{t_0}^{t_1} \left(\frac{\partial L}{\partial q^i} \frac{\partial q^i}{\partial \epsilon} + \frac{\partial L}{\partial \dot{q}^i} \frac{\partial \dot{q}^i}{\partial \epsilon} \right) dt \\ &= \sum_{i=1}^n \int_{t_0}^{t_1} \left(\frac{\partial L}{\partial q^i} - \frac{d}{dt} \frac{\partial L}{\partial \dot{q}^i} \right) \frac{\partial q^i}{\partial \epsilon} dt + \sum_{i=1}^n \underbrace{\frac{\partial L}{\partial \dot{q}^i} \frac{\partial q^i}{\partial \epsilon}}_{=0} \Big|_{t_0}^{t_1} \end{aligned}$$

where in the last equation, we performed integrated by part to the second term of the integrand in the second equation, and the boundary term is zero because variation fixes endpoints. Since the above calculation is true for any variation, we get the Euler-Lagrange equation as desired. $\square$

Next, we will do a lot of examples, applying Euler-Lagrange (EL) equation to find the equation of motion (EOM).

Example 1.17. Let $M = \mathbf{R}^3$. Consider the situation where we just have a free particle. The physics should be invariant under rotation and translation, hence the symmetry group is $G = O(3) \ltimes \mathbf{R}^3$, i.e., the *Gallilean group*. Thus the path should be invariant under the transformation

$$r \mapsto gr + r_0$$

for $g \in O(3), r_0 \in \mathbf{R}^3$. Thus the Lagrangian should be invariant under the symmetry group, i.e., should be of the form

$$L = \frac{1}{2}m\dot{r}^2.$$

Applying EL equation gives the EOM:

$$\ddot{r} = 0.$$

Example 1.18. Consider N particle interacting now, $M = \mathbf{R}^{3N}$. The Lagrangian is given by

$$L = \sum_{i=1}^N \frac{1}{2}m_i\dot{r}_i^2 - V(r) =: T - V(r).$$

The function $V(r) = V(r_1, \dots, r_N)$ is called the *potential* of the system, and the function T is called the *kinetic energy* of the system. The EOM is then

$$m_i\ddot{r}_i = -\frac{\partial V}{\partial r_i}.$$

The term $F_i := -\partial V/\partial r_i$ is defined to be the *force* exerted on the i th particle.

Example 1.19. Usually we have pairwise interactions, i.e.,

$$V(r) = \sum_{1 \leq a < b \leq N} V_{ab}(r_a - r_b).$$

Then

$$L = \frac{1}{2} \sum_{i=1}^N m_i\dot{r}_i^2 - \sum_{1 \leq a < b \leq N} V_{ab}(r_a - r_b).$$

If now V is homogeneous of degree d , i.e., $V(\lambda r) = \lambda^d V(r)$, then

$$\begin{aligned} d\bar{T} &= \frac{1}{\tau} \int_0^\tau \sum_i m_i\dot{r}_i^2 dt \\ &= -\frac{1}{\tau} \int_0^\tau \sum_i m_i\ddot{r}_i^2 \cdot r_i dt \\ &= -\frac{1}{\tau} \int_0^\tau \sum_i F_i \cdot r_i dt \\ &= \frac{1}{\tau} \int_0^\tau \sum_i r_i \cdot \frac{\partial V}{\partial r_i} dt \\ &= d\bar{V}. \end{aligned}$$

Hence $\bar{T} + \bar{V}$ is conserved. Note that in the second equation we used integration by parts.

Example 1.20 (Charged Particle in EM Field). The electromagnetic field consists of a scalar potential $\phi(r)$, and a vector potential $A(r) = (A_1, A_2, A_3)$, assuming they are all time independent for the moment. A particle is characterized by its mass m and its charge e . The Lagrangian is given by

$$L = \frac{1}{2}m\dot{r}^2 + e\left(\frac{\dot{r} \cdot A}{c} - \phi\right).$$

Then by the EL equation we get that the EOM is given by

$$m\ddot{r} = e\left(\frac{\dot{r}}{c} \times (\nabla \times A) - \frac{\partial \phi}{\partial r}\right) =: e\left(\frac{\dot{r}}{c} \times B + E\right)$$

where

$$E = -\frac{\partial \phi}{\partial r}$$

is the *electric field*, and it is to be understood component wise, i.e., E is a vector field with components $E_i = -\partial\phi/\partial r_i$.

$$B = \nabla \times A$$

is the *magnetic field*. There is a 4-dimensional interpretation; Let $M = \mathbf{R}^4$. Consider the 1-form

$$A_1 dx^1 + A_2 dx^2 + A_3 dx^3 + \phi dt$$

on M . Consider the connection on the trivial bundle $M \times \mathbf{R}$ over M . The curvature 2-form is going to be

$$F = \sum_{i=1,2,3} E_i dx^i dt + \sum_{abc=\{123,231,312\}} B_a dx^b dx^c.$$

Example 1.21 (Simple Harmonic Oscillator). Let $M = \mathbf{R}$ and $V = \omega r^2$. The EOM of this system is

$$\ddot{r} + \omega^2 r = 0.$$

The solution is given by

$$r(t) = a \cos \omega t + b \sin \omega t.$$

For a general potential $V(r)$, Taylor expanding at minimum will give a quadratic term approximation up to a constant shift, hence we can approximate any potential by harmonic oscillators at the minimum.

Example 1.22. Consider a free particle traveling on a Riemannian manifold (M, g) . The Lagrangian is going to be

$$L = \frac{1}{2} \langle \dot{\gamma}, \dot{\gamma} \rangle.$$

The equation of motion is going to be the geodesic equation in Riemannian geometry: If $\gamma(t) = x^a(t) \frac{\partial}{\partial x^a}$ in local coordinates, then the EOM is given by

$$\frac{d^2 x^a}{dt^2} + \Gamma_{bc}^a \frac{dx^b}{dt} \frac{dx^c}{dt} = 0$$

where Γ_{bc}^a are the associated christoffel symbols to the metric g_{ab} .

Example 1.23 (Rigid Body). For a rigid body, we have $M = SO(3)$, and $L = \langle x, x \rangle$, where the norm is any left-invariant metric on $SO(3)$. Let $\Omega := (L_{g^{-1}})_* \dot{g} \in \mathfrak{so}(3)$, where $g \in SO(3)$. Then

$$L = \frac{1}{2} \langle A\Omega, \Omega \rangle$$

for some matrix A , which physicists call inertial tensor. Diagonalizing A we get a bunch of eigenvectors (principal axis), and a bunch of eigenvalues (frequencies).

Chapter 2

Special Relativity

In this chapter, we review some of the most important prerequisites for general relativity. First, we will review the Newtonian equation for gravitation; Then, we discuss what went wrong with Newtonian theory, and study its generalization in flat spacetime – special relativity. Finally, we will introduce some classical field theory, including Maxwell’s theory for electromagnetism.

2.1 Newtonian Gravity

Our first step toward the investigation of gravity is to derive the Poisson’s equation, the field equation that governs Newton’s theory of gravity.

Imagine that we have some mass distribution throughout some space. And at every point $\mathbf{r}$ (with some arbitrary origin chosen), the quantity $\rho(\mathbf{r})$ gives the mass density at that point. We let $\mathbf{g}(\mathbf{r})$ to be the gravitational field at the point $\mathbf{r}$. Then the effect of gravitational field $\mathbf{g}(\mathbf{r})$ at point $\mathbf{r}$ can be calculated by adding all the contributions with source in an infinitesimal region dV . Since the region dV can be approximated by point mass, we can use Newton’s gravity formula to calculate the infinitesimal potential contribution at point $\mathbf{s}$:

$$d\mathbf{g} = -G \frac{\mathbf{r} - \mathbf{s}}{|\mathbf{r} - \mathbf{s}|^3} \rho(\mathbf{s}) dV$$

where G is Newton’s gravitational constant. Integrating all such contributions we find

$$\mathbf{g}(\mathbf{r}) = -G \int \rho(\mathbf{s}) \frac{\mathbf{r} - \mathbf{s}}{|\mathbf{r} - \mathbf{s}|^3} dV.$$

We use the identity

$$\nabla \cdot \left(\frac{\mathbf{r}}{|\mathbf{r}|^3} \right) = 4\pi\delta(\mathbf{r})$$

to derive that

$$\nabla \cdot \mathbf{g}(\mathbf{r}) = -4\pi G \int \rho(\mathbf{s}) \delta(\mathbf{r} - \mathbf{s}) dV = -4\pi G \rho(\mathbf{r}).$$

Since the gravitational field has zero curl (equivalently, gravity is a conservative force), it can be written as the gradient of a scalar potential, called the gravitational potential:

$$\mathbf{g} = -\nabla\phi.$$

Substitute this into the equation above and we get the *Poisson's equation*

$$\boxed{\nabla^2 \phi = 4\pi G \rho.} \quad (2.1)$$

Therefore, given any mass density distribution $\rho(\mathbf{r})$ of matters, we can solve (2.1) and find the gravitational potential $\phi(\mathbf{r})$ of this matter distribution. Then any particle in the space will experience such a potential due to the presence of the given matter.

Einstein's field equation is a generalization of Poisson's equation:

$$R_{\mu\nu} - \frac{1}{2} R g_{\mu\nu} = 8\pi G T_{\mu\nu}.$$

We will discuss later what the terms above actually are and why this is a generalization of Poisson's equation.

2.2 Frame of Reference

Physical processes either directly or indirectly involve the dynamics of particles or fields moving through space and time. Almost all important physical laws involve position and time in some ways. For example, Newton's second law and the wave equation obtained from Maxwell's equation:

$$\mathbf{F} = m\mathbf{a} \quad \nabla^2 \mathbf{E} - \frac{1}{c^2} \frac{\partial^2 \mathbf{E}}{\partial t^2} = 0.$$

Implicit in these statements of these fundamental physical laws is the notion that we have at hand some way of measuring or specifying or labelling each point in space and and different instants in time. So, in order to describe in a quantitative fashion the multitude of physical processes that occur in the natural world, one of the important requirements is that we be able to specify where and when events take place in space and time. By "event" we could be referring to something that occurs at an instant in time at a point in space, or, more colloquially, over a localized interval in time, and in a localized region in space. Whatever the circumstance, we specify the where and when of an event by measuring its position relative to some conveniently chosen origin, and using a clock synchronized in some agreed fashion with all other clocks, to specify the time. This combination of a means of measuring the position of events, and the time at which they occur, constitutes what is referred to as a frame of reference.

A frame of reference can be constructed in essentially any way, provided it meets the requirements that it labels in a unique fashion the position and the time of the occurrence of any event that might occur. A convenient way of imagining how this might be done is to suppose that all of space is filled with a three dimensional lattice or scaffolding something like a fishing net, perhaps. The net need not be rigid, and the spacing between adjacent points where the coordinate lines cross need not be the same everywhere.

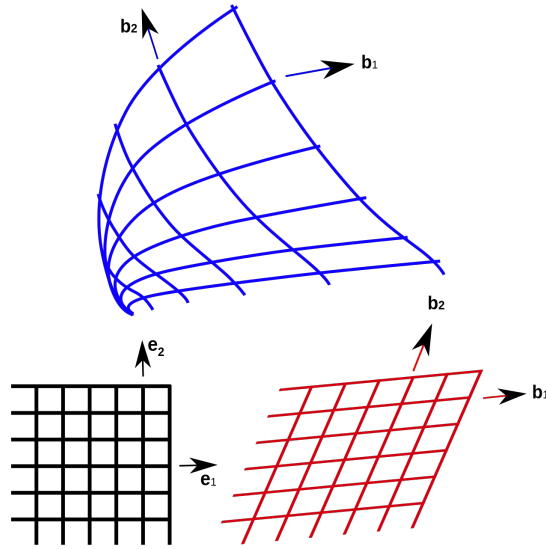


Figure 2.1: Three examples of frame of reference. Each intersection point is labeled by a triplet of three real numbers, the usual coordinates of a point in three dimensional space, with one arbitrary point on the net chosen as the origin.

To complete the picture, we also imagine that attached to the net at each intersection point is a clock. In the same way that we do not necessarily require the net to be rigid, or uniform, we do not necessarily require these clocks to run at the same rate (which is the temporal equivalent of the intersection points on the net not being equally spaced). We do not even require the clocks to be synchronized in any way. At this stage, the role of these clocks is simply to provide us with a specification of the time at which an event occurs in the neighbourhood of the site to which the clock is attached. Thus, if an event occur in space, such as a small supernova flaring up, burn marks will be left on the netting, and the clock closest to the supernova will grind to a halt, so that it reads the time at which the event occurred, while the coordinates of the crossing point closest to where the burn marks appear will give the position of the event.

We can set up any number of such reference frames, each with its own coordinate network and set of coordinate clocks. We have (almost) total freedom to set up a reference frame any way we like, including different reference frames being in motion, or even accelerating, with respect to one another, and not necessarily in the same way everywhere. But whichever reference frame we use, we can then conduct experiments whose outcomes are expressed in terms of the associated set of coordinates, and express the various laws of physics in terms of the coordinate systems used. So where does the principle of relativity come into the picture here? What this principle is saying, in its most general form, is that any physical process taking place in space and time ought to proceed in a fashion that takes no account of the reference frame that we use to describe it.

In other words, it ought to be possible to write down the laws of physics in terms of quantities that make no mention whatsoever of

any particular reference frame.

As we have just seen, a reference frame can be defined in a multitude of ways, but quite obviously it would be preferable to use the simplest possible, which brings to mind the familiar Cartesian set of coordinate axes. Thus, suppose we set up a lattice work of rods as illustrated in Fig. 2.2 in which the rods extend indefinitely in all directions.

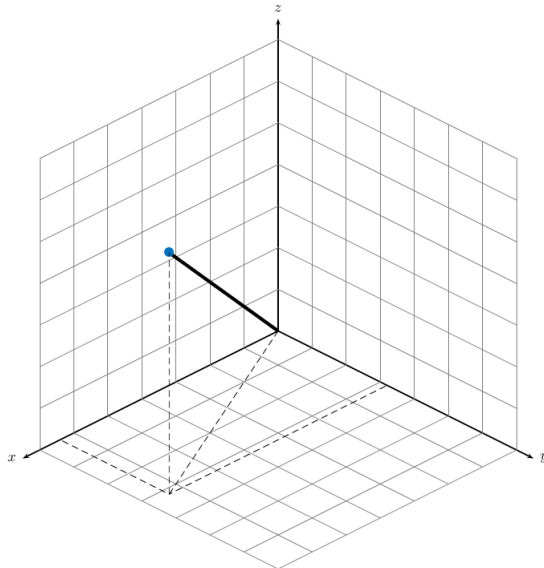


Figure 2.2

The question then arises: can we in fact do this for all of space? From the time of Euclid, and perhaps even earlier, until the 19th century, it was taken for granted that this would be possible, with the rods remaining parallel in each direction out to infinity, even though attempts to prove this from the basic axioms of Euclidean geometry never succeeded. Eventually it was realized by the mathematicians Gauss, Riemann and Lobachevsky that this idea about parallel lines never meeting, while intuitively plausible, was not in fact necessarily true. It was perfectly possible to construct geometries wherein parallel lines could either meet, or diverge, when sufficiently far extended without resulting in any mathematical inconsistencies. Such behaviour is indicative of space being intrinsically curved, and there is no apriori reason why space ought to be flat as Euclid assumed, i.e. space could possess some kind of curvature. Riemann, a student of Gauss, surmised that the curvature of space was somehow related to the force of gravity, but he was missing one important ingredient it is the curvature of space and time together that gives rise to gravity, as Einstein was able to show. Gravitational forces are thus rather peculiar forces in comparison to the other forces of nature such as electromagnetic forces or the nuclear forces in that they are not due to the action of some external influence exerting its effect within pre-existing space and time, but rather is associated with the intrinsic properties of spacetime itself.

Thus, we cannot presume that the arrangement of rods as indicated in Fig. 2.2 can be extended indefinitely throughout all of space. However, over a

sufficiently localized region of space it can be shown to be the case that such an arrangement is possible, i.e. it is possible to construct frames of reference of a particularly simple form: a lattice of mutually perpendicular rigid rods. In the presence of gravity, it is not possible to extend such a lattice throughout all space, essentially because space is curved. But for gravity free space, these rods can be presumed to extend to infinity in all directions. In addition, it is possible to associate with this lattice a collection of synchronized identical clocks positioned at the intersection of these rods. We will take this simple lattice structure as our starting point for discussing special relativity. That this can be done is a signature of what is known as flat space-time.

First of all we can specify the positions of the particle in space by determining its coordinates relative to a set of mutually perpendicular axes x, y, z . In practice this could be done by choosing our origin of coordinates to be some convenient point and imagining that rigid rulers which we can also imagine to be as long as necessary are laid out from this origin along these three mutually perpendicular directions. The position of the particle can then be read off from these rulers, thereby giving the three position coordinates (x, y, z) of the particle.

By this means we can specify where the particle is. As discussed, in order to specify when it is at a particular point in space we stretch our imagination further and imagine that in addition to having rulers to measure position, we also have at each point in space a clock, and that these clocks have all been synchronized in some way. The idea is that with these clocks we can tell when a particle is at a particular position in space simply by reading off the time indicated by the clock at that position.

According to our common sense notion of time, it would appear sufficient to have only one set of clocks filling all of space. Thus, no matter which set of moving rulers we use to specify the position of a particle, we always use the clocks belonging to this single vast set to tell us when a particle is at a particular position. In other words, there is only one time for all the position measuring set of rulers. This time is the same time independent of how the rulers are moving through space. This is the idea of universal or absolute time due to Newton. However, as Einstein was first to point out, this idea of absolute time is untenable, and that the measurement of time intervals (e.g. the time interval between two events such as two supernovae occurring at different positions in space) will in fact differ for observers in motion relative to each other. In order to prepare ourselves for this possibility, we shall suppose that for each possible set of rulers including those fixed relative to the ground, or those moving with a subatomic particle and so on, there are a different set of clocks. Thus the position measuring rulers carry their own set of clocks around with them. The clocks belonging to each set of rulers are of course synchronized with respect to each other. Later on we shall see how this synchronization can be achieved. The idea now is that relative to a particular set of rulers we are able to specify where a particle is, and by looking at the clock (belonging to that set of rulers) at the position of the particle, we can specify when the particle is at that position. Each possible collection of rulers and associated clocks constitutes what is known as a frame of reference or a reference frame.

2.3 Newton's First Law and Inertial Frame

Having established how we are going to measure the coordinates of a particle in space and time, we can now turn to considering how we can use these ideas to make a statement about the physical properties of space and time. To this end let us suppose that we have somehow placed a particle in the depths of space far removed from all other matter. It is reasonable to suppose that a particle so placed is acted on by no forces whatsoever.

The question then arises: What kind of motion is this particle undergoing? In order to determine this we have to measure its position as a function of time, and to do this we have to provide a reference frame. We could imagine all sorts of reference frames, for instance one attached to a rocket travelling in some complicated path. Under such circumstances, the path of the particle as measured relative to such a reference frame would be very complex. However, it is at this point that an assertion can be made, namely that for certain frames of reference, the particle will be travelling in a particularly simple fashion a straight line at constant speed. This is something that has not and possibly could not be confirmed experimentally, but it is nevertheless accepted as a true statement about the properties of the motion of particles in the absences of forces. In other words we can adopt as a law of nature, the following statement:

There exist frames of reference relative to which a particle acted on by no forces moves in a straight line at constant speed.

This essentially a claim that we are making about the properties of spacetime. It is also simply a statement of Newtons First Law of Motion. A frame of reference which has this property is called an *inertial frame of reference*, or just an inertial frame.

2.4 Galilean Transformation

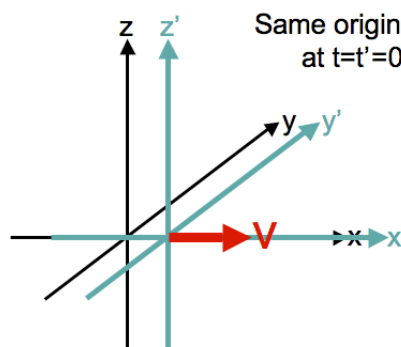


Figure 2.3

Now the question is: How to relate the coordinates of an event as observed from the point-of-view of one inertial reference frame to the coordinates of the same event as observed in some other? These transformation laws are essential if we are to compare the mathematical statements of the laws of physics in different inertial reference frames.

Consider the case where we have two inertial reference frames, one frame moving relative to the other in the x -direction with velocity v , as the Fig.2.3 illustrates.

Now, suppose the clocks in two reference frames are set such that when the origin coincide, both clocks read zero, i.e. $t = t' = 0$. If we accept the "common sense", i.e. the absolute time concept, then $t = t'$ always. Suppose now an event occurs at a point (x', y', z', t') in the green coordinate, then translating it back to the black coordinate, we see that this event occurs at

$$\begin{aligned}x &= x' + vt' \\y &= y' \\z &= z' \\t &= t' .\end{aligned}$$

These equations together are known as the ***Galilean Transformation***, and they tell us how the coordinates of an event in one inertial frame are related to the coordinates of the same event as measured in another frame moving with a constant velocity relative to the first. A simple exercise shows that we can also translate the unprimed coordinate to primed coordinate, by

$$\begin{aligned}x' &= x - vt \\y' &= y \\z' &= z \\t' &= t .\end{aligned}$$

Having proposed the existence of a special class of reference frames, the inertial frames of reference, and the Galilean transformation that relates the coordinates of events in such frames, we can now proceed further and study whether or not Newton's remaining laws of motion are indeed consistent with the principle of relativity. Perhaps it is better to illustrate this with an example.

Consider a 1-dimensional problem where two particles are connected by a spring. If the coordinate of the two particles are x_1, x_2 relative to some reference frame S , then by Newton's second law, the law of motion of the particle at x_1 is

$$m_1 \frac{d^2 x_1}{dt^2} = -k(x_1 - x_2 - l) \quad (2.2)$$

where k is the spring constant and l is the natural length of the spring, and m_1 the mass of the particle. Now consider the same point mass, but relative to a different reference frame S' , which travels with a velocity v relative to S . Then we have

$$x_1 = x'_1 + vt' \quad x_2 = x'_2 + vt' .$$

So then

$$\frac{d^2 x_1}{dt^2} = \frac{d^2 x'_1}{dt'^2} \quad x_2 - x_1 = x'_2 - x'_1 .$$

Substitute these into (2.2) we have that

$$m_1 \frac{d^2 x'_1}{dt'^2} = -k(x'_1 - x'_2 - l) .$$

In Newtonian mechanics $m_1 = m'_1$, so

$$m'_1 \frac{d^2 x'_1}{dt'^2} = -k (x'_1 - x'_2 - l).$$

This is exactly the same equation we obtained in the frame S , except now it's in S' and everything is primed. Continuing with this example, we can show that the momentum is conserved in all inertial reference frames. In frame S , the momentum is given by

$$P = m_1 \dot{x}_1 + m_2 \dot{x}_2.$$

In S' , we have

$$P' = m'_1 \dot{x}'_1 + m'_2 \dot{x}'_2 = m_1 \dot{x}_1 + m_2 \dot{x}_2 - (m_1 + m_2) v = P - (m_1 + m_2) v$$

which is also conserved. Therefore, we conclude that Newton's Laws of motion are identical in all inertial frames of reference.

The Newtonian principle of relativity had been successful until the advent of Maxwell's work in which he formulated a mathematical theory of electromagnetism which, amongst other things, provided a successful physical theory of light. Not unexpectedly, it was anticipated that the equations Maxwell derived should also obey the above Newtonian principle of relativity in the sense that Maxwell's equations should also be the same in all inertial frames of reference. Unfortunately, it was found that this was not the case. Maxwell's equations were found to assume completely different forms in different inertial frames of reference.

Maxwell's equations in vacuum read

$$\begin{aligned} \nabla \cdot \mathbf{E} &= 0 & \nabla \times \mathbf{E} &= -\frac{\partial \mathbf{B}}{\partial t} \\ \nabla \cdot \mathbf{B} &= 0 & \nabla \times \mathbf{B} &= \frac{1}{c^2} \frac{\partial \mathbf{E}}{\partial t}. \end{aligned}$$

We will show that Maxwell equations are not Galilean invariant. To simplify the problem, let us consider the case where S' is moving in the x direction with velocity u relative to S . Further, we assume

$$\mathbf{B} = B(x, t) \mathbf{k} \quad \mathbf{E} = E(x, t) \mathbf{j}.$$

That is, the electric field $\mathbf{E}$ only has nonzero component in y direction, and the strength of $\mathbf{E}$ only depends on x and t . Similarly for $\mathbf{B}$.

Now we explore how the field $\mathbf{E}'$, $\mathbf{B}'$ is like compared to $\mathbf{E}$, $\mathbf{B}$. Notice that a particle with charge q moving with a velocity $\mathbf{v}$ in S frame experiences a Lorentz force

$$\mathbf{F} = q(\mathbf{E} + \mathbf{v} \times \mathbf{B}).$$

Assuming $\mathbf{v} = v \mathbf{i}$ we get that the only nonzero component of $\mathbf{F}$ is:

$$F_y = q(E - vB).$$

Since the forces experienced by the particle in S and S' are Galilean invariant, we have that

$$q(E - vB) = q(E' - v'B').$$

Now using $v' = v - u$ we have that

$$E' - v'B' = (E' + uB') - vB' = E - vB.$$

Since v is arbitrary, setting $v = 0$ we find that

$$E = E' + uB' \quad B = B'.$$

This is the rule of how the field transform in this case. Now we consider the third Maxwell's equation, $\nabla \times \mathbf{B} = \frac{1}{c^2} \frac{\partial \mathbf{E}}{\partial t}$. In our case it becomes:

$$\frac{\partial B}{\partial x} = -\frac{1}{c^2} \frac{\partial E}{\partial t}.$$

Now we want to transform this equation into the primed coordinate S' , and see whether it remains the same form. Now, the chain rule gives

$$\frac{\partial B}{\partial x} = \frac{\partial B}{\partial x'} \frac{\partial x'}{\partial x} + \frac{\partial B}{\partial t'} \frac{\partial t'}{\partial x} \quad \frac{\partial E}{\partial t} = \frac{\partial E}{\partial t'} \frac{\partial t'}{\partial t} + \frac{\partial E}{\partial x'} \frac{\partial x'}{\partial t}.$$

Since $x' = x - ut$ and $t' = t$, we have that

$$\frac{\partial x'}{\partial x} = 1 \quad \frac{\partial x'}{\partial t} = -u \quad \frac{\partial t'}{\partial t} = 1 \quad \frac{\partial t'}{\partial x} = 0.$$

Therefore

$$\frac{\partial B}{\partial x} = \frac{\partial B}{\partial x'} \quad \frac{\partial E}{\partial t} = \frac{\partial E}{\partial t'} - u \frac{\partial E}{\partial x'}.$$

Plugging everything into the Maxwell equation, we see that the transformed equation is

$$\frac{\partial B'}{\partial x'} = -\frac{1}{c^2} \frac{\partial E'}{\partial t'} + \frac{u}{c^2} \left(\frac{\partial E'}{\partial x'} - \frac{\partial B'}{\partial t'} + u \frac{\partial B'}{\partial x'} \right).$$

We thus conclude that Maxwells equations are not invariant under Galilean transformations. This special frame S was assumed to be the one that defined the state of absolute rest as postulated by Newton, and that stationary relative to it was a most unusual entity, the ether. The ether was a substance that was supposedly the medium in which light waves were transmitted in a way something like the way in which air carries sound waves. Consequently it was believed that the behaviour of light, in particular its velocity, as measured from a frame of reference moving relative to the ether would be different from its behaviour as measured from a frame of reference stationary with respect to the ether. Since the earth is following a roughly circular orbit around the sun, then it follows that a frame of reference attached to the earth must at some stage in its orbit be moving relative to the ether, and hence a change in the velocity of light should be observable at some time during the year. From this, it should be possible to determine the velocity of the earth relative to the ether. An attempt was made to measure this velocity. This was the famous experiment of Michelson and Morley. Simply stated, they argued that if light is moving with a velocity c through the ether, and the Earth was at some stage in its orbit moving with a velocity v relative to the ether, then light should be observed to be travelling with a velocity $c' = c - v$ relative to the Earth.

Needless to say, on performing their experiment which was extremely accurate they found that the speed of light was always the same. Obviously

something was seriously wrong. Their experiments seemed to say that the earth was not moving relative to the ether, which was manifestly wrong since the earth was moving in a circular path around the sun, so at some stage it had to be moving relative to the ether. Many attempts were made to patch things up while still retaining the same Newtonian ideas of space and time. But all these attempts failed in various ways. It was Einstein who pointed the way out of the impasse, a way out that required a massive revision of our concepts of space, and more particularly, of time.

2.5 Special Relativity

Einstein's 1905 successful modification of Newtonian mechanics is called the special theory of relativity, or special relativity for short. To formulate it, Einstein assumed that the velocity of light had the same value c in all inertial frames—an assumption that from the present perspective is clearly motivated by the Michelson-Morley experiment. The assumption that the velocity of light is the same in every inertial frame requires a reexamination and ultimately the abandonment of the Newtonian idea of absolute time. One can see this most clearly by examining the idea of simultaneity. Two events are simultaneous if they occur at the same time. In Newtonian theory two events that are simultaneous in one inertial frame are simultaneous in every other inertial frame because there is a single absolute time. To see the impact on this idea of assuming the constancy of the velocity of light, consider the thought experiment illustrated by the following thought experiment:

Three observers A, B, O are riding a rocket of length L . Observer O is halfway between A and B . Now, A and B each emit light signals directed along the rocket toward O , and O receives the signal simultaneously. Which signal was emitted first? The answer depends on the inertial frame if the velocity of light is the same in all of them.

On the left of the figure above, the rocket is at rest. An observer at rest in this frame reasons as follows: The rocket is at rest and the two observers A and B are equal distances away from O . It therefore takes the same length of time for a light signal to propagate from A to O as it does from B to O . Since the signals reached O at the same instant, they were emitted simultaneously.

A different result is obtained in an inertial frame in which the rocket is moving, such as that shown on the right of the figure above. An observer at rest in this frame reasons as follows: The signals are received simultaneously by O . At earlier times when the signals were emitted B was always closer to O 's position at reception than A . Since *both signals travel with speed c* , the one from A was emitted earlier than the one from B because it has a longer distance to travel to reach O at the same instant as the one from B .

Thus, two events simultaneous in one inertial frame are not simultaneous in one moving with respect to the first if the velocity of light is the same in both. The Newtonian idea of time must be abandoned.

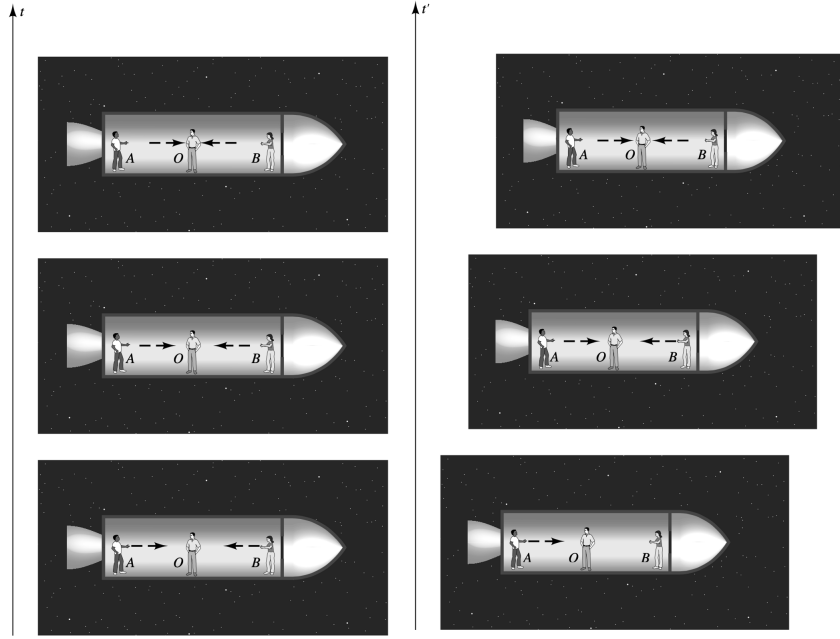


Figure 2.4

Spacetime

In view of the previous discussion of simultaneity, there is no reason to accept the assumption of Newtonian physics that the times of different inertial frames will agree. Rather, there is generally a different notion of time and simultaneity for each inertial frame. The correct geometric arena for physics is, therefore, not a separate space and absolute time but rather a four-dimensional unification of space and time called **spacetime**. Therefore, inertial frames are spanned by four Cartesian coordinates (t, x, y, z) , and a different inertial frame has a different set of four coordinates (t', x', y', z') .

To describe four-dimensional spacetime we first introduce a tool, which is so simple it appears trivial, but so powerful it is indispensable. This is the idea of a spacetime diagram. A spacetime diagram is a plot of two of the coordinate axes of an inertial frame two coordinate axes of spacetime. Since there are four axes and only two dimensions on a piece of paper, two or at most three of these axes can be drawn. It is convenient to use ct rather than t as an axis, because then both have the same dimension. A point P in spacetime can be called an **event** because an event occurs at a particular place at a particular time, that is, at a point in spacetime. A particle describes a curve in spacetime called a **world line**. It is the curve of positions of the particle at different instants, i.e., $x(t)$. The slope of the world line gives the ratio c/V^x for

$$\frac{d(ct)}{dx} = c \frac{dt}{dx} = c \left(\frac{dx}{dt} \right)^{-1} = \frac{c}{V^x}.$$

Geometry of Spacetime

The central assumption of special relativity is a geometry for spacetime. A geometry is specified by a line element that gives the distance between nearby points. It would be appropriate to begin a discussion of special relativity by

positing this line element. It would be appropriate to begin a discussion of special relativity by positing this line element. However, before doing that, consider a simple thought experiment that motivates the form of the line element and connects that with Einstein's assumption that the velocity of light is c in all inertial frames.

The thought experiment is illustrated in the figure below. Two parallel mirrors separated by a distance L are at rest in an inertial frame in which events are described by coordinates (t, x, y, z) . We take y to be the vertical direction between the mirrors and x the direction parallel to them. A light signal bounces back and forth between the mirrors; the right hand part of the figure below shows its world line in a spacetime diagram. A clock measures the time interval Δt between the event A of the departure of the light ray and the event C of its return to the same point in space. These two events are separated by coordinate intervals

$$\Delta t = 2L/c, \quad \Delta x = \Delta y = \Delta z = 0$$

in the inertial frame where the mirrors are at rest.

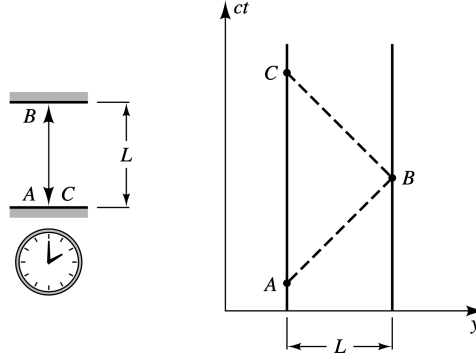


Figure 2.5

Analyze the same thought experiment in an inertial frame that is moving with speed V with respect to the (t, x, y, z) inertial frame along the negative x -direction parallel to the mirrors. Locate events in this frame by coordinates (t', x', y', z') with x' parallel to x . In this frame the mirrors are moving with speed V along the positive x' -direction, as illustrated in the figure below.

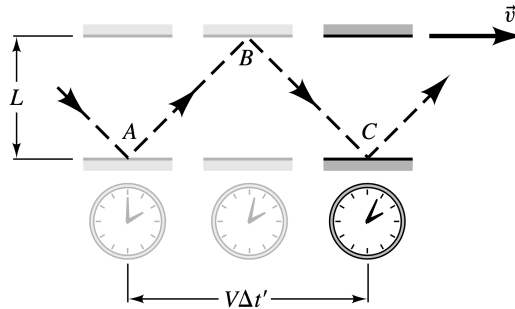


Figure 2.6

the light ray travels a distance $\Delta x' = V\Delta t'$ in the x' -direction between

between emission at A and return at C . The distance traveled in the y' -direction is L , assuming that the transverse distances are the same in both inertial frames. The total distance traveled between departure and return is, therefore,

$$2[L^2 + (\Delta x'/2)^2]^{1/2}.$$

Assuming with Einstein that the velocity of light is c in this inertial frame, the time of travel, $\Delta t'$, is this distance divided by c . Thus the coordinate intervals between A and C in this frame are

$$\Delta t' = \frac{2}{c} \sqrt{L^2 + \left(\frac{\Delta x'}{2}\right)^2}, \quad \Delta x' = V \Delta t', \quad \Delta y' = 0, \quad \Delta z' = 0.$$

Then straightforward computation gives

$$-(c\Delta t')^2 + (\Delta x')^2 = -4 \left[L^2 + (\Delta x'/2)^2 \right] + (\Delta x')^2 = -4L^2 = -(c\Delta t)^2.$$

Since $\Delta x = \Delta y = \Delta z = 0$ and $\Delta y' = \Delta z' = 0$, we can judiciously add them back into two sides of the equation above and find that

$$\begin{aligned} (\Delta s)^2 &\equiv -(c\Delta t)^2 + (\Delta x)^2 + (\Delta y)^2 + (\Delta z)^2 \\ &= -(c\Delta t')^2 + (\Delta x')^2 + (\Delta y')^2 + (\Delta z')^2. \end{aligned}$$

Therefore, we propose that the quantity

$$(\Delta s)^2 \equiv -(c\Delta t)^2 + (\Delta x)^2 + (\Delta y)^2 + (\Delta z)^2$$

is the same in both frames. The quantity $(\Delta s)^2$ is invariant under the change in inertial frames. And since the distance between points defining spacetime geometry must be the same in all systems of coordinates used to label points, and the principle of relativity requires that the line element that defines the distance should have the same form in all inertial frames, we are motivated to posit the line element

$$ds^2 = -c^2 dt^2 + dx^2 + dy^2 + dz^2 \quad (2.3)$$

as defining the geometry of 4-dimensional spacetime and the starting point of special relativity.

Light Cones

As we can see from (2.3), the geometry of spacetime is not 4-dimensional Euclidean geometry (whose invariant line element is given by $ds^2 = dx^2 + dy^2 + dz^2$). In particular, two points can be separated by distances whose square is positive, negative, or zero. We say two points are

$$\begin{aligned} \text{spacelike separated:} & \quad (\Delta s)^2 > 0 \\ \text{null separated:} & \quad (\Delta s)^2 = 0 \\ \text{timelike separated:} & \quad (\Delta s)^2 < 0. \end{aligned}$$

The locus of points that are null separated from a point P in spacetime is its **light cone**. Each point P in spacetime has a light cone. Light cones are an

important feature of the geometry of spacetime. The points that are timelike separated from P lie inside the light cone. Points that are spacelike separated from P lie outside the light cone.

The paths of light rays are straight lines in spacetime with constant slope corresponding to the speed of light, that is, along null world lines. At every point P along the world line of a light ray, the straight line is tangent to the light cone of that point. The distance between two points along a light ray is zero!

Particles with nonzero rest mass move along timelike world lines that are always inside the light cone of any point along their trajectory (see next figure). That way their velocity is always less than the speed of light at that point.

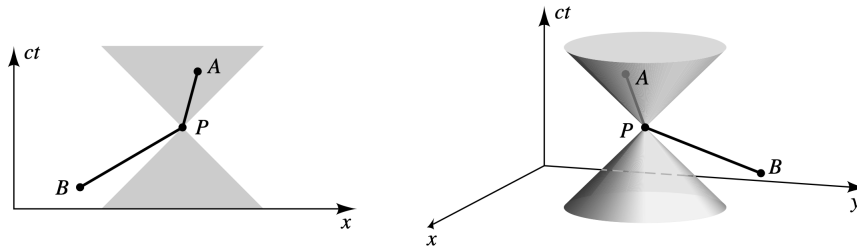


Figure 2.7: The 45° lines from point P are the set of points that are null separated from it. They are the intersection of P 's light cone with this slice. Point A is timelike separated from P , as are all points in the shaded wedges. The upper shaded wedge is the inside of the future light cone; the lower shaded wedge is the inside of the past light cone. The unshaded area is the outside of the light cone. Point B is spacelike separated from P as are all the points in the unshaded wedges. The figure at right shows the same point P but with one more spatial dimension. The light cone is the locus of points that would be traced out by a pulse of light emitted at P or converging on it. The surface of the pulse would be an expanding or contracting sphere in three spatial dimensions. In this reduced number of dimensions, it appears as the increasing circular cross section of the cone.

To measure the distance along a particle's world line, it is convenient to introduce

$$d\tau^2 = -\frac{ds^2}{c^2}.$$

Then $d\tau$ is real with units of time. Thus a clock moving along a timelike curve measures the distance τ along it. An alternative name for this distance is the **proper time** the time that would be measured by a clock carried along the world line. To see this, notice that if the coordinate frame is at rest, i.e. $dx^2 = dy^2 = dz^2 = 0$, then $d\tau^2 = dt^2$ hence $d\tau = dt$. Therefore, in the rest frame, proper time corresponds to the "usual" time measured by a clock traveled along a world line. And since ds^2 is invariant in any inertial frame, so is $d\tau^2$.

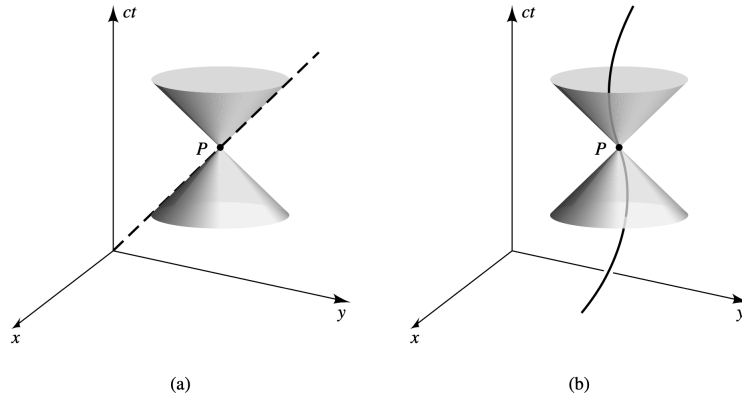


Figure 2.8: (a) The path of a light ray must be tangent to the light cone at every point along its trajectory. (b) The timelike path of a particle must lie inside the light cone at every point along its trajectory. These are the invariant ways of saying that light rays move at speed c and particles move with speed less than c in every inertial frame.

Lorentz Transformation

In this subsection we will derive Lorentz transformation, which relates two inertial frames (t', x', y', z') and (t, x, y, z) in relative motion. We present two methods, the first one is more intuitive, and the second one more mathematical.

Method I:

We stick with the idea of two inertial frames, S and S' , moving with relative speed v . For simplicity, we will start by ignoring the directions y and z which are perpendicular to the direction of motion. Both inertial frames come with Cartesian coordinates: (x', t') for S' and (x, t) for S . We want to know how these are related. The most general possible relationship takes the form

$$x' = f(x, t) \quad , \quad t' = g(x, t).$$

But this is very general. Fortunately, there are several observations that allow us to restrict the form of f, g immediately:

- The first is that the law of inertia holds; left alone in an inertial frame, a particle will travel at constant velocity. Drawn in the (x, t) plane, the trajectory of such a particle is a straight line. Since both S and S' are inertial frames, the map $(x, t) \mapsto (x', t')$ must map straight lines to straight lines; such maps are, by definition, linear. The functions f and g must therefore be of the form

$$x' = \alpha_1 x + \alpha_2 t \quad , \quad t' = \alpha_3 x + \alpha_4 t$$

where α_i can be functions of v .

- Secondly, we use the fact that S' is travelling at speed v relative to S . This means that an observer sitting at the origin, $x' = 0$ of S' moves

along the trajectory $x = vt$ in S . In other words, the points $x = vt$ must map to $x' = 0$. Together with the requirement that the transformation is linear, this restricts the transformation to be of the form

$$x' = \gamma(x - vt) \quad (2.4)$$

for some γ , which can be a function of v , which we denote by γ_v .

- Thirdly, we claim that γ_v must be an even function of v . That is, $\gamma_v = \gamma_{-v}$. Consider the frames $\tilde{S}$ and $\tilde{S}'$ which are exactly like S and S' except that we measure x -coordinate the the negative direction. This means that

$$\tilde{x} = -x \quad \tilde{x}' = -x' \quad \tilde{t} = t \quad \tilde{t}' = t'.$$

While S' is moving with velocity $+v$ relative to S , $\tilde{S}'$ is moving with velocity $-v$ with respect to $\tilde{S}$ simply because we measure things in the opposite direction. That means that

$$\tilde{x}' = \gamma_{-v}(\tilde{x} + v\tilde{t}).$$

Compare this with (2.4) and plugging in the relations in the previous equation, we have $\gamma_v = \gamma_{-v}$.

We can also look at things in the perspective of S , which moves with velocity $-v$ relative to S' . Same argument applies and we get

$$x = \gamma(x' + vt'). \quad (2.5)$$

Finally, we used the assumption that the speed of light is the same in every inertial frame. So in S , a light ray has trajectory $x = ct$, and in S' , we demand it to have the trajectory $x' = ct'$. Substitute these trajectories into (2.4) and (2.5), we have

$$ct' = \gamma(c - v)t \quad \text{and} \quad ct = \gamma(c + v)t'.$$

Solving for γ gives

$$\gamma = \sqrt{\frac{1}{1 - v^2/c^2}}. \quad (2.6)$$

Therefore S and S' are related by

$$\begin{aligned} t' &= \gamma(t - vx/c^2) \\ x' &= \gamma(x - vt) \\ y' &= y \\ z' &= z. \end{aligned}$$

The inverse transformation is obtained just by changing v to $-v'$:

$$\begin{aligned} t &= \gamma(t + vx'/c^2) \\ x &= \gamma(x' + vt') \\ y &= y' \\ z &= z'. \end{aligned}$$

Method II: Lorentz Group

As discussed before, the line element (2.3) should be invariant under the transformation that we are interested in, just like in Euclidean geometry the line element $ds^2 = dx^2 + dy^2 + dz^2$ is invariant under a translation or a rotation.

We use greek letter $\mu, \nu = 0, 1, 2, 3$ as indices of coordinate, say x^μ , where $x^0 = ct$ and $x^1 = x, x^2 = y, x^3 = z$. Usually instead of x, y, z we will let x^1, x^2, x^3 be spatial coordinates, and we use letters $i, j, k = 1, 2, 3$ to run through spatial indices.

We further introduce the diagonal matrix

$$\eta_{\mu\nu} = \text{diag}(-1, 1, 1, 1) = \begin{pmatrix} -1 & & & \\ & 1 & & \\ & & 1 & \\ & & & 1 \end{pmatrix}.$$

Then

$$\Delta s^2 = (\Delta \mathbf{x})^T \eta (\Delta \mathbf{x})$$

where $\Delta \mathbf{x} = (c\Delta t, \Delta x, \Delta y, \Delta z)$. We could use the index notation, but here we are more concerned with the matrix multiplication. Then since we have seen that Lorentz transformation must be linear in Method I, the Lorentz transformation corresponds to a matrix Λ such that

$$\mathbf{x}' = \Lambda \mathbf{x}.$$

Since the line element is invariant under Lorentz transformation, we must have

$$\Delta s^2 = (\Delta \mathbf{x})^T \eta (\Delta \mathbf{x}') = (\Delta \mathbf{x})^T \eta (\Delta \mathbf{x}') = (\Delta \mathbf{x})^T \Lambda^T \eta \Lambda (\Delta \mathbf{x}).$$

Therefore we have

$$\eta = \Lambda^T \eta \Lambda. \quad (2.7)$$

The matrices Λ satisfying (2.7) are called **Lorentz transformations** (this time formally). Lets try to understand the solutions to this. We can start by counting how many we expect. The matrix Λ has $4 \times 4 = 16$ components. Both sides of equation (2.7) are symmetric matrices, which means that the equation only provides 10 constraints on the coefficients of Λ . We therefore expect to find $16 - 10 = 6$ independent solutions.

The solution of (2.7) fall into 2 classes. The first class is very familiar. We consider the solution of the form

$$\Lambda = \left(\begin{array}{c|ccc} 1 & 0 & 0 & 0 \\ \hline 0 & & & \\ 0 & & R & \\ 0 & & & \end{array} \right)$$

where R are 3×3 matrices. Thus these matrices leave the time component intact and change the spatial component. The condition (2.7) reduces to

$$R^T R = 1.$$

We know that the R s that satisfies this equation are precisely the spatial rotations. The remaining three solutions to (2.7) are the Lorentz boosts that we discussed in Method I. The boost along the x -axis is given by

$$\Lambda = \left(\begin{array}{cc|cc} \gamma & -\gamma v/c & 0 & 0 \\ -\gamma v/c & \gamma & 0 & 0 \\ \hline 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{array} \right).$$

Two remaining solutions are boosts in y, z directions.

It can be checked that the set of all Λ s satisfying (2.7) form a group. This is called the **Lorentz group**, denoted by $O(3, 1)$. One can also check that $\det \Lambda^2 = 1$, which implies $\det \Lambda = \pm 1$. So $O(3, 1)$ splits into two parts, the part where $\det \Lambda = 1$ are denoted by $SO(3, 1)$.

We will now see how to use Lorentz group to derive the addition of velocities. For a boost with velocity v in x -direction, we can focus on the upper left part:

$$\Lambda[v] = \begin{pmatrix} \gamma & -\gamma v/c \\ -\gamma v/c & \gamma \end{pmatrix}.$$

Now if we combine two boosts, both in x -direction, we have that

$$\Lambda[v_1] \Lambda[v_2] = \begin{pmatrix} \gamma_1 & -\gamma_1 v_1/c \\ -\gamma_1 v_1/c & \gamma_1 \end{pmatrix} \begin{pmatrix} \gamma_2 & -\gamma_2 v_2/c \\ -\gamma_2 v_2/c & \gamma_2 \end{pmatrix}.$$

A little algebra shows

$$\Lambda[v_1] \Lambda[v_2] = \Lambda \left[\frac{v_1 + v_2}{1 + v_1 v_2 / c^2} \right]. \quad (*)$$

What does this piece of math tell us? Imagine now a particle moves with a constant velocity u' in the S' frame, which moves with a constant velocity v relative to the frame S . What is the velocity u of the particle in S ? Newtonian theory will simply give $u = u' + v$, but this is wrong when $u' = c$. This situation, then, amounts to applying two boosts, with $v_1 = u'$ and $v_2 = v$. Thus we get

$$u = \frac{u' + v}{1 + u'v/c^2}.$$

Notice that when $u' = c$ we get $u = c$, which is just what we want.

Is there a better way to arrive at (*)? There is, actually. Recall that when we add rotations

$$R[\theta] = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

we have

$$R[\theta_1] R[\theta_2] = R[\theta_1 + \theta_2].$$

But the nice addition rule only worked because we were clever in parameterising our rotation by an angle. In the case of Lorentz boosts, there is a similarly clever parameterisation. Instead of using the velocity v , we define the **rapidity** φ by

$$\gamma = \cosh \varphi.$$

Then

$$\sinh \varphi = \sqrt{\cosh^2 \varphi - 1} = \sqrt{\gamma^2 - 1} = \frac{v\gamma}{c}.$$

Therefore we can write a Lorentz boost by

$$\Lambda[\varphi] = \begin{pmatrix} \cosh \varphi & -\sinh \varphi \\ -\sinh \varphi & \cosh \varphi \end{pmatrix}.$$

Now we have

$$\Lambda[\varphi_1] \Lambda[\varphi_2] = \Lambda[\varphi_1 + \varphi_2].$$

Also, combining the two equation above we see that

$$\tanh \varphi = \frac{\sinh \varphi}{\cosh \varphi} = \frac{v\gamma/c}{\gamma} = \frac{v}{c}$$

or $v = c \tanh \varphi$.

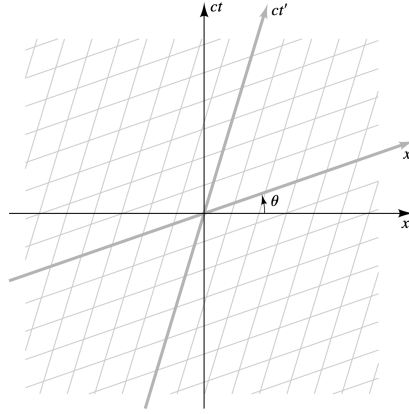


Figure 2.9: A Lorentz boost as a change of coordinates (linear transformation) on a spacetime diagram. The rapidity θ is a measure of the velocity between the two frames.

Time Dilation and the Twin Paradox

Just the few facts about the geometry of the spacetime of special relativity can be put to work to derive some its most famous consequences. First is the phenomenon of time dilation. The proper time, τ_{AB} , between any two points A and B on timelike world line can be computed by integrating the differential $d\tau$:

$$\begin{aligned} \tau_{AB} &= \int_A^B d\tau = \int_A^B [dt^2 - (dx^2 + dy^2 + dz^2)/c^2]^{1/2} \\ &= \int_{t_A}^{t_B} dt \left\{ 1 - \frac{1}{c^2} \left[\left(\frac{dx}{dt} \right)^2 + \left(\frac{dy}{dt} \right)^2 + \left(\frac{dz}{dt} \right)^2 \right] \right\}^{1/2}. \end{aligned}$$

More compactly,

$$\tau_{AB} = \int_{t_A}^{t_B} dt' [1 - v^2(t')/c^2]^{1/2}.$$

Since the integrand $[1 - v^2/c^2]^{1/2} < 1$, we have that

$$\tau_{AB} < \int_{t_A}^{t_B} dt' = t_B - t_A.$$

This is the phenomenon of **time dilation**, summarized by the slogan "moving clocks run slow". It will frequently be useful to consider the differential form

$$d\tau = dt\sqrt{1 - v^2/c^2}.$$

A crucial fact is that for more general trajectories, the proper time and the coordinate time are different (although proper time is always that measured by the clock carried by an observer along the trajectory).

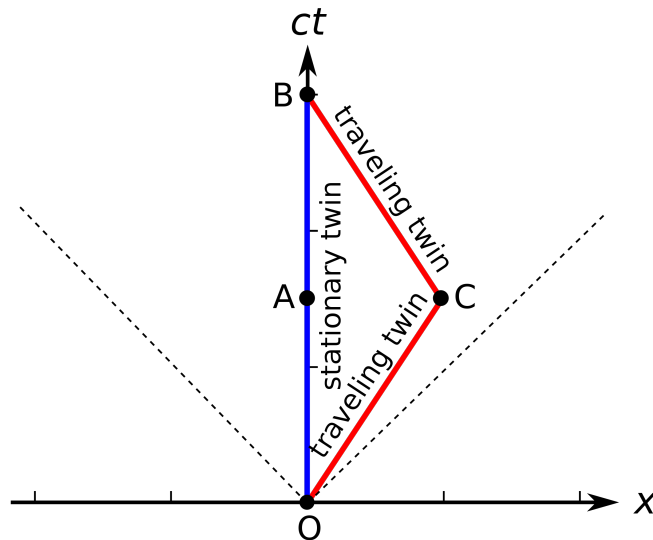


Figure 2.10: the Twin Paradox.

Consider two trajectories between event O and B , one straightline crossing a halfway point marked A , and another traveled by an observer moving away from O at a constant velocity v to a point C then back at a constant velocity $-v$ to intersect at the event B . Suppose the coordinates are given by

$$A = \left(\frac{1}{2}\Delta t, 0\right) \quad C = \left(\frac{1}{2}\Delta t, \Delta x\right) \quad B = (\Delta t, 0)$$

with $\Delta x = \frac{1}{2}v\Delta t$. Then clearly

$$\tau_{OC} = \sqrt{\left(\frac{1}{2}\Delta t\right)^2 - (\Delta x)^2} = \frac{1}{2}\sqrt{1 - v^2}\Delta t < \frac{1}{2}\Delta t.$$

Therefore

$$\tau_{OCB} = \sqrt{1 - v^2}\Delta t < \Delta t.$$

Even though two observers begin and end at the same points in spacetime, they have aged different amounts. This is the famous "twin paradox". The truth is straightforward: a nonstraight path in spacetime has a different interval than a straight path, just as a nonstraight path in space has a different length than a straight one.

Length Contraction

We will conclude this section with some relativistic effects. The first effect is the **length contraction**. The length of any object in a moving reference

frame will be contracted. To see this, notice that the length of a rod in a frame is the distance between two events happening in the same time in that frame. So say in a rest frame, the length of a rod is given by $L = x_2 - x_1$. That is, two ends of the rod are at coordinates $(0, x_1)$ and $(0, x_2)$. Then consider another frame moving with velocity v relative to the rest frame, what is the length of the same rod L measured in the moving frame? Well, $L' = x'_2 - x'_1$ and

$$L' = x'_2 - x'_1 = \gamma(x_2 - vt_2) - \gamma(x_1 - vt_1) = \gamma(x_2 - x_1) = \gamma L.$$

The second step is just the formula for Lorentz transformation derived before, and in the third step we used the fact that the two measurements in the rest frame are made simultaneously: $t_1 = t_2$. Another way to write the formula above is

$$L = \frac{L'}{\gamma}.$$

If $v \approx c$ then $1/\gamma \approx 0$. That is to say, that an object at rest might be measured to be 200 feet long; yet the same object when moving at relativistic speeds relative to the observer/measurer would have a measured length which is less than 200 ft. See [this link](#) for an interesting demonstration.

2.6 Vector, Dual Vectors, and Tensors

In this section we will introduce the tensor language for relativity.

Vectors

We will start with vectors. In spacetime, vectors are 4-dimensional, and usually called 4-vectors. This turns out to make quite a bit of difference. For example, there is no such thing as a cross product between two 4-vectors.

One of the most important things to emphasize is that each vector is located at a given point in spacetime. In $\mathbf{R}^n$ one can slide vectors from point to point. But once we introduce curvature, we lose the ability to move vectors uniquely around a manifold. Rather, to each point p in spacetime we associate the set of all possible vectors located at p , called the **tangent space** at p , or T_p . Each tangent space T_p is a vector space.

Let us imagine at each tangent space T_p we set up a basis of four vectors e_μ . Then any abstract vector $V \in T_p$ can be written as

$$V = V^\mu e_\mu$$

where $V^\mu \in \mathbf{R}$ are coefficients. They are called the components of the vector V with respect to the basis e_μ . An example of a vector in spacetime is the tangent vector to a curve. A parameterized curve or path through spacetime is specified by the coordinates as a function of the parameter, for example $x^\mu(\lambda)$. The tangent vector $V(\lambda)$ has components

$$V^\mu = \frac{dx^\mu}{d\lambda}.$$

The entire vector is $V = V^\mu e_\mu$. Under the Lorentz transformation the coordinates change by

$$x^{\mu'} = \Lambda^{\mu'}_{\nu} x^\nu.$$

Notice that we do the prime in the indices μ' , not $(x')^\mu$, because the geometrical object is the same, the only things that have changed are the components, due to a change of coordinate. Taking derivatives on both sides gives the vector transformation rule under the Lorentz transformation:

$$V^{\mu'} = \Lambda^{\mu'}_{\nu} V^{\nu}.$$

We have the component rule. But how does the basis vector change? Since the vector V even when we changed the coordinates (components are changed, but not the vector itself), we have

$$V = V^{\mu} e_{\mu} = V^{\nu'} e_{\nu'} = \Lambda^{\nu'}_{\mu} V^{\mu} e_{\nu'}.$$

But this relation must hold for every V^{μ} , so we conclude

$$e_{\mu} = \Lambda^{\nu'}_{\mu} e_{\nu'}.$$

To get the new basis $e_{\nu'}$ in terms of the old one e_{μ} , we should multiply the inverse of $\Lambda^{\nu'}_{\mu}$. But the inverse of a Lorentz transformation from the unprimed coordinate to the prime coordinate is itself a Lorentz transformation. Therefore, we introduce the following notation for the inverse of $\Lambda^{\nu'}_{\mu}$, which is the same symbol but with primed and unprimed indices switched. That is, $\Lambda^{\mu'}_{\nu}$ has an inverse $\Lambda^{\rho}_{\sigma'}$. Or

$$\Lambda^{\mu}_{\nu'} \Lambda^{\nu'}_{\rho} = \delta^{\mu}_{\rho},$$

$$\Lambda^{\sigma'}_{\lambda} \Lambda^{\lambda}_{\tau'} = \delta^{\sigma'}_{\tau'}.$$

Thus multiplying both sides of the previous equation by $\Lambda^{\mu}_{\nu'}$ we get

$$e_{\nu'} = \Lambda^{\mu}_{\nu'} e_{\mu}.$$

Dual Vectors

Let V be a finite dimensional vector space. Let V^* be its dual. If $\{e_{\mu}\} \subset V$ is a basis of the vector space, then there exists a naturally induced basis for the dual space $\{e^{\mu}\} \subset V^*$, defined by $e^{\mu}(e_{\nu}) = \delta^{\mu}_{\nu}$. Notice that as vector space, there is a natural isomorphism from V^{**} to V , given by the map

$$\varphi : V \rightarrow V^{**} \quad v \mapsto E_v$$

where

$$E_v(\phi) = \phi(v)$$

for all $\phi \in V^*$. It is easy to verify that this map $\varphi : v \mapsto E_v$ is linear. Also, φ is injective: Say $E_v = E_w$. Then

$$E_v(e^{\mu}) = e^{\mu}(v) = v^{\mu} = w^{\mu} = e^{\mu}(w) = E_w(e^{\mu}).$$

Hence $v = w$. And since $\dim V = \dim V^* = \dim(V^*)^*$, we see (by rank-nullity theorem) that φ is also surjective. Hence φ is indeed an isomorphism. Since $V^{**} \cong V$, we shall make no distinction between $v \in V$ and $E_v \in V^{**}$, and we simply denote v for E_v in V^{**} . With this notation, we see that

$$e_{\nu}(e^{\mu}) := E_{e_{\nu}}(e^{\mu}) = e^{\mu}(e_{\nu}) = \delta^{\mu}_{\nu}.$$

Now let $\omega \in V^*$, or given a tangent space T_p , let $\omega \in T_p^*$. Then we can use the same argument to see how dual vector transforms. The answer is

$$\omega_{\mu'} = \Lambda^\nu_{\mu'} \omega_\nu.$$

And for basis dual vectors,

$$e^{\rho'} = \Lambda^{\rho'}_{\sigma} e^{\sigma}.$$

This should come no surprise, since there is really no other way it can be given how the indices are placed. In spacetime, the simplest example of a dual vector is the gradient of a scalar function. This is the set of partial derivatives with respect to the spacetime coordinate, which we denote by

$$d\phi = \frac{\partial \phi}{\partial x^\mu} e^\mu.$$

You may think the notation $d\phi$ is confusing, since it is the same notation we use for a differential, or a line element. Turns out that they are really the same thing. As we will see later. But for now bear with me for a moment. The components transforms according to the chain rule:

$$\frac{\partial \phi}{\partial x^{\mu'}} = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial \phi}{\partial x^\mu} = \Lambda^\mu_{\mu'} \frac{\partial \phi}{\partial x^\mu}.$$

In the last step we have used the fact that since $x^\mu = \Lambda^\mu_{\mu'} x^{\mu'}$, the differential must transform like $dx^\mu = \Lambda^\mu_{\mu'} dx^{\mu'}$. And by chain rule $\frac{\partial x^\mu}{\partial x^{\mu'}} = \Lambda^\mu_{\mu'}$. We have seen that dual vectors act on vectors in a natural way:

$$\omega(V) = \omega_\mu e^\mu (V^\nu e_\nu) = \omega_\mu V^\nu e^\mu (e_\nu) = \omega_\mu V^\nu \delta_\nu^\mu = \omega_\mu V^\mu \in \mathbf{R}.$$

Therefore gradient can act on tangent vector in the same way:

$$d\phi(V) = \partial_\mu \phi \frac{\partial x^\mu}{\partial \lambda} = \frac{d\phi}{d\lambda}$$

by the chain rule.

Tensors

Now, a tensor T of rank (k, l) is just a multi-linear map

$$T : \underbrace{V^* \times \dots \times V^*}_k \times \underbrace{V \times \dots \times V}_l \rightarrow \mathbf{R}.$$

We can add tensors of the same rank and multiply tensor by a scalar, in the obvious ways. This makes tensors of rank (k, l) into a vector space. We can also combine a tensor T of rank (k, l) and a tensor T' of rank (k', l') into a new tensor $T \otimes T'$ of rank $(k + k', l + l')$ by

$$\begin{aligned} (T \otimes T')(\omega^1, \dots, \omega^k, \omega^{k+1}, \dots, \omega^{k+k'}, v_1, \dots, v_l, v_{l+1}, \dots, v_{l+l'}) \\ = T(\omega^1, \dots, \omega^k, v_1, \dots, v_l) \cdot T'(\omega^{k+1}, \dots, \omega^{k+k'}, v_{l+1}, \dots, v_{l+l'}). \end{aligned}$$

We can also define the components of a tensor T of rank (k, l) , after having chosen a basis $\{e_\mu\}$ for V .

$$T^{i_1, \dots, i_k}_{j_1, \dots, j_l} = T(e^{i_1}, \dots, e^{i_k}, e_{j_1}, \dots, e_{j_l}).$$

Notice that we can expand the tensor as a linear combination of basis elements:

$$\begin{aligned} T(\omega^1, \dots, \omega^k, v_1, \dots, v_l) &= T(\omega_{i_1}^1 e^{i_1}, \dots, \omega_{i_k}^k e^{i_k}, v_1^{j_1} e_{j_1}, \dots, v_l^{j_l} e_{j_l}) \\ &= \omega_{i_1}^1 \cdots \omega_{i_k}^k v_1^{j_1} \cdots v_l^{j_l} T(e^{i_1}, \dots, e^{i_k}, e_{j_1}, \dots, e_{j_l}) \\ &= T^{i_1, \dots, i_k}_{j_1, \dots, j_l} e_{i_1} \otimes \cdots \otimes e_{i_k} \otimes e^{j_1} \otimes \cdots \otimes e^{j_l} (\omega^1, \dots, \omega^k, v_1, \dots, v_l). \end{aligned}$$

Here we used the identification of double dual vectors with vectors mentioned above, and allowing vectors act on dual vectors. The element of the form

$$e_{i_1} \otimes \cdots \otimes e_{i_k} \otimes e^{j_1} \otimes \cdots \otimes e^{j_l}$$

constitute a basis for the vector space of rank (k, l) tensors. If our vector space has dimension n , then this vector space of (k, l) -tensors has dimension n^{k+l} .

Finally, the transformation of tensor components under Lorentz transformations can be derived by applying what we already know about transformation of basis vectors and dual vectors. The answer is just what you would expect from index placement:

$$T^{\mu'_1 \cdots \mu'_k}_{\nu'_1 \cdots \nu'_l} = \Lambda^{\mu'_1}_{\mu_1} \cdots \Lambda^{\mu'_k}_{\mu_k} \Lambda^{\nu_1}_{\nu'_1} \cdots \Lambda^{\nu_l}_{\nu'_l} T^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l}.$$

This is why tensor is so good! Because they can make our equations independent of coordinate transformation. Suppose we have a tensor equation, say

$$T^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l} = S^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l}.$$

Then we move the RHS to the LHS and get

$$T^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l} - S^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l} = 0.$$

This difference is a tensor, which is zero in all of its component. How does this tensor transform? We have

$$\begin{aligned} T^{\mu'_1 \cdots \mu'_k}_{\nu'_1 \cdots \nu'_l} - S^{\mu'_1 \cdots \mu'_k}_{\nu'_1 \cdots \nu'_l} &= \Lambda^{\mu'_1}_{\mu_1} \cdots \Lambda^{\mu'_k}_{\mu_k} \Lambda^{\nu_1}_{\nu'_1} \cdots \Lambda^{\nu_l}_{\nu'_l} (T^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l} - S^{\mu_1 \cdots \mu_k}_{\nu_1 \cdots \nu_l}) \\ &= \Lambda^{\mu'_1}_{\mu_1} \cdots \Lambda^{\mu'_k}_{\mu_k} \Lambda^{\nu_1}_{\nu'_1} \cdots \Lambda^{\nu_l}_{\nu'_l} (0) = 0. \end{aligned}$$

Therefore, if a tensor equation holds in one frame, it holds in every frame. Remember all those wonderful things that you first learned about in special relativity: time dilation and length contraction and twins and spaceships so on. Youll never have to worry about those again. From now on, you can guarantee that youre working with a theory consistent with relativity by ensuring two simple things:

- That you only deal with tensors.
- That the indices match up on both sides of the equation.

Its sad, but true. Its all part of growing up and not having fun anymore.

We can multiply two tensors to get a third tensor. For example,

$$U^\mu_\nu = T^{\mu\rho}_\sigma S^\sigma_{\rho\nu}$$

is a perfectly good (1,1) tensor. There is a particular type of tensor we want to introduce, the (0,2) tensors. They are linear maps that take two vectors and

give a real number. These are sometimes called inner product. For example, the inner product in Euclidean space, which we denote by $E(\cdot, \cdot)$, is given by

$$E(V, W) = V^1 W^1 + V^2 W^2 + V^3 W^3 = \delta_{ij} V^i V^j.$$

In Minkowski spacetime, this is the familar 4-vector inner product:

$$\eta(V, W) = \eta_{\mu\nu} V^\mu W^\nu.$$

We refer to two vectors whose inner product vanishes as **orthogonal**. Since inner products are scalars, which are invariant under Lorentz transformations, the basis vector of any Cartesian inertial frame, which are chosen to be orthogonal, remains orthogonal after a Lorentz transformation.

We introduce the **inverse metric** $\eta^{\mu\nu}$, a type of (2,0) tensor defined as the inverse of the metric:

$$\eta^{\mu\nu} \eta_{\nu\rho} = \eta_{\nu\rho} \eta^{\nu\mu} = \delta_\rho^\mu.$$

Operation of Tensors

One can **contract** a (k, l) tensor and make it into a $(k-1, l-1)$ tensor by summing over one upper and one lower index:

$$S^{\mu\rho}{}_\sigma = T^{\mu\nu\rho}{}_{\sigma\nu}.$$

The inverse metric can be used to **raise and lower indices** of tensors. Given a tensor $T^{\alpha\beta}{}_{\gamma\delta}$, we can use metric to define new tensors, which we choose to denote by the same letter T :

$$\begin{aligned} T^{\alpha\beta\mu}{}_\delta &= \eta^{\mu\gamma} T^{\alpha\beta}{}_{\gamma\delta} \\ T_\mu{}^\beta{}_{\gamma\delta} &= \eta_{\mu\alpha} T^{\alpha\beta}{}_{\gamma\delta} \\ T_{\mu\nu}{}^{\rho\sigma} &= \eta_{\mu\alpha} \eta_{\nu\beta} \eta^{\rho\gamma} \eta^{\sigma\delta} T^{\alpha\beta}{}_{\gamma\delta} \\ V_\mu &= \eta_{\mu\nu} V^\nu \\ \omega^\mu &= \eta^{\mu\nu} \omega_\nu. \end{aligned}$$

And this gives us the right to simultaneously raise and lower a pair of indices being contracted over:

$$A^\lambda B_\lambda = \eta^{\lambda\rho} A_\rho \eta_{\lambda\sigma} B^\sigma = \delta_\sigma^\rho A_\rho B^\sigma = A_\sigma B^\sigma.$$

A tensor is **symmetric** in any of its indices if it is unchanged under exchange of those indices. For example, we say S is symmetric in its first two indices if

$$S_{\mu\nu\rho} = S_{\nu\mu\rho}.$$

And if

$$S_{\mu\nu\rho} = S_{\mu\rho\nu} = S_{\rho\mu\nu} = S_{\nu\rho\mu} = S_{\rho\nu\mu},$$

we say S is symmetric in all of its indices. A tensor is **antisymmetric** in any of its indices if it changes sign when those indices are exchanged. Thus

$$A_{\mu\nu\rho} = -A_{\rho\nu\mu}$$

means A is antisymmetric in its first and third indices. Notice that it makes no sense to exchange an upper and a lower indices, so do not think of δ_ν^μ as a symmetric tensor. Given any tensor, we can **symmetrize** any number of its upper and lower indices: just take the sum of all permutations of the relevant indices and divide by the number of terms:

$$T_{(\mu_1\mu_2\cdots\mu_n)\rho}{}^\sigma = \frac{1}{n!} \sum_{\tau \in S_n} T_{\tau(\mu_1\cdots\mu_n)\rho}{}^\sigma$$

where S_n denotes the symmetric group of n letters and $\tau \in S_n$ is a permutation on n letters. The reason we divide by the factor $n!$ is that if T is already symmetric, then if we symmetrize it we want to get the same tensor. We can also **antisymmetrize** a tensor by summing over alternatively:

$$T_{[\mu_1\mu_2\cdots\mu_n]\rho}{}^\sigma = \frac{1}{n!} \sum_{\tau \in S_n} \text{sgn}(\tau) T_{\tau(\mu_1\cdots\mu_n)\rho}{}^\sigma$$

where $\text{sgn}(\tau)$ is 1 if τ is an even permutation, and is -1 if τ is an odd permutation. A tensor that has been antisymmetrized is antisymmetric, in the corresponding indices: If we switch two indices, then an even permutation will become an odd permutation, and an odd permutation will become an even permutation, so $\text{sgn}(\tau)$ will change sign for all $\tau \in S_n$. Therefore this result in a total of a minus sign.

Sometimes we want to (anti)symmetrize indices that are not next to each other. In this case we use vertical bars to denote indices not included in the sum:

$$T_{(\mu|\nu|\rho)} = \frac{1}{2} (T_{\mu\nu\rho} + T_{\rho\nu\mu}).$$

For a (1,1) tensor $X^\mu{}_\nu$, the **trace** is a scalar, often denoted by leaving off the indices, which is simply the contraction

$$X = X^\mu{}_\mu.$$

For a (0,2) tensor $Y_{\lambda\nu}$, to compute the trace, we must first raise an index:

$$Y^\mu{}_\nu = \eta^{\mu\lambda} Y_{\lambda\nu}$$

then we do the sum:

$$Y = Y^\lambda{}_\lambda = \eta^{\mu\nu} Y_{\mu\nu}.$$

Gradient, Normal Vectors, and Flux

I want to conclude this section by pointing out the connection between gradient, differentials, one forms, normal vectors, and flux. First, given a choice of coordinate x^α (we can do this in flat spacetime, but it will be more difficult on manifolds). Then consider the gradient dx^μ of the scalar function x^μ , then we have

$$dx^\mu = \frac{\partial x^\mu}{\partial x^\nu} e^\nu = \delta^\mu{}_\nu e^\nu = e^\mu.$$

So it turned out that our basis for the covector are exactly dx^μ s. This turns out to still be true on curved manifolds. Then the gradient of an arbitrary scalar function ϕ now becomes

$$d\phi = \frac{\partial \phi}{\partial x^\mu} e^\mu = \frac{\partial \phi}{\partial x^\mu} dx^\mu.$$

Holy crap! This is the chain rule for differentiating ϕ , which we think of as a infinitesimal change $d\phi$ is given by the sum of infinitesimal dx^μ s times the respective rate of change $\partial_\mu\phi$. But really this is an informal notion of gradient written as a linear combination of dual vector basis dx^μ , where the coefficients are given by exactly $\partial_\mu\phi$.

Now the key point: The gradient $d\phi$ defines a constant surface $\phi = c$ for some constant c , since it defines uniquely the normal vector to this surface. And the flux of a vector $\mathbf{N}$ on this normal surface is given by its contraction with the gradient. To see this, we consider a level surface

$$\phi(x^\mu) = c$$

for some constant c . Now let $x^\mu(\lambda)$ be a parametrized curve on this level surface $\phi = c$. Clearly its tangent vector is given by

$$V^\mu = \frac{dx^\mu}{d\lambda}.$$

These are also the tangent vectors of this level surface. We can construct a vector $\nabla\phi$, whose component is given by

$$(\nabla\phi)^\mu = \eta^{\mu\nu}(\partial_\nu\phi) = \eta^{\mu\nu}\partial_\nu\phi.$$

We claim that $\nabla\phi$ and V are orthogonal. Indeed,

$$\eta_{\mu\nu}(\nabla\phi)^\mu V^\nu = \eta_{\mu\nu}(\eta^{\mu\lambda}\partial_\lambda\phi)V^\nu = \delta^\lambda_\nu\partial_\lambda\phi V^\nu = \partial_\nu\phi V^\nu.$$

But

$$\partial_\nu\phi V^\nu = \frac{\partial\phi}{\partial x^\nu} \frac{dx^\nu}{d\lambda} = \frac{d\phi}{d\lambda} = \frac{d}{dt}(c) = 0.$$

This completes the proof. Thus we can define a normal vector $\nabla\phi$ of the constant surface $\phi = c$. Usually we want to normalize it, so define

$$\mathbf{n} = \frac{\nabla\phi}{|\eta_{\mu\nu}(\nabla\phi)^\mu(\nabla\phi)^\nu|^{1/2}}.$$

Recall that when we study electrodynamic, to calculate the flux of a vector field $\mathbf{F}$ on a surface S with unit normal $\mathbf{n}$, we simply compute $\mathbf{F} \cdot \mathbf{n}$ and integrate over the surface. Thus the flux is given by

$$\Phi = \int_S \mathbf{F} \cdot \mathbf{n} dS.$$

But sometimes $\mathbf{F} \cdot \mathbf{n}$ alone will be sufficient to determine the flux, since all we have to do is integrate this scalar over a given surface. In 4-dimensional spacetime, we use the same notion. So suppose we have a 4-vector $\mathbf{N}$. How do we compute its flux in the x^α direction? The normal vector of the constant surface $x^\alpha = c$ is ∇x^α . It is already normalized, since

$$(\nabla x^\alpha)^\lambda = \eta^{\lambda\gamma}\partial_\gamma x^\alpha = \eta^{\lambda\gamma}\delta^\alpha_\gamma = \eta^{\lambda\alpha}.$$

Thus

$$|\eta_{\mu\nu}(\nabla x^\alpha)^\mu(\nabla x^\alpha)^\nu| = |\eta_{\mu\nu}\eta^{\mu\alpha}\eta^{\nu\alpha}| = |\eta^{\alpha\alpha}| = 1.$$

So $\mathbf{n} = (\nabla x^\alpha)$. And we have

$$\mathbf{N} \cdot (\nabla x^\alpha) = \eta_{\mu\nu}(\nabla x^\alpha)^\mu N^\nu = \eta_{\mu\nu}\eta^{\mu\alpha}N^\nu = \delta^\alpha_\nu N^\nu = N^\alpha.$$

So the flux of $\mathbf{N}$ in x^α direction is just N^α . Also notice that

$$dx^\alpha(\mathbf{N}) = \partial_\mu x^\alpha N^\mu = \delta^\alpha_\mu N^\mu = N^\alpha.$$

So the flux of $\mathbf{N}$ in x^α direction can be obtained by $dx^\alpha(\mathbf{N})$.

2.7 Relativistic Dynamics

From now on we set $c = 1$. Let us start with the world line of a single particle, which we usually think of the path being a parametrized curve $x^\mu(\lambda)$. Let us take proper time τ as our paramter: $\lambda = \tau$. And define the **4-velocity**:

$$U^\mu = \frac{dx^\mu}{d\tau}.$$

The four velocity is normalized:

$$\eta_{\mu\nu}U^\mu U^\nu = -1.$$

To see this, we notice that the line element is given by

$$ds^2 = \eta_{\mu\nu}dx^\mu dx^\nu.$$

Therefore

$$\left(\frac{ds}{d\tau}\right)^2 = \eta_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} = \eta_{\mu\nu}U^\mu U^\nu = -1.$$

The last step is simply by definition $d\tau^2 = -ds^2/c^2 = -ds^2$ since $c = 1$. We can give a more explicit form of 4-velocity. Since by chain rule,

$$U^\mu = \frac{dx^\mu}{d\tau} = \frac{dx^\mu}{dt} \frac{dt}{d\tau} = \frac{dx^\mu}{dt} \left(\frac{d\tau}{dt}\right)^{-1} = \frac{dx^\mu}{dt} \frac{1}{\sqrt{1-v^2}} = \gamma \frac{dx^\mu}{dt}.$$

Therefore,

$$\mathbf{U} = (\gamma, \gamma V^x, \gamma V^y, \gamma V^z) = (\gamma, \gamma \mathbf{V}).$$

Newtons first law of motion holds in special relativistic mechanics as well as nonrelativistic mechanics. In the absence of forces, a body is at rest or moves in a straight line at constant speed. This is summarized by

$$\frac{d\mathbf{U}}{d\tau} = 0.$$

The objective of relativistic mechanics is to introduce the analog of Newtons second law $\mathbf{F} = m\mathbf{a}$. There is nothing from which this law can be derived, but plausibly it must satisfy certain properties: (1) It must satisfy the principle of relativity, i.e., take the same form in every inertial frame; (2) It must reduces to the equation above in the absence of forces; (3) it must reduce to $\mathbf{F} = m\mathbf{a}$ in any inertial frame when the speed of the particle is much less than the speed of light. The choice

$$m \frac{d\mathbf{U}}{d\tau} = \mathbf{F}$$

naturally suggests itself. Here m is the rest mass of the particle. This is the correct law of motion for special relativistic mechanics and the special relativistic generalization of Newtons second law. By introducing the four-acceleration four-vector

$$\mathbf{a} = \frac{d\mathbf{U}}{d\tau}$$

we recover $\mathbf{F} = m\mathbf{a}$. We can also define the 4-momentum:

$$p^\mu = mU^\mu$$

where m is the rest mass of the particle. One can see that

$$\eta_{\mu\nu}p^\mu p^\nu = -m^2.$$

Componentwise, we have that

$$\mathbf{p} = (\gamma m, \gamma m \mathbf{V}).$$

For small speeds $v := |\mathbf{V}| \ll 1$, we have that

$$p^t = p^0 = \frac{m}{1-v^2} \approx m + \frac{1}{2}mv^2 + O(v^2).$$

This reduces to the rest energy plus the momentum of the particle. Fixing the dimension by adding c , we see that the **rest energy** is given by

$$E_0 = mc^2.$$

So p^0 is nothing but the energy. We can rewrite $\eta_{\mu\nu}p^\mu p^\nu = -m^2$ and get the energy:

$$E = \sqrt{m^2 + p^2}.$$

2.8 Energy Momentum Tensor

Although p^μ provides a complete description of the energy and momentum of an individual particle, we often need to deal with extended systems with huge numbers of particles. We instead describe this system as a fluid - a continuum characterized by macroscopic quantities such as density, pressure, viscosity, and so on.



Figure 2.11: An example of dust.

We will first start off with a simple example of **dust**. Dust may be defined as a collection of particles that have mass and energy, and at rest with respect to each other. In a dust cloud, different dust element may have different rest frames (the frames with respect to which it is at rest), but for each individual element we can define a rest frame. This is where we start to describe the dust. How can we characterize dust? First we may want to know how many dust particles are there per unit volume. We call this quantity the **number density**. Let us define n_0 to be the number density in the chosen rest frame (which has dimension $1/m^3$). But we also want to know how to characterize

the number density *outside of the chosen rest frame, i.e. in some other inertial frame*. Let us consider a frame moving at speed v with respect to the rest frame. The number of dust elements stays the same, but the volume is contracted. So suppose the volume were originally $\Delta x \Delta y \Delta z$, then contraction will reduce this volume to $\Delta x \Delta y \Delta z \sqrt{1 - v^2}$. So we conclude that

$$n = \gamma n_0.$$

Also, there will be a flow of the dust through space. We can define a flux that describes the number of particles crossing unit area per unit time. This is a *three-vector*, which can only be the number density times the 3-velocity:

$$\mathbf{n} = n \mathbf{v} = \gamma n_0 \mathbf{v}.$$

We can combine these two pieces together into a 4-vector $\mathbf{N}$, by

$$\mathbf{N} = (n, n \mathbf{v}) = n_0 (\gamma, \gamma \mathbf{v}) = n_0 \mathbf{U}.$$

Here $\mathbf{U}$ is the overall 4-velocity field of the dust (but this is also the constant 4-velocity of individual particles since they are at rest with respect to each other). This is really cool, since now we have a geometrical object. Every little element in the dust have a trajectory through spacetime, and we can attach a 4-velocity to it. And if we know the number density of the dust in some rest frame, we can describe the dust in other frames. By our previous discussion on gradient and flux, we see that the flux of $\mathbf{N}$ in the $x^0 = t$ direction is given by

$$(dt)_\beta N^\beta = \delta^0_\beta N^\beta = N^0 = n.$$

So the flux of this vector through time is actually density. Now, the flux out of a certain volume must come at the expense of the density of the dust inside the volume. Mathematically, this is the conservation law given by

$$\frac{\partial n}{\partial t} = -\nabla \cdot \mathbf{n}$$

where on the right hand side we are just doing dot product of the 3-vectors. Writing this in a different way, we see that

$$\partial_\alpha N^\alpha = 0.$$

Or we can integrate the equation of conservation and use divergence theorem to get the integral form of conservation law:

$$\frac{\partial}{\partial t} \int_V n dV = - \int_{\partial V} \mathbf{n} \cdot d\mathbf{S}.$$

One important place where conservation law holds is in electromagnetism. Recall that we can define a 4-current

$$\mathbf{J} = (\rho, \mathbf{j})$$

where ρ is the charge density and $\mathbf{j}$ is the current density. Then the conservation of charge reads

$$\frac{\partial \rho}{\partial t} = -\nabla \cdot \mathbf{j} \quad \text{or} \quad \partial_\alpha J^\alpha = 0.$$

Now, electric and magnetic field are difficult to describe using geometric object perfect for spacetime. So if we want to describe electric and magnetic field using a geometric object, the first thing we think of is a 4-vector. But a 4-vector has only 4 components whereas $\mathbf{E}$ and $\mathbf{B}$ have a total of 6 component. So a 4-vector is not enough. It does not make too much sense either to make two 4-vectors and just ignore 2 redundant component. What about a 2-tensor? It has 16 components. That's too much. What if we make the tensor symmetric? Then there are only 10 independent component. But that is still too much. What about an antisymmetric tensor? An antisymmetric tensor must have zeros in the diagonal entries since $T^{\alpha\beta} = -T^{\beta\alpha}$. Thus the independent component are $10 - 4 = 6$, which is exactly the number we need. Therefore, we can define an antisymmetric tensor $F^{\mu\nu}$ that captures all the information we need to know about the $\mathbf{E}$ and $\mathbf{B}$ fields:

$$F^{\mu\nu} = \begin{pmatrix} 0 & E_x & E_y & E_z \\ -E_x & 0 & B_z & -B_y \\ -E_y & -B_z & 0 & B_x \\ -E_z & B_y & -B_x & 0 \end{pmatrix}.$$

We can also lower the indices by

$$F_{\mu\nu} = \eta_{\mu\alpha}\eta_{\nu\beta}F^{\alpha\beta}.$$

Then Maxwell's equations become

$$\begin{aligned}\partial_\nu F^{\mu\nu} &= 4\pi J^\mu \\ \partial_{[\mu} F_{\nu\lambda]} &= 0.\end{aligned}$$

Notice that Maxwell's equation already implies that the current has zero divergence:

$$\begin{aligned}4\pi\partial_\mu J^\mu &= \partial_\mu\partial_\nu F^{\mu\nu} \\ &= \partial_\nu\partial_\mu F^{\mu\nu} \\ &= -\partial_\nu\partial_\mu F^{\mu\nu} = 0 \\ &= -\partial_\mu\partial_\nu F^{\mu\nu} = 0.\end{aligned}$$

This is actually a trick we will use from time to time: If S is a symmetric tensor and A is an antisymmetric tensor, then

$$S_{\alpha\beta}A^{\alpha\beta} = 0.$$

Now, let us go back to the dust example. We can also consider the energy and momentum of dust. Suppose each dust particle has a rest mass m . Then in the chosen rest frame of a dust element we have its rest energy density:

$$\rho_0 = mn_0.$$

In a frame moving with speed v relative to the rest frame, we have the energy density in this frame being

$$\rho = (\gamma m)(\gamma n_0) = \gamma^2 mn_0 = \gamma^2 \rho_0.$$

If this ρ were a component of a 4-vector, there is no way by a Lorentz transformation (which will only pick up a factor of γ) we pick up a factor of γ^2 . This is not a component of a 4-vector. Of course this is not a scalar, which is invariant under Lorentz transformation. Conclusion: ρ is the component of a tensor. What is this tensor? Well, we assembled ρ by combining energy, the timelike component of $\mathbf{p}$, and the number density, the timelike component of $\mathbf{N}$. Therefore, $\rho = p^t N^t$ suggests that the tensor T such that $T^{00} = \rho$ is given by

$$T = \mathbf{N} \otimes \mathbf{p} = n_0 m \mathbf{U} \otimes \mathbf{U} = \rho_0 \mathbf{U} \otimes \mathbf{U}.$$

Or in component form:

$$T^{\mu\nu} = \rho_0 U^\mu U^\nu.$$

Actually there is a physical definition of this tensor T :

$$T^{\alpha\beta} = \text{flux of 4-momentum } \mathbf{p} \text{ in the } x^\beta \text{ direction.}$$

This is the **energy-momentum tensor**. So for example: $T^{00} = \rho_0 U^0 U^0 = \gamma^2 \rho_0 = \rho$. So this is the flux of momentum in the t direction. The flux of energy in time direction is energy density. This should be no surprise, since the flux of the vector $\mathbf{N}$ in time is given by

$$(dt)_\beta N^\beta = N^0 = n$$

which is the number density. We describe the flux of a thing in time as just being its density.¹ What about the other components? T^{0i} is the flux of p^t in x^i direction, so this is energy flux. T^{i0} is the flux of p^i in the t direction. So this is momentum density, as we discussed. T^{ij} is the flux of p^i in the x^j direction, which is just momentum flux. If we write out these components, we get

$$T^{0i} = \gamma^2 \rho_0 v^i \quad T^{i0} = \gamma^2 \rho_0 v^i \quad T^{ij} = \gamma^2 \rho_0 v^i v^j.$$

Thus the energy-momentum tensor for dust is symmetric. Actually, in all physical cases, the energy-momentum tensors are always symmetric.[?]

Another example is the **perfect fluid**. A perfect fluid is such that there exists a rest frame for some fluid element that no energy flow is present in the rest frame ($T^{0i} = 0$), and no lateral stresses ($T^{ij} = 0$ when $i \neq j$). Moreover, the fluid is characterized completely by the energy density ρ of the fluid, and its pressure p , which is an isotropic spatial stress (same in all direction). So in this frame,

$$T^{\mu\nu} = \begin{pmatrix} \rho & 0 & 0 & 0 \\ 0 & p & 0 & 0 \\ 0 & 0 & p & 0 \\ 0 & 0 & 0 & p \end{pmatrix} = \text{diag}(\rho, p, p, p).$$

We would like, of course, a formula that is good in any frame. For dust we have $T^{\mu\nu} = \rho_0 U^\mu U^\nu$, so we might begin by guessing $(\rho + p)U^\mu U^\nu$. Then in

¹This is not as trivial as it sounds. We define the **flux** to of something across a surface to be the amount of something crossing a unit area of that surface in a unit time. And things are more subtle when we discuss time flux. For detailed explanation, see [?] section 4.2-4.5, where the time flux being the energy density is explained in section 4.3.

this rest frame, we get

$$\begin{pmatrix} \rho + p & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

This is not quite right, but not too bad. We see that all we need is to add

$$\begin{pmatrix} -p & 0 & 0 & 0 \\ 0 & p & 0 & 0 \\ 0 & 0 & p & 0 \\ 0 & 0 & 0 & p \end{pmatrix}$$

to get the right answer. Fortunately, this has an obvious covariant generalization, namely $p\eta^{\mu\nu}$. Thus, the general form of the energy-momentum tensor for a perfect fluid is

$$T^{\mu\nu} = (\rho + p)U^\mu U^\nu + p\eta^{\mu\nu}.$$

We can have complete confidence in the result. Since this result reduces to $\text{diag}(\rho, p, p, p)$ in the rest frame, and this is a tensorial expression, we know that $T^{\mu\nu}$ above must be the right expression in any frame. This is the good thing about tensor! If it is true in one frame, it is true in all frame.

Therefore, you can now recall the Newtonian equation of gravitation:

$$\nabla^2 \phi = 4\pi G \rho.$$

You can immediately see where the problem is: since ρ is just component of a tensor, this is not a healthy theory. We have to make this equation tensorial. So we say, since the component of the energy-momentum tensor plays a huge role in gravity, we must replace the RHS of this equation by some constant times $T^{\mu\nu}$. And on the LHS we have some sort of second order derivative, which also needs a tensorial generalization. So our Einstein field equation should be something like

$$[\nabla^2(\text{a potential like object})]_{\mu\nu} \propto T_{\mu\nu}.$$

And that potential like object turned out to be the metric of spacetime. We will spend the next chapter developing the tools for generalizing the LHS.

One of the most important property of the energy-momentum tensor is that it is conserved:

$$\partial_\mu T^{\mu\nu} = 0.$$

This is a set of 4 equations, one for each value of ν . When $\nu = 0$, the equation is just conservation of energy, and when $\nu = k$ for $k = 1, 2, 3$, the equation is just conservation of momentum in the k th direction.

Chapter 3

Geometry of Spacetime

Let us review what we have done so far. Einstein developed the theory of special relativity based on two postulates:

1. The laws of physics are the same in all inertial frames.
2. The speed of light is the same in all inertial frames.

Because of postulate 2, we are forced to conclude that our spacetime has a very different geometrical structure than that we used to think. Namely, a line element takes the form $ds^2 = -dt^2 + dx^2 + dy^2 + dz^2$, instead of the intuitive Euclidean line element $ds^2 = dx^2 + dy^2 + dz^2$. This our spacetime with this geometrical structure is called the *Minkowski spacetime*. Now, according to postulate 1, all physical laws should be take the same form in all inertial frames, meaning they should be invariant under Lorentz transformations. Therefore, it is natural to review all the physical laws that had been previously developed and see if they fit into the relativistic framework. If yes, then great; if not, we need to modify them.

The first physical theory that Einstein looked into was Maxwell's theory for electricity and magnetism. They are described by four beautiful equations: Maxwell's equations. We have seen that Maxwell's equations are not invariant under Galilean transformation. But luckily, they are invariant under Lorentz transformations (For details see [this link](#)). Therefore, Maxwell's theory of electricity and magnetism fits perfectly into Einstein's relativity framework.

The next thing Einstein looked into was Newtonian theory of motion, which is governed by the law $\mathbf{F} = m\mathbf{a}$. Although this theory is not Lorentz invariant, but as we seen in the previous chapter, we can define the 4-velocity $U^\mu = dx^\mu/d\tau$ and the acceleration to be $\mathbf{a} = d\mathbf{U}/d\tau$. Then we can generalize Newton's second law into $m d\mathbf{U}/d\tau = \mathbf{F}$, where $\mathbf{F}$ is the 4-force. Therefore, We have successfully generalized Newtonian dynamics, which fits into relativistic framework. This is nice. And it shows that *special relativity describes accelerated motions just fine*.

But there is one thing missing: gravity. As we have seen before, Newton's theory of gravity implies Poisson's equation

$$\nabla^2 \phi = 4\pi G\rho.$$

And we have discussed in the last chapter that this theory does not work well in relativistic framework. Therefore, a remedy of Newton's theory of gravity

is needed. You may think this is not very hard, since after all Newton's formulation of gravity looks almost the same as the Columb force:

$$\mathbf{F}_{\text{grav}} = G \frac{m_1 m_2}{r^2} \mathbf{e}_r \quad \mathbf{F}_{\text{columb}} = k \frac{q_1 q_2}{r^2} \mathbf{e}_r.$$

And you may think since Maxwell's theory of electromagnetism fits perfectly into special relativity, we may develop a theory of gravity that is similar to the theory raised by Maxwell, and since they are of almost identical nature, we can fit gravity into the picture with not much effort. Physicists tried that, but failed, unfortunately. The reason is that the columb force depends on charges of particles (positive or negative), while the mass of a particle is always positive. And this seemingly tiny subtlety destroyed this attempt.

At last, it was Einstein who saved the day, once again. Einstein suggested that if we could not fit gravity into the framework of special relativity, we need to establish a new theory of gravity, which again revolutionized our understanding of spacetime. Einstein proposed that our spacetime is actually a 4-manifold instead of Minkowski spacetime. And the motivation behind this revolutionary proposition was the Principle of Equivalence.

3.1 Principle of Equivalence

The principle of equivalence has different forms. The first one is the **Weak Equivalence Principle** (WEP). Then Einstein generalized this to **Einstein's Equivalence Principle** (EEP). We first start with WEP and then move to EEP.

To understand WEP, we must first understand the difference between inertial mass and gravitational mass. So say we have two balls in deep space, where the gravity is absent. How do we determine which one is heavier? We can apply a force $F = 1N$ to both balls. Then the ball accelerates, and we can measure the magnitude of the acceleration. Then by using Newton's second law, we get that the mass of individual balls is equal to the ratio of force over acceleration. This quantity that we calculated, the so called "mass", is the **inertial mass**. It characterizes how a body will react to a given force, i.e. how hard it is to make this body move. We denote this mass by m_I . If we are on earth, then there is a gravitational field, which is approximately uniform, $g = 9.8m/s^2$, pointing towards the center of the earth. Then we have a natural intuition of how heavy the ball is: namely how much gravitational force the ball experiences. If we can measure this gravitational force, and we divide by g , we get the so called **gravitational mass**, which we denote by m_G . It characterizes how much gravitational force a certain object experiences. Now since the gravitational force of an object with gravitational mass m_G is given by $\mathbf{F} = m_G \mathbf{g}$, and by Newton's second law, we have $\mathbf{F} = m_I \mathbf{a}$, then equating these two expressions (assuming the object only experience the force of gravity) we have

$$\mathbf{a} = \left(\frac{m_G}{m_I} \right) \mathbf{g}.$$

One often compare gravity with the force of a particle with charge q experiences in an electric field $\mathbf{E}$, so we might as well compare these two expressions. The

acceleration that the charged particle experiences is

$$\mathbf{a} = \left(\frac{q}{m} \right) \mathbf{E}$$

where m is the inertial mass of the particle. So somehow you can think of m_G as the gravitational charge, and the ratio m_G/m_I as the counterpart of q/m in electromagnetic theory. Galileo long ago showed that (apocryphally by dropping weights off the Leaning Tower of Pisa) that the response of matter to gravity is universal - every object falls at the same rate in the gravitational field, independent of the composition of the object. In other words, two objects starting with the same initial state fall at the same rate. That is, m_G/m_I is a constant. By adjusting the value of the gravitational constant G in $\mathbf{F} = G \frac{m_1 m_2}{r^2} \mathbf{e}_r$ we may assume $m_G/m_I = 1$, or

$$m_G = m_I.$$

But one should understand that m_G and m_I are of very different nature: One is due to the effect of gravity, the other is observed due to acceleration. Einstein suggested that these two masses are actually the same thing: They have different causes, but they have the same effect. This is best characterized by the following thought experiment.

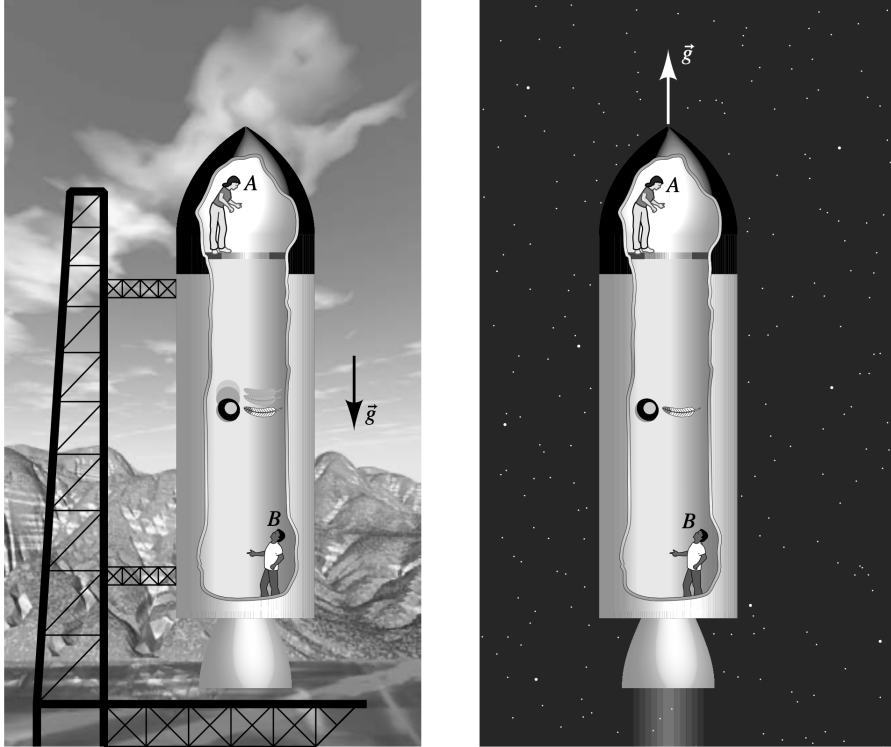


Figure 3.1: An observer inside a closed laboratory cannot distinguish whether they are in one situation or the other.

Consider an experimenter in a small, closed laboratory at rest on the surface of the Moon or other source of gravitational force as illustrated in the figure above. The laboratory is small enough that the gravitational field is

uniform to any accuracy the experimenter can test. The experimenter can carry out experiments with various objects. For example, if a cannonball and a feather are dropped, they will fall to the floor of the laboratory with the same acceleration the local acceleration of gravity g because of the equality of gravitational and inertial mass. Consider the same laboratory in empty space, far from any source of gravitational force, not at rest, but accelerated upward with an acceleration g , as also illustrated in the figure. An experimenter inside who drops a cannonball and a feather will observe that they fall to the floor of the laboratory with equal acceleration the same result as for the laboratory at rest in a gravitational field. By this, or any other mechanical experiment with particles, the experimenter inside cannot tell whether the laboratory is unaccelerated in a uniform gravitational field or accelerated in empty space. The two laboratories are equivalent as far as these experiments are concerned.

The power of the equivalence principle derives from its assertion that it applies to all laws of physics. As an example, if we accept the equivalence principle, we must also accept that light falls in a gravitational field with the same acceleration as material bodies. It is not obvious otherwise how to calculate the effect of gravity on light. There is no Newtonian equation of motion: no $F = m_I a$. Here is how the equivalence principle forces one to the conclusion that light falls in a gravitational field.

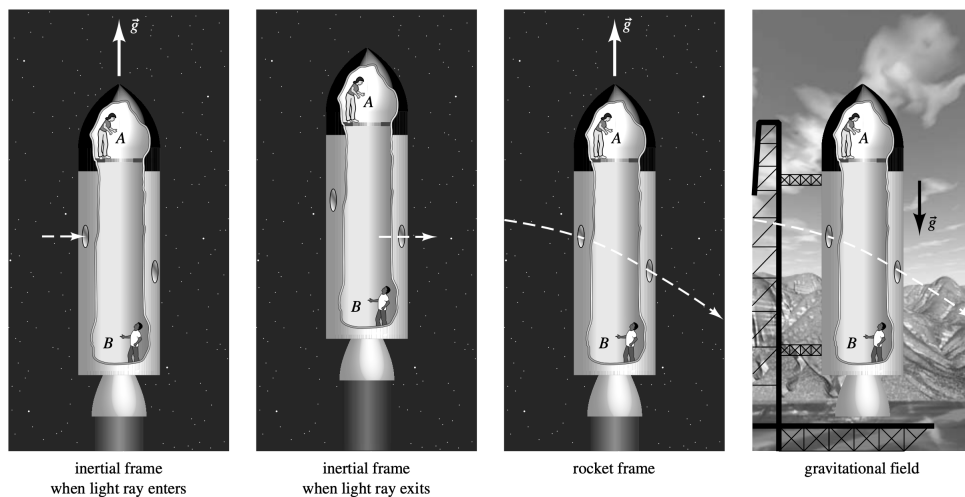


Figure 3.2: A light ray in a gravitational field must fall with the same acceleration as other objects. Gravity bends light!

In empty space, a light ray will move on a straight line in an inertial frame. Suppose a light ray is observed from a laboratory accelerating transversely to its direction of propagation (Fig.3.2). In the laboratory frame, the light ray will exit at a position below where it entered because the laboratory has accelerated upward in the time the light ray crosses. Thus, in the laboratory frame the light ray will accelerate downward with the acceleration of the laboratory. From the equivalence principle, one can deduce that the same behavior occurs in a uniform gravitational field; i.e., the light ray accelerates downward with the local acceleration of gravity.

To be careful, we must limit our claims about the impossibility of distinguishing gravity from uniform acceleration by restricting our attention to

small enough regions in *spacetime* (space and time). If the lab is big enough, the gravitational field would change from place to place in an observable way, while the effect of acceleration would always fall straight down. But when the spacetime volume is small, then uniform acceleration over a small enough time interval barely affects the moving speed of a reference frame (in this case the lab). Therefore, this lab behaves almost like an inertial frame if the lab occupies a small enough spatial volume and if we consider a small enough interval of time. Therefore, the WEP can be stated as follows:

Weak Equivalence Principle (WEP). Experiments in a sufficiently small freely falling laboratory, over a sufficiently short time, give results that are indistinguishable from those of the same experiments in an inertial frame in empty space.

Inertial frame is something we deal with a lot in SR. And we know what physical laws would be like in inertial frames. So Einstein generalizes to WEP to the following:

Einstein Equivalence Principle (EEP). In small enough regions of spacetime, the laws of physics reduce to those of special relativity; it is impossible to detect the existence of a gravitational field by means of local experiments.

Guess what? We can actually do better than this. We can show that there always exists a **local Lorentz frame**. That means given a set of coordinate system x^α , with metric $g_{\alpha\beta}$ (the component can vary from place to place, need not be constant). Then there exists a set of coordinates $x^{\alpha'}$, and a coordinate transformation $L : \{x^\alpha\} \rightarrow \{x^{\alpha'}\}$, such that the spacetime is Lorentz in the vicinity of some event p . By Lorentz we mean the metric in the coordinate $x^{\alpha'}$ has a representation close to $\eta_{\alpha\beta}$. So suppose that our old coordinate can be written as functions of the new coordinate:

$$x^\alpha = x^\alpha(x^{\alpha'}).$$

And this gives rise to the coordinate transformation matrix, which is just the Jacobian:

$$L^\alpha_{\mu'} = \frac{\partial x^\alpha}{\partial x^{\mu'}}.$$

The goal is to show that we can find a coordinate system such that

$$g_{\mu'\nu'} = L^\alpha_{\mu'} L^\beta_{\nu'} g_{\alpha\beta} = \eta_{\mu'\nu'} \quad (3.1)$$

over a small region containing p . We can expand the second expression of the equation above in a Taylor series centered at p . Since we are given the metric g , and we are free to choose coordinate transformation and their derivatives to be whatever we want, we can compare the degrees of freedom offered by the coordinate transformation (which we select) to the constraints imposed by the metric and its derivatives. If we expand the metric:

$$g_{\alpha\beta} = g_{\alpha\beta}(p) + (x^{\gamma'} - x_p^{\gamma'}) \partial_{\gamma'} g_{\alpha\beta}(p) + \frac{1}{2} (x^{\gamma'} - x_p^{\gamma'}) (x^{\sigma'} - x_p^{\sigma'}) \partial_{\gamma'} \partial_{\sigma'} g_{\alpha\beta}(p) + \dots$$

And we also have

$$L^\alpha_{\mu'} = L^\alpha_{\mu'}(p) + (x^{\gamma'} - x_p^{\gamma'}) \partial_{\gamma'} L^\alpha_{\mu'}(p) + \frac{1}{2} (x^{\gamma'} - x_p^{\gamma'}) (x^{\sigma'} - x_p^{\sigma'}) \partial_{\gamma'} \partial_{\sigma'} L^\alpha_{\mu'}(p) + \dots$$

So now $g_{\alpha\beta}(p), \partial g_{\alpha\beta}(p), \partial^2 g_{\alpha\beta}(p)$ are constraints, which we have no control over. And $L^\alpha_{\mu'}, \partial L^\alpha_{\mu'}, \partial^2 L^\alpha_{\mu'}$ are specifiable, so they are our degrees of freedom. Now let us start counting! Plug the expansions into (3.1), we have that (very schematically):

$$L^\alpha_{\mu'} L^\beta_{\nu'} g_{\alpha\beta} = L^\alpha_{\mu'}(p) L^\beta_{\nu'}(p) g_{\alpha\beta}(p) + (x^{\gamma'} - x_p^{\gamma'}) A(\partial g_{\alpha\beta}(p), \partial L^\alpha_{\mu'}) \quad (3.2)$$

$$+ (x^{\gamma'} - x_p^{\gamma'})(x^{\sigma'} - x_p^{\sigma'}) B(\partial^2 g_{\alpha\beta}(p), \partial^2 L^\alpha_{\mu'}). \quad (3.3)$$

Here, $A(\partial g_{\alpha\beta}(p), \partial L^\alpha_{\mu'})$ is some function involving the first derivatives of the metric and L , and similarly for B . We want to make this expression look like $\eta_{\mu'\nu'}$ as possible. Let us look at this at the 0th order: My constraints are $g_{\alpha\beta}(p)$. Since g is a metric, it is a symmetric 4×4 matrix,¹ thus we have 10 constraints to satisfy. And our degrees of freedom are:

$$L^\alpha_{\mu'}(p) = \frac{\partial x^\alpha}{\partial x^{\mu'}}.$$

This is NOT symmetric. So we have 16 degrees of freedom to satisfy the 10 constraints. We can easily make the 0th order equal to $\eta_{\mu'\nu'}$, and we still have $16 - 10 = 6$ degrees of freedom. What about the 1st order? My constraints are $\partial_{\gamma'} g_{\alpha\beta}(p)$, where $\gamma = 0, 1, 2, 3$. So this gives me $4 \times 10 = 40$ constraints to satisfy. The things we are allowed to choose are

$$\partial_{\gamma'} L^\alpha_{\mu'}(p) = \frac{\partial^2 x^\alpha}{\partial x^{\gamma'} \partial x^{\mu'}}(p).$$

This is symmetric on μ' and γ' , and $\alpha = 0, 1, 2, 3$. So we have a total of $4 \times 10 = 40$ degrees of freedom. This is a perfect match. So we can make the first order equal to zero, since I have as many constraints as my degrees of freedom! However, we are not so luck when we come to the 2nd order. The constraints are:

$$\partial_{\gamma'} \partial_{\sigma'} g_{\alpha\beta}(p).$$

This is symmetric on γ', σ' , and symmetric on α, β . So we have a total of $10 \times 10 = 100$ constraints to satisfy. And our degrees of freedom are:

$$\partial_{\gamma'} \partial_{\sigma'} L^\alpha_{\mu'}(p) = \frac{\partial^3 x^\alpha}{\partial x^{\gamma'} \partial x^{\sigma'} \partial x^{\mu'}}(p).$$

The number of degrees of freedom turns out to be (combinatorics...)

$$4 \times \frac{n(n+1)(n+2)}{3!} \Big|_{n=4} = 80.$$

Thus we cannot make the second order derivative of (3.2) vanish. However, we can make (3.2) equal to $\eta_{\mu'\nu'}$ in the 0th order (which means it is exactly $\eta_{\mu'\nu'}$

¹A metric is always symmetric. For

$$g_{\alpha\beta} dx^\alpha dx^\beta = \frac{1}{2} (g_{\alpha\beta} dx^\alpha dx^\beta + g_{\beta\alpha} dx^\beta dx^\alpha) = \frac{1}{2} (g_{\alpha\beta} + g_{\beta\alpha}) dx^\alpha dx^\beta.$$

So so only the symmetric part of $g_{\alpha\beta}$ would survive the sum. As such we may as well take $g_{\alpha\beta}$ to be symmetric in its definition.

at p), and make its first derivative vanish. But this is already good enough, since this calculation implies that we can write

$$g_{\mu'\nu'} = \eta_{\mu'\nu'} + O((\Delta x)^2).$$

The term $O((\Delta x)^2)$ vanishes quickly when $\Delta x \rightarrow 0$. Therefore, just as the EEP stated, we can find a local Lorentz frame in a sufficiently small region around every event p in the spacetime. And the geometric in this small region is effectively Minkowskian!

Now, when a gravitational field is present, in a sufficiently small region of spacetime, we can construct locally inertial (Lorentz) frames. Thus to generalize the theory of gravity, the solution is to retain the notion of inertial frames, but to discard the hope that they can be uniquely extended throughout space and time. Instead, we can define locally inertial frames at every point. We therefore need to invoke a mathematical framework in which physical theories can be consistent with these conclusions. The solution will be to imagine the spacetime has a curved geometry, and that gravitation is a manifestation of this curvature.² The appropriate mathematical structure for this is that of a differentiable manifold.

3.2 Smooth Manifolds

We assume the reader is familiar with basic point set topology. Here we quickly review some definitions from topology. A **topology** on a set $\mathcal{M}$ is a collection of subsets of $\mathcal{M}$ that includes $\emptyset, \mathcal{M}$, and are closed under arbitrary unions and finite intersections. These collection of subsets are called **open sets** of $\mathcal{M}$. The set $\mathcal{M}$ equipped with a topology is called a **topological space**. A map

$$f : \mathcal{M} \rightarrow \mathcal{N}$$

between two topological spaces is called **continuous**, if for every open set $U \subseteq \mathcal{N}$, the preimage $f^{-1}(U) \subseteq \mathcal{M}$ of U under f is open in $\mathcal{M}$. If f is bijective and the inverse f^{-1} is also continuous, then f is called a **homeomorphism**. The space $\mathcal{M}, \mathcal{N}$ are then called **homeomorphic**. Two space being homeomorphic means that they are equivalent from the topological pointview. This is just like when we say two sets are equivalent in the sense that there exists a bijective between them, or two vector space being equivalent in the sense that they are linearly isomorphic.

Suppose $\mathcal{M}$ is a Hausdorff, second countable space (that means for every pair of distinct points there are disjoint open sets containing them, and there exists a countable basis for the topology of $\mathcal{M}$). Then we call $\mathcal{M}$ a **topological manifold** of dimension n if for every point $p \in \mathcal{M}$ we can find

- An open subset $U \subseteq \mathcal{M}$ containing p ;
- An open subset $V \subset \mathbf{R}^n$;

²Though this is very surprising, we can make some sense into it. Before we discussed that the response of matter to gravity is independent of the composition of the object, and they have the same acceleration provided they have the same initial state (e.g. where they are released). This suggests that the effect of gravity does not depend on matters, but on spacetime itself. So we may guess that gravity is actually an intrinsic property of the spacetime, unlike other forces.

- A homeomorphism $\varphi : U \rightarrow V$.

The pair (U, φ) above is called a **chart** or coordinate system of $\mathcal{M}$. A collection of charts $\mathcal{A} = \{(U_\alpha, \varphi_\alpha)\}$ of $\mathcal{M}$ is called an **atlas** if the set of U_α covers M . When $U_\alpha \cap U_\beta \neq \emptyset$, i.e. when two charts intersect, we can define the chart transition map:

$$\tau_{\alpha,\beta} : \varphi_\alpha(U_\alpha \cap U_\beta) \rightarrow \varphi_\beta(U_\alpha \cap U_\beta) \quad \tau_{\alpha,\beta} = \varphi_\beta \circ \varphi_\alpha^{-1}.$$

Since $\tau_{\alpha,\beta}$ is a composition of two homeomorphisms, it is a homeomorphism itself.

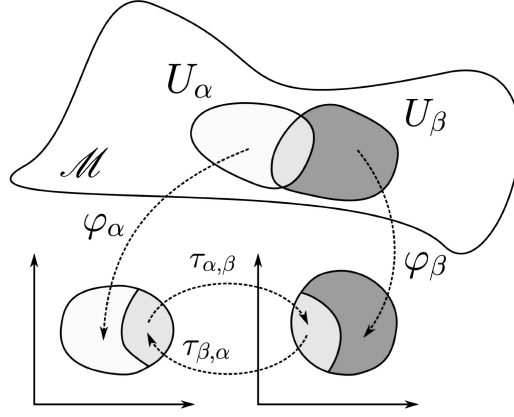


Figure 3.3: Chart transition maps.

A map $f : V \rightarrow W$, where $V \subseteq \mathbf{R}^m$ and $W \subseteq \mathbf{R}^n$ is called **smooth** if its component functions are smooth, i.e. infinitely differentiable. If f is bijective and the inverse f^{-1} is also smooth, then f is called a **diffeomorphism**. Then V, W are called **diffeomorphic**. Two charts $(U_\alpha, \varphi_\alpha)$ and (U_β, φ_β) are called **smoothly compatible** if either $U_\alpha \cap U_\beta = \emptyset$ or the transition map $\tau_{\alpha,\beta}$ is a diffeomorphism (which implies that $\tau_{\beta,\alpha} = (\tau_{\alpha,\beta})^{-1}$ is a diffeomorphism as well). If any pair of charts in the atlas $\mathcal{A}$ of $\mathcal{M}$ is smoothly compatible, we say $\mathcal{M}$ is a **smooth manifold** or a C^∞ manifold.

Given a smooth manifold $\mathcal{M}$, we can build local coordinates on $\mathcal{M}$. Suppose we have a chart (U, φ) with $U \subseteq \mathcal{M}$. Then we claim that we can build coordinate U by labeling the points in U uniquely (see the discussion on reference frames in the last chapter). We have a homeomorphism $\varphi : U \rightarrow V$ where $V \subseteq \mathbf{R}^n$ is some open set in $\mathbf{R}^n$. Then, since homeomorphism φ is a bijective map, for every point $p \in U$ there exists a unique n -tuple

$$\varphi(p) = (x^\mu) \in V \subseteq \mathbf{R}^n \quad \mu = 1, 2, \dots, n.$$

For different values of p we have different values of (x^μ) . Therefore, for each μ , we can think of x^μ as a function $x^\mu : U \subseteq \mathcal{M} \rightarrow \mathbf{R}^n$ by the rule $x^\mu(p)$ is the μ th coordinate of the n -tuple $\varphi(p)$. Or alternatively, we can think of labeling the points p on $U \subset \mathcal{M}$ uniquely by the tuple (x^μ) . This gives the local coordinates of a manifold $\mathcal{M}$. Since U and V are homeomorphic, we identify them sometimes, and think about coordinate lines x^μ on manifolds instead.

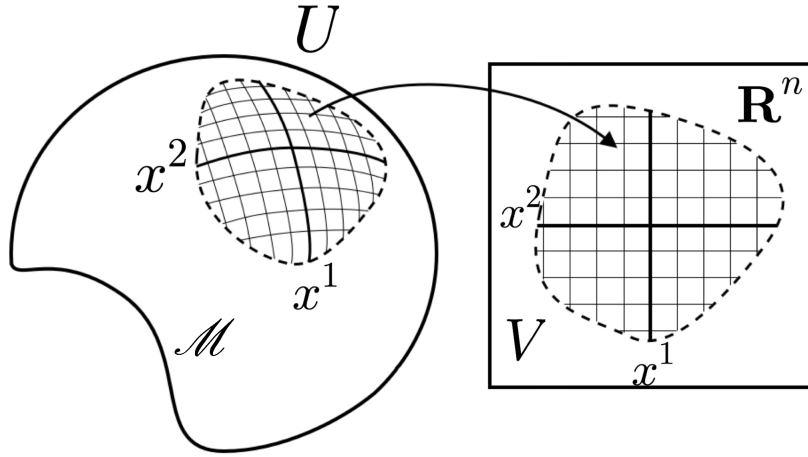


Figure 3.4: Local coordinate system

3.3 Tangent Vectors

Having constructed this groundwork, we can now proceed to introduce various kinds of structure on manifolds. We begin by vectors and tangent spaces. We mentioned these concepts in the last chapter. One point we stressed was the notion of a tangent space T_p at a point $p \in \mathcal{M}$, which is the set of all vectors at a single point on our (spacetime) manifold $\mathcal{M}$.

Let us think about how to construct the tangent space T_p at a point $p \in \mathcal{M}$, using only things that are intrinsic to $\mathcal{M}$. We want to define tangent vectors in an intrinsic way because in relativity often we are given a manifold $\mathcal{M}$ where the embedding of $\mathcal{M}$ into larger Euclidean space is not available. So there is no way we can get aid from anything other than the manifold itself. The first thing we may consider is the tangent vector to a curve $\gamma : \mathbf{R} \rightarrow \mathcal{M}$ passing through p at the point p . Then we may want to define T_p to be the collection of all tangent vectors to these curves at the point p . But this is obviously cheating, since we do not have a notion of what a "tangent vector" to a curve means. But we are on the right track, we just have to make this coordinate independent. Since each curve defines a directional derivative on the set of smooth functions³ $f : \mathcal{M} \rightarrow \mathbf{R}$, namely every curve $\gamma(\lambda)$ maps $f \rightarrow d(f \circ \gamma)/d\lambda(\lambda_p)$, we can sort of identify these operation of directional derivatives to a tangent vector. So let us think about what directional derivatives need to satisfy in general.

To this end, we define $\mathcal{F}$ to be the set of all smooth functions $f : \mathcal{M} \rightarrow \mathbf{R}$. We define a tangent vector v at point $p \in \mathcal{M}$ to be a map $v : \mathcal{F} \rightarrow \mathbf{R}$ such that

1. v is linear: For all $f, g \in \mathcal{F}$ and $a, b \in \mathbf{R}$ we have $v(af + bg) = av(f) + bv(g)$.
2. v obeys the Leibniz rule: $v(fg) = f(p)v(g) + g(p)v(f)$.

It is clear that the set of all tangent vectors T_p at point p form a vector space, since addition and scalar multiplication are well-defined: $(v_1 + v_2)(f) =$

³A function $f : \mathcal{M} \rightarrow \mathbf{R}$ is called smooth if for every $p \in \mathcal{M}$, there exists a smooth chart (U, φ) such that $p \in U$ and the composite map $f \circ \varphi^{-1} : V \rightarrow \mathbf{R}$ is smooth.

$v_1(f) + v_2(f)$ and $(av)(f) = av(f)$. Now we claim that

$$\dim T_p = \dim \mathcal{M} = n.$$

We will show this claim by finding a suitable basis for T_p . Let (O, ψ) be a coordinate chart such that $p \in O$. For $\mu = 1, 2, \dots, n$, define the **partial derivative operator** $\partial_\mu : \mathcal{F} \rightarrow \mathbf{R}$ by

$$\partial_\mu f = \frac{\partial}{\partial x^\mu} (f \circ \psi^{-1}) \Big|_{\psi(p)}.$$

Here, $(x^1, \dots, x^n)$ are the Cartesian coordinates of $\mathbf{R}^n$. Then it is easy to verify that

- ∂_μ are indeed tangent vectors (satisfying 1 and 2 above);
- ∂_μ are linearly independent.

We left this easy exercise to the reader. To show that they span T_p , we use the following result from calculus: If $F : \mathbf{R}^n \rightarrow \mathbf{R}$ is smooth (C^∞), then for each $a = (a^1, \dots, a^n) \in \mathbf{R}^n$ there exists C^∞ functions H_μ such that for all $x \in \mathbf{R}^n$ we have

$$F(x) = F(a) + (x^\mu - a^\mu)H_\mu(x) \quad H_\mu(a) = \frac{\partial F}{\partial x^\mu} \Big|_{x=a}.$$

This is just a generalized version of the chain rule $dF = \partial_\mu F dx^\mu$. Now, we apply the result here, letting $F = f \circ \psi^{-1}$ and $a = \psi(p)$. Then, for all $q \in O$, we have that

$$f(q) = f(p) + [x^\mu(\psi(q)) - x^\mu(\psi(p))]H_\mu(\psi(q)).$$

Let $v \in T_p$, we wish to show that v is a linear combination of ∂_μ s. To do this, apply v to f , using the equation above, the linearity, the Leibniz rule, and the fact that v applied to a constant such as $f(p)$ vanishes. We obtain

$$\begin{aligned} v(f) &= v[f(p)] + [x^\mu(\psi(q)) - x^\mu(\psi(p))]_{q=p} v(H_\mu \circ \psi) + (H_\mu \circ \psi(p)) v[x^\mu(\psi(q)) - x^\mu(\psi(p))] \\ &= [H_\mu \circ \psi(p)] v(x^\mu \circ \psi). \end{aligned}$$

But by our definition, we have that

$$H_\mu \circ \psi(p) = \frac{\partial F}{\partial x^\mu}(\psi(p)) \frac{\partial}{\partial x^\mu} (f \circ \psi^{-1}) \Big|_{\psi(p)} = \partial_\mu f.$$

Therefore, we have

$$v(f) = v^\mu \partial_\mu(f)$$

where

$$v^\mu = v(x^\mu \circ \psi).$$

Since this is true for any $f \in \mathcal{F}$, we have that

$$v = v^\mu \partial_\mu.$$

This proves the claim. Therefore, indeed we have $\dim T_p = \dim \mathcal{M} = n$ and the basis of T_p is given by the partial derivatives ∂_μ s defined on local coordinates.

Given a new basis $x^{\mu'}$, each of them is a function of the old coordinates, $x^{\mu'}(x^\mu)$, and the transformation rule of the basis tangent vectors is simply

$$\partial_{\mu'} = \frac{\partial x^\nu}{\partial x^{\mu'}} \partial_\nu.$$

Given a vector V , we can write it in terms of two different set of basis:

$$V^\mu \partial_\mu = V^{\mu'} \partial_{\mu'} = V^{\mu'} \frac{\partial x^\mu}{\partial x^{\mu'}} \partial_\mu$$

Thus

$$V^{\mu'} = \frac{\partial x^{\mu'}}{\partial x^\mu} V^\mu.$$

Notice that $\partial x^{\mu'}/\partial x^\mu$ is the inverse of the matrix $\partial x^\mu/\partial x^{\mu'}$. Observe that this rule is compatible with the Lorentz transformation rule of vectors that we derived earlier, but much more general. So how exactly are derivations related to the tangent vectors of curves? Say we have a curve $\gamma : \mathbf{R} \rightarrow \mathcal{M}$. And say $\lambda_p \in \mathbf{R}$ such that $\gamma(\lambda_p) = p \in \mathcal{M}$. Let (U, φ) be a chart containing p and let $f \in \mathcal{F}$. Then we can identify the tangent vector of γ at p with a tangent vector $T \in T_p$, such that

$$T(f) = \left. \frac{d}{d\lambda} (f \circ \gamma) \right|_{\lambda=\lambda_p}.$$

The map $\varphi \circ \gamma : \mathbf{R} \rightarrow \mathbf{R}^n$ is a curve in $\mathbf{R}^n$, with components $x^\mu(\lambda)$. Then by the chain rule (see the figure below), we have

$$T(f) = \left. \frac{d}{d\lambda} (f \circ \gamma) \right|_{\lambda=\lambda_p} = \left. \frac{dx^\mu}{d\lambda} \right|_{\lambda=\lambda_p} \cdot \left. \frac{\partial}{\partial x^\mu} (f \circ \varphi^{-1}) \right|_{\varphi(p)} = \frac{dx^\mu}{d\lambda} \partial_\mu f.$$

Therefore, the tangent vector T is a linear combination of the partial derivatives with coefficients $dx^\mu/d\lambda$. This is very easy to remember. So what are the partial derivatives ∂_μ anyway? Given an arbitrary curve, its tangent vector is of the form

$$T = \frac{dx^\nu}{d\lambda} \partial_\nu.$$

We see that

$$\partial_\mu = \delta^\nu_\mu \partial_\nu.$$

Therefore, if we have a curve, with the condition

$$\frac{dx^\nu}{d\lambda} = \delta^\nu_\mu,$$

then its tangent vector should be the vector ∂_μ . Clearly, a curve γ that satisfies the condition above is the curve whose image $\varphi \circ \gamma$ that is constant in all the components $x^\nu(\lambda)$ for all $\nu \neq \mu$, and $x^\mu(\lambda)$ is varying (actually, it is parametrized by λ itself). Recalling from our discussion on local coordinate in the last section, we can think of this curve on $\mathcal{M}$ as the one fixing all the grids x^ν for $\nu \neq \mu$. So this curve is a grid line in the μ -direction, as the figure below suggests.

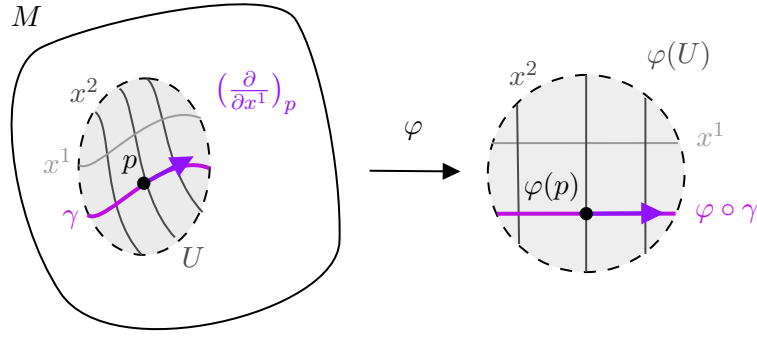


Figure 3.5: A curve γ on M having $(\partial/\partial x^1)_p$ as its tangent vector at p , with $\varphi : U \rightarrow \varphi(U) \subset \mathbf{R}^n$ a diffeomorphism between open sets.

Since a vector $v \in T_p$ is a map $\mathcal{F} \rightarrow \mathbf{R}$ by taking derivative, it should be clear that a **vector field** defines a map $\mathcal{F} \rightarrow \mathcal{F}$ by taking derivatives. Given two vector fields X, Y , we can define their **commutator** $[X, Y]$ by

$$[X, Y](f) = X(Y(f)) - Y(X(f)).$$

It is easy to check that the commutator is linear and obeys the Libniz rule. So the commutator of two vector fields gives a third vector field. The components can be written as

$$[X, Y]^\mu = X^\lambda \partial_\lambda Y^\mu - Y^\lambda \partial_\lambda X^\mu.$$

This is actually a well-defined tensor. The commutator is sometimes called the **Lie bracket**. Note that since partial derivatives commutes, the commutator of the vector fields given by the partial derivatives of the coordinate functions ∂_μ always vanishes.

3.4 Tensors

Now we follow the steps we took in the last chapter, and study dual vectors. Given a tangent space T_p , the **cotangent space** T_p^* , which is dual to T_p , is the collection of linear maps $T_p \rightarrow \mathbf{R}$. In the last chapter, we defined the gradient of a function f . But now our function f is a map $\mathcal{M} \rightarrow \mathbf{R}$. There is a natural way to let the gradient df act on a vector v and give a real number, so the definition is

$$df(v) = v(f).$$

It is easy to check that df is linear. Thus $df \in T_p^*$. We know that $\dim T_p^* = \dim T_p = n$, so we want to find a basis for T_p^* . Clearly the natural basis is the set of dual vectors $\omega^\mu(\partial_\nu) = \delta^\mu_\nu$. But notice that

$$dx^\mu(\partial_\nu) = \partial_\nu x^\mu = \frac{\partial x^\mu}{\partial x^\nu} = \delta^\mu_\nu.$$

Therefore the gradient dx^μ are an appropriate set of basis dual vectors (you should check spanning and linear independence). The transformation law is just what you would expect:

$$dx^{\mu'} = \frac{\partial x^{\mu'}}{\partial x^\nu} dx^\nu.$$

It follows immediately that if $\omega = \omega_\mu dx^\mu \in T_p^*$, then the transformation rule is given by

$$\omega_{\mu'} = \frac{\partial x^\nu}{\partial x^{\mu'}} \omega_\nu.$$

The definition of a tensor is exactly the same as that in a flat space: a (k, l) tensor is a linear map that takes k dual vectors and l vectors to $\mathbf{R}$. The component of a tensor in a coordinate basis can be obtained by

$$T^{\mu_1 \dots \mu_k}_{\nu_1 \dots \nu_l} = T(dx^{\mu_1}, \dots, dx^{\mu_k}, \partial_{\nu_1}, \dots, \partial_{\nu_l}).$$

This is equivalent to

$$T = T^{\mu_1 \dots \mu_k}_{\nu_1 \dots \nu_l} \partial_{\mu_1} \otimes \dots \otimes \partial_{\mu_k} \otimes dx^{\nu_1} \otimes \dots \otimes dx^{\nu_l}.$$

Again we have the transformation rule as expected:

$$T^{\mu'_1 \dots \mu'_k}_{\nu'_1 \dots \nu'_l} = \frac{\partial x^{\mu'_1}}{\partial x^{\mu_1}} \dots \frac{\partial x^{\mu'_k}}{\partial x^{\mu_k}} \frac{\partial x^{\nu_1}}{\partial x^{\nu'_1}} \dots \frac{\partial x^{\nu_l}}{\partial x^{\nu'_l}} T^{\mu_1 \dots \mu_k}_{\nu_1 \dots \nu_l}.$$

However, it is usually easier to transform a tensor by taking the identity of basis vectors and basis dual vectors as partial derivatives and gradients at face value, and simply substitute in the coordinate transformation. For example, let our coordinate system be $x^1 = x$ and $x^2 = y$ and the tensor be

$$S_{\mu\nu} = \begin{pmatrix} 1 & 0 \\ 0 & x^2 \end{pmatrix}.$$

So we write

$$S = S_{\mu\nu} (dx^\mu \otimes dx^\nu) = (dx)^2 + x^2 (dy)^2.$$

Now suppose our coordinate transformation is

$$x' = \frac{2x}{y} \quad y' = \frac{y}{2}.$$

Then we have

$$x = x'y' \quad y = 2y'.$$

Differentiating, we get

$$dx = y'dx' + x'dy' \quad dy = 2dy'.$$

Plugging in:

$$S = (y')^2 (dx')^2 + x'y' (dx'dy' + dy'dx') + [(x')^2 + 4(x'y')^2] (dy')^2.$$

Therefore,

$$S_{\mu'\nu'} = \begin{pmatrix} (y')^2 & x'y' \\ x'y' & (x')^2 + 4(x'y')^2 \end{pmatrix}.$$

For the most part, the various tensor operations we defined in flat space are unaltered in a more general setting: contraction, symmetrization, and so on. There are three important exceptions: partial derivatives, the metric, and the Levi-Civita tensor. Let us look at the partial derivative first.

In general, the partial derivative is not a tensor. We can see that by considering the partial derivative of a dual vector, $\partial_\mu W_\nu$. If we change to a new coordinate we have

$$\begin{aligned}\frac{\partial}{\partial x^{\mu'}} W_{\nu'} &= \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial}{\partial x^\mu} \left(\frac{\partial x^\nu}{\partial x^{\nu'}} W_\nu \right) \\ &= \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^\nu}{\partial x^{\nu'}} \left(\frac{\partial}{\partial x^\mu} W_\nu \right) + W_\nu \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial}{\partial x^\mu} \frac{\partial x^\nu}{\partial x^{\nu'}}.\end{aligned}$$

The second term in the last line should not be there if $\partial_\mu W_\nu$ were to transform as a (0,2) tensor. This is because the derivative of the transformation matrix does not vanish, as it did for Lorentz transformations in flat space. Since differentiation is obviously important in physics, we will have to invent new tensorial operation to make place of the partial derivative. In fact we will invent several: the exterior derivative, and covariant derivative, and the Lie derivative.

One of the most important tensors is the **metric tensor**. A metric tensor $g_{\mu\nu}$ is a symmetric, nondegenerate (0,2) tensor. Nondegenerate means the determinant $g = |g_{\mu\nu}|$ does not vanish. This allows us to define the inverse metric $g^{\mu\nu}$ via

$$g^{\mu\nu} g_{\nu\sigma} = g_{\lambda\sigma} g^{\lambda\mu} = \delta^\mu_\sigma.$$

In our discussion of path lengths in special relativity, we introduced the line element

$$ds^2 = \eta_{\mu\nu} dx^\mu dx^\nu.$$

Of course now we know that dx^μ is really a basis dual vector, it becomes natural to use the term metric and line element interchangeably, and write

$$ds^2 = g_{\mu\nu} dx^\mu dx^\nu.$$

Usually there is a set of equivalent expressions for inner product of two vectors V^μ and W^ν :

$$g_{\mu\nu} V^\mu W^\nu = g(V, W) = ds^2(V, W).$$

As we will see, the metric tensor contains all the information we need to describe the curvature of the manifold. And we will also see that the constancy of the metric components is sufficient for a space to be flat. Although this is not the case when our spacetime manifold $\mathcal{M}$ is curved, we can always construct locally inertial coordinates, as we discussed in the first section. The associated basis vectors then constitute a local Lorentz frame. This is the rigorous statement of the fact that small enough regions of spacetime look like flat (Minkowski) spacetime.

3.5 Differential Forms and Integration on Manifolds

As we mentioned before, there are three important non-tensorial objects. Now we study the second one, the Levi-Civita symbol. This is defined as

$$\epsilon_{\mu_1 \mu_2 \dots \mu_n} = \begin{cases} +1 & \text{if } \mu_1 \mu_2 \dots \mu_n \text{ is an even permutation of } 012 \dots (n-1) \\ -1 & \text{if } \mu_1 \mu_2 \dots \mu_n \text{ is an odd permutation of } 012 \dots (n-1) \\ 0 & \text{otherwise.} \end{cases}$$

By definition, the Levi-Civita symbol has components specified above in any coordinate system. So this is not a tensor, which is why we call it a symbol, for it does not change under coordinate transformation. But its behavior can be related to that of an ordinary tensor by noting that given any $n \times n$ matrix $M^\mu_{\mu'}$, the determinant $|M| = \det M$ obeys

$$\epsilon_{\mu'_1 \mu'_2 \dots \mu'_n} |M| = \epsilon_{\mu_1 \mu_2 \dots \mu_n} M^{\mu_1}_{\mu'_1} M^{\mu_2}_{\mu'_2} \dots M^{\mu_n}_{\mu'_n}.$$

This is just equivalent to the usual formula in terms of matrices of cofactors, as you should check. Now if we let $M^\mu_{\mu'} = \partial x^\mu / \partial x^{\mu'}$, we have

$$\epsilon_{\mu'_1 \mu'_2 \dots \mu'_n} = \left| \frac{\partial x^{\mu'}}{\partial x^\mu} \right| \epsilon_{\mu_1 \mu_2 \dots \mu_n} \frac{\partial x^{\mu_1}}{\partial x^{\mu'_1}} \frac{\partial x^{\mu_2}}{\partial x^{\mu'_2}} \dots \frac{\partial x^{\mu_n}}{\partial x^{\mu'_n}}.$$

We have used the fact that $\frac{\partial x^{\mu'}}{\partial x^\mu}$ is the inverse of the matrix $\frac{\partial x^\mu}{\partial x^{\mu'}}$. And $|M^{-1}| = |M|^{-1}$. So the Levi-Civita symbol transforms similar to the tensor transformation, except for the determinant factor in the front. How can we make a tensor out of this? Consider how the metric transforms:

$$g_{\mu'\nu'} = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^\nu}{\partial x^{\nu'}} g_{\mu\nu}.$$

If we take the determinant on both sides, we find

$$|g'| = \left| \frac{\partial x^{\mu'}}{\partial x^\mu} \right|^{-2} |g|.$$

So the *determinant* of the metric is not a tensor either. "You are not a tensor? Neither am I! Maybe we can hook up and be a tensor!" Therefore, we can define the **Levi-Civita tensor** as

$$\epsilon_{\mu_1 \mu_2 \dots \mu_n} = \sqrt{|g|} \epsilon_{\mu_1 \mu_2 \dots \mu_n}.$$

This is now a real tensor. Now let us shift gears and study a special class of tensors, called **differential forms**. A differential p -form is a complete antisymmetric $(0, p)$ tensor. Thus, scalars are automatically 0-forms, and dual vectors are automatically 1-forms. We also have the 4-form $\epsilon_{\mu\nu\rho\sigma}$. The space of all p -forms are denoted by Λ^p . The space of all p -form over a manifold $\mathcal{M}$ are denoted by $\Lambda^p(\mathcal{M})$. It can be shown that the number of linearly independent p -forms on a n -dimensional vector space is $\frac{n!}{p!(n-p)!}$. Also, there are no p -forms for $p > n$, since all the components will automatically be zero by antisymmetry.

Given a p -form A and a q -form B , we can form a $(p+q)$ form by the operation of **wedge product** $A \wedge B$ by taking the antisymmetric tensor product:

$$(A \wedge B)_{\mu_1 \dots \mu_{p+q}} = \frac{(p+q)!}{p!q!} A_{[\mu_1 \dots \mu_p} B_{\mu_{p+1} \dots \mu_{p+q}]}$$

For example, the wedge product of two 1-forms is

$$(A \wedge B)_{\mu\nu} = 2A_{[\mu} B_{\nu]} = A_\mu B_\nu - A_\nu B_\mu.$$

Note that

$$A \wedge B = (-1)^{pq} B \wedge A.$$

So you can alter the order of a wedge product, if you are careful with signs. You are free to suppress indices when using forms, since we know that all the indices are downstairs, and the tensors are completely antisymmetric.

The *exterior derivative* operator d allows us to differentiate a p -form field and get a $(p+1)$ -form field. This is defined as an appropriately normalized, antisymmetrized partial derivative:

$$(dA)_{\mu_1 \dots \mu_{p+1}} = (p+1) \partial_{[\mu_1} A_{\mu_2 \dots \mu_{p+1}]}$$

The simplest example is the gradient, which is the exterior derivative of a 0-form:

$$(d\phi)_\mu = \partial_\mu \phi.$$

Exterior derivatives obey a modified version of the Leibniz rule when applied to the product of a p -form ω and a q -form η :

$$d(\omega \wedge \eta) = (d\omega) \wedge \eta + (-1)^p \omega \wedge (d\eta).$$

There are so many interesting things we can say about differential forms: [10] mentioned closed and exact forms, and Hodge duality. But this is not the main point right now. We will turn to integration after mentioning the interesting observation that the Maxwell's equation $\partial_{[\mu} F_{\nu\lambda]} = 0$ can be written neatly into

$$dF = 0.$$

The other Maxwell's equation can be written as

$$d(*F) = *J.$$

Now, integration. In $\mathbf{R}^n$, the change of variable formula tells us that

$$\int_{\mathbf{R}^n} d^n x' = \int_{\mathbf{R}^n} \left| \frac{\partial x^{\mu'}}{\partial x^\mu} \right| d^n x.$$

There is a beautiful explanation of this formula from the pointview of differential forms, which arises from the following fact: On a n -dimensional manifold $\mathcal{M}$, the integrand is properly understood as an n -form. In other words, an integral over an n -dimensional region is a map from an n -form to real numbers: $\int_\Sigma : \omega \rightarrow \mathbf{R}$.

The realization that the integrand is actually an n -form is not as trivial as it might seem. We have to explain that in what sense is the integrand antisymmetric, and why it is a $(0, n)$ tensor. Integrals can be written as $\int f(x) d\mu$, where $f(x)$ is a scalar function and $d\mu$ is the volume element or volume measure. The volume element assign every region a real number, namely the volume of that region. For an infinitesimal region, we can approximate it by parallelepipeds, specified by n vectors that defines its edge. One volume element then, should be a map from n vectors to real numbers. Linearity is obvious from the geometry. Therefore our volume element $d\mu$ is a $(0, n)$ tensor. It is antisymmetric because we are taking the orientation of the volume into consideration.

Now the idea is to identify the volume element $d^n x$ as an antisymmetric tensor density constructed with wedge products:

$$d^n x = dx^0 \wedge \cdots \wedge dx^{n-1}.$$

But this is a density, not a tensor! You may be confused, since forms are tensors. Well, certainly if we have two functions $f, g \in \mathcal{F}$, then df, dg are 1-forms and $df \wedge dg$ is a 2-form. But then function appearing above are coordinate functions themselves. When we change coordinates, we replace the one form dx^μ with a new set $dx^{\mu'}$. You see the funny business - ordinarily a coordinate transformation changes components, not one forms themselves. The RHS of the equation above is coordinate dependent. Let us see this in action. We can write according to the definition of wedge product:

$$dx^0 \wedge \cdots \wedge dx^{n-1} = \frac{1}{n!} \epsilon_{\mu_1 \cdots \mu_n} dx^{\mu_1} \wedge \cdots \wedge dx^{\mu_n} \quad (3.4)$$

since both the wedge product and the Levi-Civita symbol are completely antisymmetric. The factor $1/n!$ takes care of the overcounting introduced by summing over the permutations of the indices. Under a coordinate transformation ϵ remains the same, while the one form transforms according to the chain rule, leading to

$$\begin{aligned} \epsilon_{\mu_1 \cdots \mu_n} dx^{\mu_1} \wedge \cdots \wedge dx^{\mu_n} &= \epsilon_{\mu_1 \cdots \mu_n} \frac{\partial x^{\mu_1}}{\partial x^{\mu'_1}} \cdots \frac{\partial x^{\mu_n}}{\partial x^{\mu'_n}} dx^{\mu'_1} \wedge \cdots \wedge dx^{\mu'_n} \\ &= \left| \frac{\partial x^\mu}{\partial x^{\mu'}} \right| \epsilon_{\mu'_1 \cdots \mu'_n} dx^{\mu'_1} \wedge \cdots \wedge dx^{\mu'_n}. \end{aligned}$$

Multiplying by the Jacobian on both sides and using (3.4) gives

$$d^n x' = \left| \frac{\partial x^{\mu'}}{\partial x^\mu} \right| d^n x.$$

Therefore, we indeed have that the volume element $d^n x$ transforms as a density, not a tensor. The integral of a function over a certain volume should not change if we change coordinates (think about integrating $f : \mathbf{R}^3 \rightarrow \mathbf{R}$ over a volume. The value of the integral should not depend on whether we use spherical coordinates or anything else). And we want to define integrals using only the intrinsic properties of the manifold, with no embeddings. But again we know how to construct a tensor out of this. The quantity

$$\sqrt{|g|} d^n x = \sqrt{|g|} dx^0 \wedge dx^1 \cdots \wedge dx^{n-1}$$

is a well-defined tensor. Therefore, given a scalar function ϕ over an n -manifold, its integral is defined to be

$$\int \phi(x) \sqrt{|g|} d^n x.$$

Given explicit forms for $\phi(x)$ and $\sqrt{|g|}$, such an integral can be evaluated by the usual methods of multivariable calculus. The metric determinant serves to automatically take care of the correct transformation properties.

3.6 Covariant Derivatives

In the previous discussion, we have seen that the partial derivative operator ∂_μ is not a tensor. However, equations like $\partial_\mu T^{\mu\nu} = 0$ must be generalized to curved space somehow. So we would like to have a covariant derivative operator, which reduces to partial derivative in flat space with inertial coordinates, but transforms as a tensor on an arbitrary manifold.

In flat space in inertial coordinates, the partial derivative operator ∂_μ is a map from (k, l) tensor fields to $(k, l + 1)$ tensor fields. It acts linearly on its components and obeys the Leibniz rule on tensor products. These are the properties we still want to be true for a covariant derivative. In addition to that there are other useful assumptions that we want to make (as we will see). So let us define the **covariant derivative** operator ∇ to be a map from (k, l) vector fields to $(k, l + 1)$ vector fields such that

1. It is linear: $\nabla(T + S) = \nabla T + \nabla S$;
2. It obeys Leibniz rule: $\nabla(T \otimes S) = (\nabla T) \otimes S + T \otimes (\nabla S)$;
3. It commutes with contractions: $\nabla_\mu(T^\lambda_{\lambda\rho}) = (\nabla T)_\mu^\lambda{}_{\lambda\rho}$;
4. It reduces to partial derivative on scalars: $\nabla_\mu\phi = \partial_\mu\phi$.

So let us investigate how this operator should look like. If ∇ is going to obey the Leibniz rule, it can always be written as the partial derivative plus some linear transformation (as a correction). We are not going to prove this statement, for that we refer to the method shown in [?]. Thus given a vector V^ν , the covariant derivative of this vector on the μ th direction is given by the partial derivative plus a set of $n \times n$ matrix coefficient $(\Gamma_\mu)^\rho_\sigma$, known as the **connection coefficients**, which is simply denoted by $\Gamma_{\mu\sigma}^\rho$. That is,

$$\nabla_\mu V^\nu = \partial_\mu V^\nu + \Gamma_{\mu\lambda}^\nu V^\lambda.$$

The first term accounts for the change in the component of the vector V^ν , and the second term $\Gamma_{\mu\lambda}^\nu V^\lambda$ accounts for the fact that the coordinates may change (not the vector components) from place to place, and this effect is a linear transformation in the first order.

Since we demand this object to be a tensor, it must change according to tensor transformation rules:

$$\nabla_{\mu'} V^{\nu'} = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \nabla_\mu V^\nu.$$

The LHS can be expanded (since we know how partial derivative transforms, see discussions before) as

$$\begin{aligned} \nabla_{\mu'} V^{\nu'} &= \partial_{\mu'} V^{\nu'} + \Gamma_{\mu'\lambda'}^{\nu'} V^{\lambda'} \\ &= \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \partial_\mu V^\nu + \frac{\partial x^\mu}{\partial x^{\mu'}} V^\nu \frac{\partial}{\partial x^\mu} \frac{\partial x^{\nu'}}{\partial x^\nu} + \Gamma_{\mu'\lambda'}^{\nu'} \frac{\partial x^{\lambda'}}{\partial x^\lambda} V^\lambda. \end{aligned}$$

Meanwhile, the RHS is given by definition

$$\frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \nabla_\mu V^\nu = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \partial_\mu V^\nu + \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \Gamma_{\mu\lambda}^\nu V^\lambda.$$

If we equate these two expressions, the first terms on both sides cancel out, and we change the dummy index $\nu \rightarrow \lambda$:

$$\Gamma_{\mu'\lambda'}^{\nu'} \frac{\partial x^{\lambda'}}{\partial x^\lambda} V^\lambda + \frac{\partial x^\mu}{\partial x^{\mu'}} V^\lambda \frac{\partial}{\partial x^\mu} \frac{\partial x^{\nu'}}{\partial x^\lambda} = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \Gamma_{\mu\lambda}^\nu V^\lambda.$$

Since this must be true for any vector V^λ , we can eliminate that on both sides and multiply on both sides by $\partial x^\lambda / \partial x^{\sigma'}$ and relabel $\sigma' \rightarrow \lambda'$, which gives

$$\Gamma_{\mu'\lambda'}^{\nu'} = \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^\lambda}{\partial x^{\lambda'}} \frac{\partial x^{\nu'}}{\partial x^\nu} \Gamma_{\mu\lambda}^\nu - \frac{\partial x^\mu}{\partial x^{\mu'}} \frac{\partial x^\lambda}{\partial x^{\lambda'}} \frac{\partial^2 x^{\nu'}}{\partial x^\mu \partial x^\lambda}.$$

This is not the tensor transformation law. But that's okay, since the connection coefficients are not components of a tensor. They are purposefully constructed to be nontensorial, but in such a way that the combination $\nabla_\mu V^\nu$ transforms as a tensor. The extra terms in the transformation of the partials and the connection coefficients exactly cancel.

We now investigate how the covariant derivative operator acts on 1-forms ω_ν . Once again this can be written as

$$\nabla_\mu \omega_\nu = \partial_\mu \omega_\nu + \tilde{\Gamma}_{\mu\nu}^\lambda \omega_\lambda.$$

We can find what the connection coefficient is by computing

$$\begin{aligned} \nabla_\mu (\omega_\lambda V^\lambda) &= (\nabla_\mu \omega_\lambda) V^\lambda + \omega_\lambda (\nabla_\mu V^\lambda) \\ &= (\partial_\mu \omega_\lambda) V^\lambda + \tilde{\Gamma}_{\mu\lambda}^\sigma \omega_\sigma V^\lambda + \omega_\lambda (\partial_\mu V^\lambda) + \omega_\lambda \Gamma_{\mu\rho}^\lambda V^\rho. \end{aligned}$$

But since $\omega_\lambda V^\lambda$ is a scalar, this must also be given by the partial derivative:

$$\begin{aligned} \nabla_\mu (\omega_\lambda V^\lambda) &= \partial_\mu (\omega_\lambda V^\lambda) \\ &= (\partial_\mu \omega_\lambda) V^\lambda + \omega_\lambda (\partial_\mu V^\lambda). \end{aligned}$$

These two expressions are equal (since both ω_σ and V^λ are arbitrary) if and only if

$$\tilde{\Gamma}_{\mu\lambda}^\sigma = -\Gamma_{\mu\lambda}^\sigma.$$

Thus the covariant derivative on a 1-form is given by

$$\boxed{\nabla_\mu \omega_\nu = \partial_\mu \omega_\nu - \Gamma_{\mu\nu}^\lambda \omega_\lambda.}$$

Now we can deduce how the covariant derivative acts on a tensor. For each upper index we introduce a term with a single $+\Gamma$ and for each lower index with a term $-\Gamma$:

$$\begin{aligned} \nabla_\sigma T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} &= \partial_\sigma T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} + \\ &\quad + \Gamma_{\sigma\lambda}^{\mu_1} T^{\lambda \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} + \Gamma_{\sigma\lambda}^{\mu_2} T^{\mu_1 \lambda \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} + \dots \\ &\quad - \Gamma_{\sigma\nu_1}^\lambda T^{\mu_1 \mu_2 \dots \mu_k}_{\lambda \nu_2 \dots \nu_l} - \Gamma_{\sigma\nu_2}^\lambda T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \lambda \dots \nu_l} - \dots \end{aligned}$$

To define a covariant derivative, then, we need to put a connection on our manifold, which is specified in some coordinate system by a set of coefficients

$\Gamma_{\mu\nu}^\lambda$. So there are a total of $n^3 = 64$ of them. Evidently, we can define a large number of connections on any manifold, and each of them implies a distinct notion of covariant differentiation. But in GR it turns out that every metric defined a unique connection, which we will stick to. Let us see how that works.

The first thing to notice is that the difference of two connections

$$S_{\mu\nu}^\lambda = \Gamma_{\mu\nu}^\lambda - \hat{\Gamma}_{\mu\nu}^\lambda$$

is a tensor. To see this, notice that

$$\begin{aligned}\nabla_\mu V^\lambda - \hat{\nabla}_\mu V^\lambda &= \partial_\mu V^\lambda + \Gamma_{\mu\nu}^\lambda V^\nu - \partial_\mu V^\lambda - \hat{\Gamma}_{\mu\nu}^\lambda V^\nu \\ &= S_{\mu\nu}^\lambda V^\nu.\end{aligned}$$

The LHS is a tensor, and V^ν is arbitrary. So the claim follows. Next, given a connection $\Gamma_{\mu\nu}^\lambda$, we can immediately form another connection simply by permuting the lower indices. These coefficients also transforms like they should, since the derivatives appeared in the last term of the connection transformation equation can be permuted. There is thus a tensor we can associate with any given connection, known as the **torsion tensor**, defined by

$$T_{\mu\nu}^\lambda = \Gamma_{\mu\nu}^\lambda - \Gamma_{\nu\mu}^\lambda = 2\Gamma_{[\mu\nu]}^\lambda.$$

A connection that is symmetric on its lower indices is known as **torsion-free**, i.e. the torsion with respect to this connection vanishes. We can now define a unique connection on a manifold with a metric $g_{\mu\nu}$ by introducing two additional properties:

- It is torsion-free: $\Gamma_{\mu\nu}^\lambda = \Gamma_{(\mu\nu)}^\lambda$;
- It is metric compatible: $\nabla_\rho g_{\mu\nu} = 0$.

This implies a couple of nice properties. First, we have

$$\nabla_\lambda \epsilon_{\mu\nu\rho\sigma} = 0 \quad \nabla_\rho g^{\mu\nu} = 0.$$

Second, a metric compatible covariant derivative commutes with raising and lowering of indices:

$$g_{\mu\lambda} \nabla_\rho V^\lambda = \nabla_\rho (g_{\mu\lambda} V^\lambda) = \nabla_\rho V_\mu.$$

The claim is that there is exactly one torsion-free connection on a given manifold that is compatible with some given metric on that manifold. We show this by deriving a unique expression of the connection coefficients in terms of the metric. To this end, we expand out the equation of metric compatibility for three different permutation of the indices:

$$\begin{aligned}\nabla_\rho g_{\mu\nu} &= \partial_\rho g_{\mu\nu} - \Gamma_{\rho\mu}^\lambda g_{\lambda\nu} - \Gamma_{\rho\nu}^\lambda g_{\mu\lambda} = 0 \\ \nabla_\mu g_{\nu\rho} &= \partial_\mu g_{\nu\rho} - \Gamma_{\mu\nu}^\lambda g_{\lambda\rho} - \Gamma_{\mu\rho}^\lambda g_{\nu\lambda} = 0 \\ \nabla_\nu g_{\rho\mu} &= \partial_\nu g_{\rho\mu} - \Gamma_{\nu\rho}^\lambda g_{\lambda\mu} - \Gamma_{\nu\mu}^\lambda g_{\rho\lambda} = 0.\end{aligned}$$

Subtracting the second and third of these from the first, and use the symmetry of the connection to obtain

$$\partial_\rho g_{\mu\nu} - \partial_\mu g_{\nu\rho} - \partial_\nu g_{\rho\mu} + 2\Gamma_{\mu\nu}^\lambda g_{\lambda\rho} = 0.$$

This gives

$$\Gamma_{\mu\nu}^{\sigma} = \frac{1}{2}g^{\sigma\rho}(\partial_{\mu}g_{\nu\rho} + \partial_{\nu}g_{\rho\mu} - \partial_{\rho}g_{\mu\nu}).$$

This formula is one of the most important in this subject. Commit it to memory. You can check the RHS of this equation transforms like a connection. This is called the **Levi-Civita connection**.

You see that when $g_{\mu\nu} = \delta_{\mu\nu}$, the connection vanishes and the covariant derivative reduces to the ordinary partial derivative. It is worth mentioning that the coefficients of the Levi-Civita connection in flat space vanishes in Cartesian coordinates, but not in curvilinear coordinate systems. The example of polar coordinates are given in [10]. Conversely, even in a curved space it is still possible to make the Levi-Civita connection vanish at a given point. This is because, as we argued in the beginning of this chapter, the first derivative of the metric can always be made to be equal to zero at a point.

We also need to note that converting partial derivatives into covariant derivatives is not always necessary to construct a well-defined tensor. In particular, the exterior derivative and the commutator are both well-defined tensors in terms of the partial derivatives, essentially because both involve an antisymmetrization that cancels the nontensorial piece of the partial derivative transformation law. The same feature implies that they could be equally well-defined in terms of (torsion-free) covariant derivatives. Thus, we have, for a 1-form ω_{μ} and vector fields X^{μ}, Y^{μ} , that

$$(d\omega)_{\mu\nu} = \partial_{[\mu}\omega_{\nu]} = \nabla_{[\mu}\omega_{\nu]} \\ [X, Y]^{\mu} = X^{\lambda}\partial_{\lambda}Y^{\mu} - Y^{\lambda}\partial_{\lambda}X^{\mu} = X^{\lambda}\nabla_{\lambda}Y^{\mu} - Y^{\lambda}\nabla_{\lambda}X^{\mu}. \quad (3.5)$$

3.7 Parallel Transport and Geodesics

Now that we know how to take covariant derivatives, let us think about what it means. We think of a derivative as a way of quantifying how fast something is changing. An ordinary function defines a number at each point in spacetime, and it is straightforward to compare two different numbers, so we shouldn't be surprised that the partial derivative of functions remain valid on arbitrary manifolds. But in the case of tensors, it is a map from vectors and dual vectors to real numbers, and it is not immediately clear how to compare such maps on different points. Can we think of covariant derivative as somehow measuring the rate of change of tensors? And if so, with respect to what? The answer is, *the covariant derivative quantifies the instantaneous rate of change of a tensor field in comparison to what the tensor would be if it were "parallel transported"*. In other words, a connection defines a specific way of keeping a tensor constant along some path, on the basis of which we can compare nearby tensors. Parallel transport is supposed to be the curved-space generalization of the concept of "keeping the vector constant" as we move along a path, similarly for a tensor of arbitrary rank. Given a curve $x^{\mu}(\lambda)$, the requirement of constancy of a tensor $T^{\mu_1\mu_2\cdots\mu_k}_{\nu_1\nu_2\cdots\nu_l}$ along this curve in flat space is simply that the components be constant:

$$\frac{d}{d\lambda}T^{\mu_1\mu_2\cdots\mu_k}_{\nu_1\nu_2\cdots\nu_l} = \frac{dx^{\mu}}{d\lambda}\frac{\partial}{\partial x^{\mu}}T^{\mu_1\mu_2\cdots\mu_k}_{\nu_1\nu_2\cdots\nu_l} = 0.$$

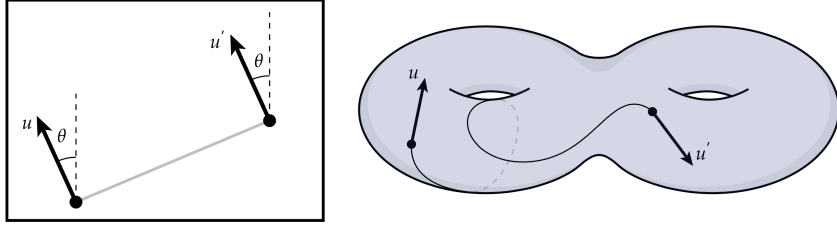


Figure 3.6: Parallel-transport a vector in flat and curved space along a path. The vector can change, and this is due to the intrinsic property of the curved space, known as the curvature, as we will see.

To make this properly tensorial we simply replace the partial derivative by a covariant one, and define the **directional covariant derivative** to be

$$\frac{D}{d\lambda} = \frac{dx^\mu}{d\lambda} \nabla_\mu.$$

This is a map, defined only along the path, from (k, l) tensors to (k, l) tensors. We then define the **parallel transport** of the tensor T along the path $x^\mu(\lambda)$ to be the requirement that the covariant derivative of T along the path vanishes:

$$\left(\frac{D}{d\lambda} T \right)^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} \equiv \frac{dx^\sigma}{d\lambda} \nabla_\sigma T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} = 0.$$

This is a well-defined tensor equation, known as the equation of parallel transport. For a vector it takes the form

$$\frac{d}{d\lambda} V^\mu + \Gamma^\mu_{\sigma\rho} \frac{dx^\sigma}{d\lambda} V^\rho = 0.$$

This is a first-order differential equation. Given an initial value (i.e. the vector at some point along the path), there will be a unique continuation of the vector to other points along the path that solves the equation above. We say such a vector is parallel-transported (and of course this generalizes to tensors).

Obviously parallel transport depends on the connection, and different connections lead to different answers. If the connection is metric-compatible, the metric is always parallel transported with respect to it:

$$\frac{D}{d\lambda} g_{\mu\nu} = \frac{dx^\sigma}{d\lambda} \nabla_\sigma g_{\mu\nu} = 0.$$

It follows that the inner product of two parallel-transported vectors is preserved. If V^μ and W^ν are parallel-transported along a curve $x^\sigma(\lambda)$, we have

$$\frac{D}{d\lambda} (g_{\mu\nu} V^\mu W^\nu) = \left(\frac{D}{d\lambda} g_{\mu\nu} \right) V^\mu W^\nu + g_{\mu\nu} \left(\frac{D}{d\lambda} V^\mu \right) W^\nu + g_{\mu\nu} V^\mu \left(\frac{D}{d\lambda} W^\nu \right) = 0.$$

Thus parallel transport with respect to a metric-compatible connection preserves the norm of vectors, the sense of orthogonality, and so on.

Next, we discuss geodesics. A geodesic is the curved-space generalization of the notion of a straight line in Euclidean space. We all know that a straight line is the path of the shortest distance between two points. But there is an

equally good definition: it is a path that parallel-transport its own tangent vector. It will turn out that these two definition coincide if and only if the connection is the Levi-Civita connection.

We will start with this definition that a **geodesic** is a path that parallel-transport its own tangent vector. The tangent vectors to a path $x^\mu(\lambda)$ is $dx^\mu/d\lambda$. The condition that it be parallel-transported is thus

$$\frac{D}{d\lambda} \frac{dx^\mu}{d\lambda} = 0.$$

Alternatively, we can write

$$\boxed{\frac{d^2 x^\mu}{d\lambda^2} + \Gamma_{\rho\sigma}^\mu \frac{dx^\rho}{d\lambda} \frac{dx^\sigma}{d\lambda} = 0.} \quad (3.6)$$

This is the **geodesic equation**. This is another equation you should memorize. We can see that in flat space with $g_{\mu\nu} = \delta_{\mu\nu}$ this reduces to $d^2 x^\mu d\lambda^2 = 0$, which is the equation for a straight line.

There is also a second definition. We define the geodesic connecting two points to be the curve that extremizes the distance between two points. Let us do the case for a timelike path, without loss of generality. We therefore consider the proper time functional,

$$\tau = \int \left(-g_{\mu\nu} \frac{dx^\mu}{d\lambda} \frac{dx^\nu}{d\lambda} \right)^{1/2} d\lambda$$

where the integral is over the path. The integral is of the form $\tau = \int \sqrt{-f} d\lambda$, so the variation looks like

$$\delta\tau = \int \delta \sqrt{-f} d\lambda = - \int \frac{1}{2} (-f)^{-1/2} (\delta f) d\lambda.$$

So it makes things easier if we parametrize the path using proper time τ itself. So that the tangent vector is the 4-velocity U^μ , and

$$f = g_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} = g_{\mu\nu} U^\mu U^\nu = -1.$$

Then we have

$$\delta\tau = -\frac{1}{2} \int \delta f d\lambda.$$

So finding the stationary points of proper time functional is equivalent to finding the stationary points of the simpler integral

$$I = \frac{1}{2} \int f d\tau = \frac{1}{2} \int g_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} d\tau.$$

Stationary points of I will of course obey the Euler-Lagrange equations. But it would be simpler to consider directly the change in the integral under infinitesimal variations of path

$$\begin{aligned} x^\mu &\rightarrow x^\mu + \delta x^\mu \\ g_{\mu\nu} &\rightarrow g_{\mu\nu} + (\partial_\sigma g_{\mu\nu}) \delta x^\sigma. \end{aligned}$$

The second line comes from the Taylor expansion in curved spacetime, which uses partial derivatives, not the covariant derivative, since we are simply thinking of the component $g_{\mu\nu}$ as functions on spacetime in some specific coordinate system. Then we see that

$$\delta I = \frac{1}{2} \int \left[\partial_\sigma g_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} \delta x^\sigma + g_{\mu\nu} \frac{d(\delta x^\mu)}{d\tau} \frac{dx^\nu}{d\tau} + g_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{d(\delta x^\nu)}{d\tau} \right] d\tau.$$

The last two terms can be integrated by path. For example:

$$\begin{aligned} \frac{1}{2} \int \left[g_{\mu\nu} \frac{dx^\mu}{d\tau} \frac{d(\delta x^\nu)}{d\tau} \right] d\tau &= -\frac{1}{2} \int \left[g_{\mu\nu} \frac{d^2 x^\mu}{d\tau^2} + \frac{dg_{\mu\nu}}{d\tau} \frac{dx^\mu}{d\tau} \right] \delta x^\nu d\tau \\ &= -\frac{1}{2} \int \left[g_{\mu\nu} \frac{d^2 x^\mu}{d\tau^2} + \partial_\sigma g_{\mu\nu} \frac{dx^\sigma}{d\tau} \frac{dx^\mu}{d\tau} \right] \delta x^\nu d\tau. \end{aligned}$$

The boundary term vanishes since we require δx^μ to vanish at the endpoints of the path. Then after some rearranging dummy indices, we find the variation to be

$$\delta I = - \int \left[g_{\mu\sigma} \frac{d^2 x^\mu}{d\tau^2} + \frac{1}{2} (\partial_\mu g_{\nu\sigma} + \partial_\nu g_{\sigma\mu} - \partial_\sigma g_{\mu\nu}) \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} \right] \delta x^\sigma d\tau.$$

We want δI to vanish for any variation δx^σ , since we are searching for stationary points, this implies

$$g_{\mu\sigma} \frac{d^2 x^\mu}{d\tau^2} + \frac{1}{2} (\partial_\mu g_{\nu\sigma} + \partial_\nu g_{\sigma\mu} - \partial_\sigma g_{\mu\nu}) \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} = 0.$$

Multiplying by the inverse metric $g^{\rho\sigma}$ finally gives

$$\frac{d^2 x^\rho}{d\tau^2} + \frac{1}{2} g^{\rho\sigma} (\partial_\mu g_{\nu\sigma} + \partial_\nu g_{\sigma\mu} - \partial_\sigma g_{\mu\nu}) \frac{dx^\mu}{d\tau} \frac{dx^\nu}{d\tau} = 0.$$

This is precisely the geodesic equation, but with the specific choice of Levi-Civita connection. Thus, on a manifold with metric, extremals of the length functional are curves that parallel transport their tangent vector with respect to the Levi-Civita connection associated with that metric.

3.8 The Riemann Curvature Tensor

In a curved space, the parallel transport of a vector around a closed loop will result in a transformation of the vector. The resulting transformation depends on the intrinsic nature of the curved space, called the curvature. In flat space, if we parallel-transport a vector around a closed loop, the resulting vector will be exactly the same as the vector that we started with. This is not the case if we consider, say, the sphere $\mathbf{S}^2$. The picture below shows that this action will result in a transformation of the vector.

So how to characterize this transformation on vectors? Since spacetime looks flat in sufficiently small regions, our loop will be specified by two infinitesimal vectors δx^μ and δx^ν . We imagine parallel-transporting a vector V^μ along this closed loop, as Fig.3.8 suggests. We know the action of parallel transport is independent of coordinates, so there should be some tensor that tells us how the vector changes when it comes back to its starting point. It

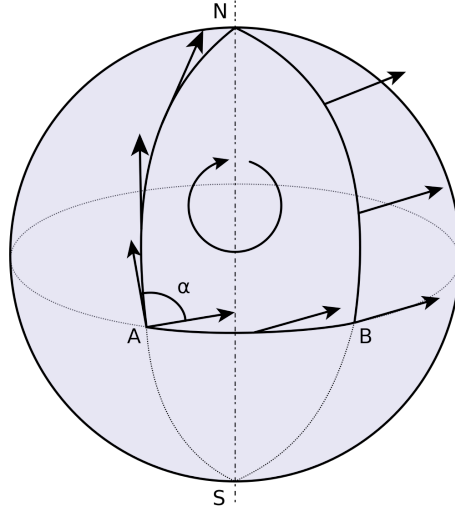


Figure 3.7: Parallel-transport a vector along the closed loop $A \rightarrow B \rightarrow N \rightarrow A$, will result in a deviation by an angle α compared to what the vector was before at the point A .

will be a linear transformation on a vector, and therefore involve one upper and one lower index. But it will also depend on the two vectors δx^μ and δx^ν . Therefore the tensor we want should involve two additional lower indices to contract with δx^μ and δx^ν . Furthermore, the tensor should be antisymmetric in these two indices, since interchanging the vectors corresponds to travel along the loop in the opposite direction. We therefore expect the change δV^ρ experienced by this vector when parallel-transported around the loop should be of the form

$$\delta V^\rho = R^\rho_{\sigma\mu\nu} V^\sigma \delta x^\mu \delta x^\nu$$

where $R^\rho_{\sigma\mu\nu}$ is a (1,3) tensor known as the **Riemann curvature tensor** or simply curvature tensor. It is antisymmetric in the last two indices:

$$R^\rho_{\sigma\mu\nu} = -R^\rho_{\sigma\nu\mu}.$$

We would like to know what $R^\rho_{\sigma\mu\nu}$ is in terms of the intrinsic properties of the space, possibly involving the metric, and the connection (by the way, usually the connection coefficients $\Gamma^\rho_{\nu\sigma}$ are called the **Christoffel symbol**). Using some *very schematic notation*, we have that

$$\delta V = V_A - V_{A'} = [(V_C - V_D) - (V_B - V_A)] - [(V_C - V_B) - (V_D - V_{A'})].$$

This notation of course makes no sense, since in curved space subtracting two vectors is simply not defined. But here the idea is that subtracting really means comparing. Since the covariant derivative of a tensor in a certain direction measures how much the tensor changes relative to what it would have been if it had been parallel transported (for the covariant derivative of a tensor in a direction along which it is parallel-transported is zero), we can write the equation above into a form that makes sense:

$$\delta V^\rho = \delta x^\mu \delta x^\nu \nabla_\mu \nabla_\nu V^\rho - \delta x^\nu \delta x^\mu \nabla_\nu \nabla_\mu V^\rho = \delta x^\mu \delta x^\nu [\nabla_\mu, \nabla_\nu] V^\rho.$$

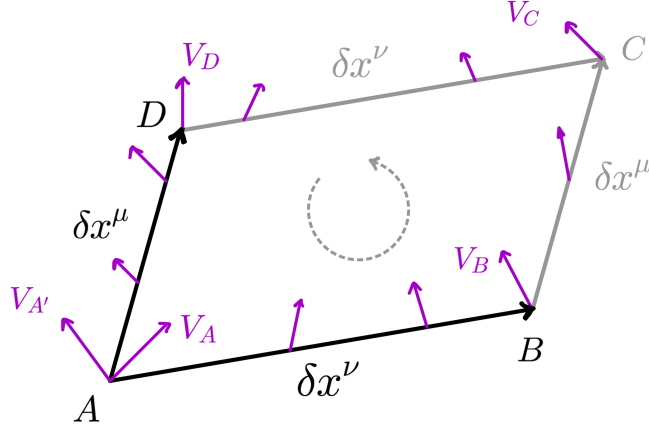


Figure 3.8: We parallel-transporting a vector V along this infinitesimal closed loop. The path is given by $A \rightarrow B \rightarrow C \rightarrow D \rightarrow A$. In the starting point the vector is labeled to be V_A , and similarly for other endpoints. The transformed vector after the entire trip back at A is labeled by $V_{A'}$.

Now the goal is clear. The last term on the RHS is really just $R^\rho_{\sigma\mu\nu} V^\sigma \delta x^\mu \delta x^\nu$. So we want to compute the expression $[\nabla_\mu, \nabla_\nu] V^\rho$. Our calculation below works under the assumption that the connection is torsion-free. The result is slightly different if this is not the case, see [10] for details. We proceed:

$$\begin{aligned}
& [\nabla_\mu, \nabla_\nu] V^\rho \\
&= \nabla_\mu \nabla_\nu V^\rho - \nabla_\nu \nabla_\mu V^\rho \\
&= \nabla_\mu [\partial_\nu V^\rho + \Gamma^\rho_{\nu\sigma} V^\sigma] - (\mu \leftrightarrow \nu) \\
&= \partial_\mu [\partial_\nu V^\rho + \Gamma^\rho_{\nu\sigma} V^\sigma] - \Gamma^\lambda_{\mu\nu} [\partial_\lambda V^\rho + \Gamma^\rho_{\lambda\sigma} V^\sigma] + \Gamma^\rho_{\mu\lambda} [\partial_\nu V^\lambda + \Gamma^\lambda_{\nu\sigma} V^\sigma] - (\mu \leftrightarrow \nu) \\
&= \partial_\mu \partial_\nu V^\rho + \underbrace{\partial_\mu (\Gamma^\rho_{\nu\sigma}) V^\sigma + \Gamma^\rho_{\nu\sigma} \partial_\mu V^\sigma}_{\partial_\mu (\Gamma^\rho_{\nu\sigma} V^\sigma)} - \Gamma^\lambda_{\mu\nu} \partial_\lambda V^\rho - \Gamma^\rho_{\mu\nu} \Gamma^\lambda_{\lambda\sigma} V^\sigma + \Gamma^\rho_{\mu\lambda} \partial_\nu V^\lambda + \Gamma^\rho_{\mu\lambda} \Gamma^\lambda_{\nu\sigma} V^\sigma \\
&\quad - \partial_\nu \partial_\mu V^\rho - \underbrace{\partial_\nu (\Gamma^\rho_{\mu\sigma}) V^\sigma + \Gamma^\rho_{\mu\sigma} \partial_\nu V^\sigma}_{\partial_\nu (\Gamma^\rho_{\mu\sigma} V^\sigma)} + \Gamma^\lambda_{\nu\mu} \partial_\lambda V^\rho + \Gamma^\rho_{\nu\mu} \Gamma^\lambda_{\lambda\sigma} V^\sigma - \Gamma^\rho_{\nu\lambda} \partial_\mu V^\lambda - \Gamma^\rho_{\nu\lambda} \Gamma^\lambda_{\mu\sigma} V^\sigma \\
&= \left[\partial_\mu \Gamma^\rho_{\nu\sigma} - \partial_\nu \Gamma^\rho_{\mu\sigma} + \Gamma^\rho_{\mu\lambda} \Gamma^\lambda_{\nu\sigma} - \Gamma^\rho_{\nu\lambda} \Gamma^\lambda_{\mu\sigma} \right] V^\sigma.
\end{aligned}$$

This gives

$$R^\rho_{\sigma\mu\nu} = \partial_\mu \Gamma^\rho_{\nu\sigma} - \partial_\nu \Gamma^\rho_{\mu\sigma} + \Gamma^\rho_{\mu\lambda} \Gamma^\lambda_{\nu\sigma} - \Gamma^\rho_{\nu\lambda} \Gamma^\lambda_{\mu\sigma}.$$

Notice that the antisymmetry of $R^\rho_{\sigma\mu\nu}$ in the last two indices is immediate from this formula. Now, as usual, it is straightforward to compute the action of the commutator on a tensor of arbitrary rank. This is given by

$$\begin{aligned}
[\nabla_\mu, \nabla_\nu] T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} &= R^{\mu_1}_{\lambda\rho\sigma} T^{\lambda \mu_2 \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} + R^{\mu_2}_{\lambda\rho\sigma} T^{\mu_1 \lambda \dots \mu_k}_{\nu_1 \nu_2 \dots \nu_l} + \dots \\
&\quad - R^\lambda_{\nu_1 \rho \sigma} T^{\mu_1 \mu_2 \dots \mu_k}_{\lambda \nu_2 \dots \nu_l} - R^\lambda_{\nu_2 \rho \sigma} T^{\mu_1 \mu_2 \dots \mu_k}_{\nu_1 \lambda \dots \nu_l} + \dots
\end{aligned}$$

Now we give a criterion to determine whether a space is flat or not. We know that even for a flat space, the metric $g_{\mu\nu}$ can take horrible forms depending on the coordinate chosen. Also, as we have seen, the Christoffel symbol is not

a good indicator of the flatness of the space, since in curved space we can also make the Christoffel symbol vanish. But not we have the following result:

Theorem 3.1. *If a coordinate system exists in which the component of the metric are constant, the Riemann curvature tensor will vanish. Conversely, if the Riemann curvature tensor vanishes, we can always construct a coordinate system in which the metric components are constant.*

The proof is given in [10], though not very rigorous. For a rigorous version, you can check [2]. We will omit the proof here. But the punchline is, given a metric $g_{\mu\nu}$, we can calculate the Riemann curvature tensor associated with that metric, and if the curvature vanishes, then the space is flat; otherwise, it has curvature.

The Riemann curvature tensor has lots of symmetries. This reduces the independent component of the Riemann curvature tensor. Since the tensor is of the form $R^\rho_{\sigma\mu\nu}$, you may guess at first that there are n^4 independent components. But there are a lot of symmetry and antisymmetry going on here, as we will see.

We define

$$R_{\rho\sigma\mu\nu} = g_{\rho\lambda} R^\lambda_{\sigma\mu\nu}.$$

Let us consider the component of this tensor in locally inertial coordinates, which exists by our previous argument in the beginning of this chapter. Then at the origin of the coordinates, since $\partial_\alpha g_{\beta\gamma} = 0$ for all α, β, γ in this coordinate, the Christoffel symbols vanishes. We have, in the locally inertial coordinate centered at the point $p \in \mathcal{M}$, that

$$\begin{aligned} R_{\rho\sigma\mu\nu}(p) &= g_{\rho\lambda}(\partial_\mu \Gamma^\lambda_{\nu\sigma} - \partial_\nu \Gamma^\lambda_{\mu\sigma}) \\ &= \frac{1}{2} g_{\rho\lambda} g^{\lambda\tau} (\partial_\mu \partial_\nu g_{\sigma\tau} + \partial_\mu \partial_\sigma g_{\tau\nu} - \partial_\mu \partial_\tau g_{\nu\sigma} - \partial_\nu \partial_\mu g_{\sigma\tau} - \partial_\nu \partial_\sigma g_{\tau\mu} + \partial_\nu \partial_\tau g_{\mu\sigma}) \\ &= \frac{1}{2} g_{\rho\lambda} g^{\lambda\tau} (\partial_\mu \partial_\sigma g_{\tau\nu} - \partial_\mu \partial_\tau g_{\nu\sigma} - \partial_\nu \partial_\sigma g_{\tau\mu} + \partial_\nu \partial_\tau g_{\mu\sigma}) \\ &= \frac{1}{2} (\partial_\mu \partial_\sigma g_{\rho\nu} - \partial_\mu \partial_\rho g_{\nu\sigma} - \partial_\nu \partial_\sigma g_{\rho\mu} + \partial_\nu \partial_\rho g_{\mu\sigma}). \end{aligned}$$

For this expression we can notice immediately several properties of $R_{\rho\sigma\mu\nu}$: It is antisymmetric in its first two indices:

$$R_{\rho\sigma\mu\nu} = -R_{\sigma\rho\mu\nu}. \quad (3.7)$$

It is antisymmetric on its last two indices, which we already know:

$$R_{\rho\sigma\mu\nu} = -R_{\rho\sigma\nu\mu}. \quad (3.8)$$

It is invariant under interchange of the first pair of indices with the second:

$$R_{\rho\sigma\mu\nu} = R_{\mu\nu\rho\sigma}. \quad (3.9)$$

We can also show that the sum of cyclic permutations of the last three indices vanishes (just plug in the result above):

$$R_{\rho\sigma\mu\nu} + R_{\rho\mu\nu\sigma} + R_{\rho\nu\sigma\mu} = 0.$$

Using (3.8), we see that this is equivalent to

$$R_{\rho[\sigma\mu\nu]} = 0.$$

An immediate consequence of the equation above is that

$$R_{[\rho\sigma\mu\nu]} = \pm R_{\rho[\sigma\mu\nu]} \pm R_{\sigma[\rho\mu\nu]} \pm R_{\mu[\rho\sigma\nu]} \pm R_{\nu[\rho\sigma\mu]} = 0.$$

These equations are derived in a very special reference frame. But that's okay. Since they are all perfectly valid tensor equations, if they are true in one frame, then they are true in every frame. Thus we can have full confidence in the generality of the equations above.

How do these symmetries affect the number of independent components of the Riemann tensor? Some arguments of combinatorics, which can be found in [10], shows that the number of independent components of the Riemann tensor $R_{\sigma\mu\nu}^\rho$ is

$$\frac{n^2(n^2 - 1)}{12}.$$

Here $n = \dim \mathcal{M}$. So when $n = 1$, there is no curvature. When $n = 4$, the number of independent components is 20. But if you go back to the argument before when we were showing the existence of locally inertial coordinates, this is exactly the degrees of freedom in the second derivatives of the metric that we could not set to zero. This should reinforce your confidence that the Riemann tensor is an appropriate measure of curvature.

We now show an important identity, called the ***Bianchi identity***, which says that

$$\nabla_{[\lambda} R_{\rho\sigma]\mu\nu} = 0. \quad (3.10)$$

To show this, we once again go back to the locally inertial frame, and we see that by substituting our previous calculation, we have that

$$\nabla_\lambda R_{\rho\sigma\mu\nu} = \partial_\lambda R_{\rho\sigma\mu\nu} = \frac{1}{2} \partial_\lambda (\partial_\mu \partial_\sigma g_{\rho\nu} - \partial_\mu \partial_\rho g_{\nu\sigma} - \partial_\nu \partial_\sigma g_{\rho\mu} + \partial_\nu \partial_\rho g_{\mu\sigma}).$$

Then we find that

$$\nabla_\lambda R_{\rho\sigma\mu\nu} + \nabla_\rho R_{\sigma\lambda\mu\nu} + \nabla_\sigma R_{\lambda\rho\mu\nu} = 0$$

simply by plugging in. Then by antisymmetry in the first two indices of the Riemann tensor (3.7), we find that the expression above is equivalent to $2\nabla_{[\lambda} R_{\rho\sigma]\mu\nu} = 0$. This proves the claim.

Now, we define the ***Ricci tensor*** as:

$$\boxed{R_{\mu\nu} = R^\lambda_{\mu\lambda\nu}.$$

The Ricci tensor is symmetric, which can be shown by rewriting its expressions:

$$\begin{aligned} R_{\mu\nu} &= \partial_\alpha \Gamma_{\mu\nu}^\alpha - \partial_\nu \Gamma_{\alpha\mu}^\alpha + \Gamma_{\alpha\beta}^\alpha \Gamma_{\nu\mu}^\beta - \Gamma_{\nu\beta}^\alpha \Gamma_{\alpha\mu}^\beta \\ &= \partial_\alpha \Gamma_{\mu\nu}^\alpha - \partial_\nu \partial_\mu (\ln |g|^{1/2}) + \Gamma_{\alpha\beta}^\alpha \Gamma_{\nu\mu}^\beta - \Gamma_{\nu\beta}^\alpha \Gamma_{\alpha\mu}^\beta \end{aligned}$$

and this is symmetric under the change of indices $\mu \leftrightarrow \nu$. To justify the second expression in the second step, we find have that

$$\Gamma_{\alpha\mu}^\alpha = \frac{1}{2}g^{\alpha\rho}(\partial_\alpha g_{\mu\rho} + \partial_\mu g_{\rho\alpha} - \partial_\rho g_{\alpha\mu}) = \frac{1}{2}g^{\alpha\rho}\partial_\mu g_{\rho\alpha}.$$

The first and the third term in the sum cancels after the exchange of indices $\mu \leftrightarrow \rho$ on the third term. Next, we use Jacobi's formula

$$\partial_\mu \ln(\det A) = \text{tr}(A^{-1}\partial_\mu A)$$

which gives

$$\partial_\mu \det g = (\det g) \text{tr}(g^{-1}\partial_\mu g) = (\det g)g^{\alpha\lambda}\partial_\mu g_{\alpha\lambda}.$$

Then finally by chain rule we have

$$\partial_\mu \ln(\sqrt{\det g}) = \frac{1}{\sqrt{\det g}} \frac{1}{2} \frac{1}{\sqrt{\det g}} \partial_\mu \det g = \text{tr}(g^{-1}\partial_\mu g) = \frac{1}{2}g^{\alpha\lambda}\partial_\mu g_{\alpha\lambda} = \Gamma_{\alpha\mu}^\alpha.$$

Thus

$$R_{\mu\nu} = R_{\nu\mu}.$$

The trace of the Ricci tensor is the ***Ricci scalar***:

$$R = R^\mu{}_\mu = g^{\mu\nu} R_{\mu\nu}.$$

The Ricci tensor and scalar contain all of the information about traces of the Riemann tensor, leaving the trace-free parts. These are captured by the Weyl tensor, which is discussed in [10].

An especially useful form of Bianchi identity (3.10) comes from contracting twice on the cyclic sum:

$$\begin{aligned} 0 &= g^{\nu\sigma} g^{\mu\lambda} (\nabla_\lambda R_{\rho\sigma\mu\nu} + \nabla_\rho R_{\sigma\lambda\mu\nu} + \nabla_\sigma R_{\lambda\rho\mu\nu}) \\ &= \nabla^\mu R_{\rho\mu} - \nabla_\rho R + \nabla^\nu R_{\rho\nu}. \end{aligned}$$

This can be simplified into

$$\nabla^\mu R_{\rho\mu} = \frac{1}{2} \nabla_\rho R. \quad (3.11)$$

We define the ***Einstein tensor*** as

$$G_{\mu\nu} = R_{\mu\nu} - \frac{1}{2} R g_{\mu\nu}.$$

We see that the Bianchi identity (3.11) is equivalent to

$$\nabla^\mu G_{\mu\nu} = 0.$$

3.9 Geodesic Deviation

The Riemann tensor shows up as a consequence of curvature in one more way: geodesic deviation. We know in flat geometry parallel lines remain parallel forever. This need not be true when the space is curved (for example, on a sphere parallel geodesics can meet at the north and south pole). We would like to quantify this behavior for an arbitrary curved space.

Consider a one-parameter family of geodesics, $\gamma_s(t)$. That is, for each $s \in \mathbf{R}$, γ_s is a geodesic parametrized by t . The collection of these curves defines a smooth 2-dimensional submanifold Σ . Let $x^\mu(s, t)$ denote the set of points of Σ on $\mathcal{M}$. The map $(s, t) \rightarrow \gamma_s(t)$ is smooth with a smooth inverse. Thus we can let s, t as coordinates of the submanifold Σ . The vector field

$$T^\mu = \frac{\partial x^\mu}{\partial t} = \frac{dx_s^\mu(t)}{dt}$$

are tangent vectors to the geodesics. The deviation vectors

$$S^\mu = \frac{\partial x^\mu}{\partial s}$$

is another vector field that we can define. This name derives from the informal notion that S^μ points from one geodesic toward the neighboring ones.

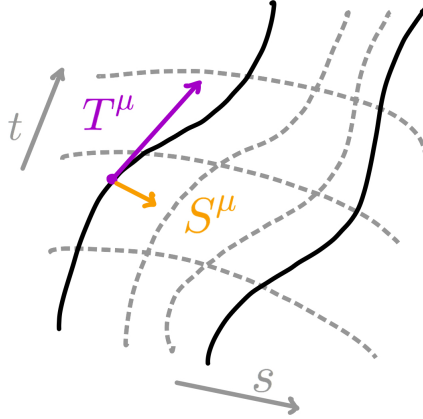


Figure 3.9

The idea that S^μ points from one geodesic to the next inspires us to define the relative velocity of geodesics, given by

$$V^\mu = (\nabla_T S)^\mu = T^\rho \nabla_\rho S^\mu$$

and the relative acceleration of geodesics, given by

$$A^\mu = (\nabla_T V)^\mu = T^\rho \nabla_\rho V^\mu.$$

Since S, T are basis vectors adapted to a coordinate system, their commutator vanishes (since in a coordinate system the basis vector will simply be $\partial_s = S$ and $\partial_t = T$):

$$[S, T] = 0$$

which we discussed in section 2.3. We then have from (3.5) that

$$S^\rho \nabla_\rho T^\mu = T^\rho \nabla_\rho S^\mu.$$

Now, we compute

$$\begin{aligned} A^\mu &= T^\rho \nabla_\rho (T^\sigma \nabla_\sigma S^\mu) \\ &= T^\rho \nabla_\rho (S^\sigma \nabla_\sigma T^\mu) \\ &= (T^\rho \nabla_\rho S^\sigma) (\nabla_\sigma T^\mu) + T^\rho S^\sigma \nabla_\rho \nabla_\sigma T^\mu \\ &= (T^\rho \nabla_\rho S^\sigma) (\nabla_\sigma T^\mu) + T^\rho S^\sigma (\nabla_\sigma \nabla_\rho T^\mu + R^\mu{}_{\nu\rho\sigma} T^\nu) \\ &= (T^\rho \nabla_\rho S^\sigma) (\nabla_\sigma T^\mu) + S^\sigma \nabla_\sigma (T^\rho \nabla_\rho T^\mu) - (S^\sigma \nabla_\sigma T^\rho) \nabla_\rho T^\mu + R^\mu{}_{\nu\rho\sigma} T^\nu T^\rho S^\sigma \\ &= R^\mu{}_{\nu\rho\sigma} T^\nu T^\rho S^\sigma. \end{aligned}$$

In the last step, the first term and the third term cancels out under the change of indices $\rho \leftrightarrow \sigma$ on the third term, and the second term vanishes because $T^\rho \nabla_\rho T^\mu = 0$, for T^μ is the tangent vectors to a geodesic, which parallel-transport its tangent vector. The result,

$$A^\mu = \frac{D^2}{dt^2} S^\mu = R^\mu{}_{\nu\rho\sigma} T^\nu T^\rho S^\sigma$$

is called the ***geodesic deviation equation***. It says that the relative acceleration between two neighboring geodesics is proportional to the curvature.

Chapter 4

Einstein's Field Equation

The study of gravity falls naturally into two pieces: how the gravitational field influences the behavior of matter, and how matter determines the gravitational field. In Newtonian gravity, these two elements consists of the expression for acceleration of a body in a gravitational potential ϕ , and Poisson's differential equation for potentials in terms of matter density ρ and Newton's gravitational constant G :

$$\mathbf{a} = -\nabla\phi \quad (4.1)$$

$$\nabla^2\phi = 4\pi G\rho. \quad (4.2)$$

In general relativity, the analogous statements will describe how the curvature of spacetime acts on matter to manifest itself as gravity, and how energy and momentum influence spacetime to create curvature.

In the previous chapter, we motivated our discussion of manifolds by introducing EEP: "In small enough regions of spacetime, the laws of physics reduce to those of special relativity; it is impossible to detect the existence of a gravitational field by local experiments." The EEP arises from the idea that gravity is universal; it affects all particles in the same way. This feature of universality led Einstein to propose that what we experience as gravity is a manifestation of the curvature of spacetime. The idea is simply that something so universal as gravitation could be most easily described as a fundamental feature of the background on which the matter fields propagate, as opposed to as a conventional force. Also, this is supported by the similarity between the undetectability of gravity in local regions and our ability to find locally inertial coordinates ($g_{\mu\nu}(p) = \eta_{\mu\nu}(p), \partial_\mu g_{\mu\nu} = 0$) on a manifold.

4.1 Physics in Curved Spacetime

The abstract philosophizing above translates directly into a simple recipe for generalizing laws of physics to curved spacetime, which may be stated as follows:

1. Take a law of physics, valid in inertial coordinates in flat spacetime.
2. Write it in a coordinate-invariant (tensorial) form.
3. Assert that the resulting law remains true in curved spacetime.

Operationally, this recipe usually amounts to taking an agreed-upon law in flat space and replacing the Minkowski metric $\eta_{\mu\nu}$ to the more general metric $g_{\mu\nu}$, and replacing partial derivatives ∂_μ by covariant derivatives ∇_μ .

As our first example, we consider the motion of free-falling particles. In flat space such particles move in straight lines; in equations, this is expressed as the vanishing of the second derivative of the parametrized path $x^\mu(\lambda)$:

$$\frac{d^2 x^\mu}{d\lambda^2} = 0.$$

This is not, in general, a tensor equation. Although $dx^\mu/d\lambda$ are the components of a vector, the second derivative are not. But we note that

$$\frac{d^2 x^\mu}{d\lambda^2} = \frac{dx^\nu}{d\lambda} \partial_\nu \frac{dx^\mu}{d\lambda}.$$

Therefore the generalization of the equation above in curved spacetime is simply

$$\frac{dx^\nu}{d\lambda} \nabla_\nu \frac{dx^\mu}{d\lambda} = 0.$$

We see that this is the geodesic equation:

$$\frac{d^2 x^\mu}{d\lambda^2} + \Gamma_{\rho\sigma}^\mu \frac{dx^\rho}{d\lambda} \frac{dx^\sigma}{d\lambda} = 0.$$

In general relativity, therefore, free particles move along geodesics. We have mentioned this before, but now you have a better idea why it is true. Another easy example is the law of energy-momentum conservation in flat spacetime:

$$\partial_\mu T^{\mu\nu} = 0.$$

We can easily generalize this to the following law in curved spacetime:

$$\nabla_\mu T^{\mu\nu} = 0.$$

It is one thing to generalize an equation from flat to curved spacetime; it is something altogether different to argue that the result describes gravity. To do so, we can show how the usual results of Newtonian gravity fits into the picture. We can define the Newtonian limit by three requirement:

- The particles are moving slowly (with respect to the speed of light), that is $\frac{dx^i}{d\tau} \ll \frac{dt}{d\tau}$;
- The gravitational field is weak: So that it can be considered as a perturbation of flat space, or we can write the metric as $g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu}$ where $|h_{\mu\nu}| \ll 1$;
- The field is static, i.e. not changing with time. That is, $\partial_0 g_{\mu\nu} = 0$.

Since the particle is moving slowly, the geodesic equation becomes

$$\frac{d^2 x^\mu}{d\tau^2} + \Gamma_{00}^\mu \left(\frac{dt}{d\tau} \right)^2 = 0. \quad (4.3)$$

Since the field is static ($\partial_0 g_{\mu\nu} = 0$), the Christoffel symbol Γ_{00}^μ simplify:

$$\Gamma_{00}^\mu = \frac{1}{2} g^{\mu\lambda} (\partial_0 g_{\lambda 0} + \partial_0 g_{0\lambda} - \partial_\lambda g_{00}) = -\frac{1}{2} g^{\mu\lambda} \partial_\lambda g_{00}.$$

And since the metric is of the form

$$g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu} \quad |h_{\mu\nu}| \ll 1$$

We can easily compute ¹ that the inverse metric is of the form

$$g^{\mu\nu} = \eta^{\mu\nu} - h^{\mu\nu}$$

to the first order in h . Here $h^{\mu\nu} = \eta^{\mu\rho} \eta^{\nu\sigma} h_{\rho\sigma}$. Putting it all together, to first order in h we find that

$$\Gamma_{00}^\mu = -\frac{1}{2} \eta^{\mu\lambda} \partial_\lambda h_{00}.$$

The geodesic equation (4.3) therefore becomes

$$\frac{d^2 x^\mu}{d\tau^2} = \frac{1}{2} \eta^{\mu\lambda} \partial_\lambda h_{00} \left(\frac{dt}{d\tau} \right)^2.$$

Using $\partial_0 h_{00} = 0$, the $\mu = 0$ component is just

$$\frac{d^2 t}{d\tau^2} = 0.$$

That is, $dt/d\tau$ is a constant. This makes sense since $dt/d\tau = (1 - v^2)^{-1/2}$ with $v^2 \ll 1$. To examine the spacelike components of the geodesic equation, recall that the spacelike component of $\eta^{\mu\nu}$ are just those of a 3×3 identity matrix. Thus

$$\frac{d^2 x^i}{d\tau^2} = \frac{1}{2} \left(\frac{dt}{d\tau} \right)^2 \partial_i h_{00}.$$

Dividing both sides by $(dt/d\tau)^2$ gives

$$\frac{d^2 x^i}{dt^2} = \frac{1}{2} \partial_i h_{00}.$$

This begins to look a great deal like Newton's theory of gravitation. Notice that we can recover (3.1) if we make

$$h_{00} = -2\phi \quad g_{00} = -(1 + 2\phi). \quad (4.4)$$

¹Let $\delta(\cdot)$ denote an infinitesimal variation of something. Note that $\delta(g^{-1}g) = \delta(g^{-1})g + g^{-1}\delta g = 0$. This gives that

$$\delta(g^{-1}) = -g^{-1}(\delta g)g^{-1}.$$

Now since $\delta g_{\mu\mu} = h_{\mu\nu}$, and $g^{\mu\nu} = \eta^{\mu\nu} + \delta g^{\mu\nu}$, substitute everything into the equation above gives

$$\delta g^{\mu\nu} = -(\eta^{\mu\alpha} + \delta g^{\mu\alpha}) h_{\alpha\beta} (\eta^{\beta\nu} + \delta g^{\beta\nu}) = -\eta^{\mu\alpha} h_{\alpha\beta} \eta^{\beta\nu} = -h^{\mu\nu}$$

to the first order. Therefore $g^{\mu\nu} = \eta^{\mu\nu} - h^{\mu\nu}$.

4.2 Einstein's Equation

Now let us generalize the Poisson's equation (3.2) to the curved spacetime, which is known as the Einstein's field equation. In (3.2), we have $\nabla^2 = \delta^{ij}\partial_i\partial_j$ is the Laplacian in space and ρ the mass density. We now informally guess what the Einstein field equation looks like (in fact this is more or less what Einstein did). When we were discussing energy-momentum tensor, we said that the RHS of the Poisson's equation must be replaced by a term involving $T_{\mu\nu}$. And the LHS of Poisson's equation should be the second derivative of a potential-like object. We now guess that this potential like object should be replaced by the metric $g_{\mu\nu}$ of spacetime, since in (4.4), we have related ϕ with a component of $g_{\mu\nu}$. The field equation then, informally, should take the form

$$[\nabla^2 g]_{\mu\nu} \propto T_{\mu\nu}.$$

So what exactly should be on the LHS? One may consider $\square g_{\mu\nu} = \nabla^\lambda \nabla_\lambda g_{\mu\nu}$. But this is automatically zero by metric compatibility. Thus $T_{\mu\nu}$ would be zero, so no. A second guess is that since $R_{\sigma\mu\nu}^\rho$ is constructed from the Christoffel symbols and their derivatives, and the Christoffel symbols are constructed from the metric and its first derivatives, so Riemann curvature tensor contains the second derivatives of $g_{\mu\nu}$. It doesn't have the right number of indices, but we can contract it to form the Ricci tensor $R_{\mu\nu}$, which does, and this is symmetric, just like $T_{\mu\nu}$. So it is tempting to guess that our field equation is

$$R_{\mu\nu} = \kappa T_{\mu\nu}$$

for some constant κ . But conservation of energy requires

$$\nabla^\mu T_{\mu\nu} = 0.$$

But Bianchi's identity (3.11) gives

$$\nabla^\mu R_{\mu\nu} = \frac{1}{2}\nabla_\nu R.$$

But our proposed field equation implies $R = \kappa g^{\mu\nu} T_{\mu\nu} = \kappa T$. So we have

$$\nabla_\mu T = \partial_\mu T = 0.$$

So this would imply T is constant throughout spacetime. This is highly unlikely, since $T \neq 0$ in matter. We have to try harder. A third guess: Einstein's tensor

$$G_{\mu\nu} = R_{\mu\nu} - \frac{1}{2}Rg_{\mu\nu}$$

contains the second derivative of the metric, has the right index, and always obeys $\nabla^\mu G_{\mu\nu} = 0$, as we discussed. Therefore, we propose

$$G_{\mu\nu} = \kappa T_{\mu\nu}. \tag{4.5}$$

Now it remains to fix the constant κ , and see whether the result reproduces gravity as we know it. Notice that contracting both sides of (4.5) gives

$$R = -\kappa T.$$

Using this we can rewrite (4.5) as

$$R_{\mu\nu} = \kappa(T_{\mu\nu} - \frac{1}{2}Tg_{\mu\nu}). \quad (4.6)$$

This is the same equation as (4.5), we just write it this way so that it simplifies our following computations a little bit. We would like to see if it predicts Newtonian gravity in the weak-field, time-independent, slowly-moving-particles limit. To this end we consider a perfect fluid source of energy-momentum, which are used to model idealized distributions of matter, such as the interior of a star or an isotropic universe. The energy momentum tensor is therefore

$$T_{\mu\nu} = (\rho + p)U_\mu U_\nu + pg_{\mu\nu}.$$

Here U^μ is the 4-velocity and ρ, p are the rest-frame energy and momentum densities. If we fix the unit by plugging in c , we see that $T_{\mu\nu} = \rho c^2 U_\mu U_\nu + \dots$, and clearly the ρ term dominates as we are considering the slowly-moving-particle limit. So it suffices to consider the energy-momentum tensor of dust:

$$T_{\mu\nu} = \rho U_\mu U_\nu.$$

Thus the "fluid" we are considering is some massive body, such as the Earth or the Sun. We will work in the fluid rest frame, in which

$$U^\mu = (U^0, 0, 0, 0).$$

The timelike component can be fixed by appealing to the normalization condition $g_{\mu\nu}U^\mu U^\nu = g_{00}(U^0)^2 = -1$ where the weak-field limit implies

$$\begin{aligned} g_{00} &= -1 + h_{00} \\ g^{00} &= -1 - h_{00} \end{aligned}$$

as we discussed before. Then to the first order of $h_{\mu\nu}$ we get

$$U^0 = 1 + \frac{1}{2}h_{00}.$$

Then we have

$$T_{00} = \rho U_0 U_0 = \rho \left(-1 + \frac{1}{2}h_{00} \right)^2 \approx \rho(1 + h_{00}) \approx \rho$$

and since $U_0 = -U^0$, we have

$$T = g^{\mu\nu}T_{\mu\nu} = g^{\mu\nu}\rho U_\mu U_\nu = \rho U^\mu U_\mu = \rho U^0 U_0 = -\rho \left(1 + \frac{1}{2}h_{00} \right)^2 \approx -\rho(1+h_{00}) \approx -\rho.$$

Then we plug this into the 00 component of (4.6), and we get

$$R_{00} = \frac{1}{2}\kappa\rho.$$

To find the explicit expression involving metric, we need to evaluate $R_{00} = R_{0\lambda 0}^\lambda$. In fact we only need R_{0i0}^i , for $R_{000}^0 = 0$ by antisymmetry.

$$R_{0j0}^i = \partial_j \Gamma_{00}^i - \partial_0 \Gamma_{j0}^i + \Gamma_{j\lambda}^i \Gamma_{00}^\lambda - \Gamma_{0\lambda}^i \Gamma_{j0}^\lambda.$$

The second term is a time derivative, which vanishes for static fields. The third and fourth terms are of the form $\Gamma^2 \approx O(h^2)$, thus can be neglected. We are left with $R_{0i0}^i = \partial_j \Gamma_{00}^i$, from which we have

$$R_{00} = R_{0i0}^i = \partial_i \left[\frac{1}{2} g^{i\lambda} (\partial_0 g_{\lambda 0} + \partial_0 g_{0\lambda} - \partial_\lambda g_{00}) \right] = -\frac{1}{2} \delta^{ij} \partial_i \partial_j h_{00} = -\frac{1}{2} \nabla^2 h_{00}.$$

And this is equal to $\frac{1}{2} \kappa \rho$. We see that in the Newtonian limit the field equation reduces to

$$\nabla^2 h_{00} = -\kappa \rho.$$

Since in (4.4) we set $h_{00} = -2\phi$, it seems that if we set

$$\kappa = 8\pi G$$

then we can reproduce Poisson's equation (3.2). Therefore, the ***Einstein's field equation*** is given by $G_{\mu\nu} = 8\pi G T_{\mu\nu}$, or

$$R_{\mu\nu} - \frac{1}{2} R g_{\mu\nu} = 8\pi G T_{\mu\nu}.$$

You may think, is this really it? Well, we will derive this equation again by an action principle later, which is more rigorous. But it is good to know how to guess. But this is not the end of the story. You can add any divergence free tensor to the LHS of the equation above and still have a good field equation. In particular we can add $g_{\mu\nu}$ to the LHS. But the unit is a little bit off so we have to insert a constant Λ to make everything comes out right. That is,

$$G_{\mu\nu} + \Lambda g_{\mu\nu} = 8\pi T_{\mu\nu}.$$

This Λ is called the cosmological constant. But we can also move the term $\Lambda g_{\mu\nu}$ to the RHS and think of this term as an additional source of energy. So let us define

$$T_{\mu\nu}^\Lambda = -\frac{\Lambda}{8\pi G} g_{\mu\nu}.$$

Then we have $G_{\mu\nu} = 8\pi G (T_{\mu\nu} + T_{\mu\nu}^\Lambda)$, and $T_{\mu\nu}^\Lambda$ is just a perfect fluid with

$$\rho = \frac{\Lambda}{8\pi G} \quad p = -\frac{\Lambda}{8\pi G}.$$

Such a energy-momentum tensor actually arises from quantum field theories. This represents a form of zero point energy.

4.3 *Einstein-Hilbert Action

4.4 Linearized Gravity

In the previous section we derived Einstein's field equation (EFE), which has the form

$$G_{\alpha\beta} = 8\pi G T_{\alpha\beta}.$$

This is a set of differential equations for the spacetime given a source². Unfortunately, this is hard to solve, as we can think of $G_{\alpha\beta} = \mathcal{D}_{\alpha\beta}^2(g_{\mu\nu})$, where $\mathcal{D}_{\alpha\beta}^2$ is some very complicated second order nonlinear operator. Even with the perfect fluid source $T_{\alpha\beta} = (\rho + p)U_\alpha U_\beta + pg_{\alpha\beta}$, the EFE can be extremely hard to solve in general. We also have the constraint of normalization $g_{\mu\nu}U^\mu U^\nu = -1$ to satisfy. There are three major ways we can solve EFE, which we list below:

1. We consider weak field limit, where $g_{\alpha\beta} \sim \eta_{\alpha\beta}$ so we can work in nearly flat space, which is the topic of this chapter.
2. We consider symmetric solutions, which is considered in [10]
3. We solve EFE numerically. This will not be discussed in this notes.

In linearized gravity theory, we are assuming the gravity is weak such that the metric takes the following form:

$$g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu} \quad |h_{\mu\nu}| \ll 1$$

and $\eta_{\mu\nu} = \text{diag}(-1, 1, 1, 1)$ is in the canonical form. This gives us the freedom to throw away all terms of the form $O(h^2)$ since they are small. We discussed in the last chapter that the inverse metric takes the form

$$g^{\mu\nu} = \eta^{\mu\nu} + h^{\mu\nu}.$$

We first discuss several properties of such "nearly Lorentz" spaces. We consider a coordinate transformation:

$$g_{\mu'\nu'} = L^\alpha_{\mu'} L^\beta_{\nu'} g_{\alpha\beta} \quad L^\alpha_{\mu'} = \frac{\partial x^\alpha}{\partial x^{\mu'}}.$$

In flat space, the Lorentz transformation $\Lambda^\alpha_{\mu'}$ plays a special role. So let us see (just for fun) how the Lorentz transformation acts on the metric. We apply Lorentz transformation in nearly flat spacetime, we get

$$\begin{aligned} g_{\mu'\nu'} &= \Lambda^\alpha_{\mu'} \Lambda^\beta_{\nu'} (\eta_{\alpha\beta} + h_{\alpha\beta}) \\ &= \Lambda^\alpha_{\mu'} \Lambda^\beta_{\nu'} \eta_{\alpha\beta} + \Lambda^\alpha_{\mu'} \Lambda^\beta_{\nu'} h_{\alpha\beta} \\ &= \eta_{\mu'\nu'} + \Lambda^\alpha_{\mu'} \Lambda^\beta_{\nu'} h_{\alpha\beta} \\ &= \eta_{\mu'\nu'} + h_{\mu'\nu'}. \end{aligned}$$

So the Lorentz transformation applied to the nearly flat spacetime leaves the metric of the flat space unchanged, and the perturbation of the background transforms just as any tensor would transform in flat space. But we are not working in flat space: Two geodesics that are initially parallel may not remain parallel. But in many way it's close enough that we can borrow many mathematical tools used in flat spacetime. In particular, in this framework, we can regard $h_{\alpha\beta}$ as an ordinary tensor field living in the $\eta_{\alpha\beta}$ metric. Fundamentally it is not, but we can imagine this sometimes, as many mathematical tools just work perfectly in this fiction.

²By source we mean the form of the energy-momentum tensor $T_{\alpha\beta}$ will be given. However, we will need to solve $T_{\alpha\beta}$ and the metric together.

Another thing to notice is how we raise indices:

$$h^{\alpha\beta} = g^{\alpha\mu} g^{\beta\nu} h_{\mu\nu} = \eta^{\alpha\mu} \eta^{\beta\nu} h_{\mu\nu}$$

to the first order of h . One final detail comes when we consider a coordinate shift

$$x^{\alpha'} = x^\alpha + \xi^\alpha(x^\beta)$$

where $\xi^\alpha(x^\beta)$ is a little offset that is a function of coordinates. Then our coordinate tranformation matrix will be

$$L^{\alpha'}_{\beta} = \frac{\partial x^{\alpha'}}{\partial x^\beta} = \delta^{\alpha'}_{\beta} + \partial_\beta \xi^\alpha.$$

We require

$$|\partial_\beta \xi^\alpha| \ll 1.$$

So the transformation above is an infinitesimal coordinate transformation. Then it is not hard to recover from the relation $L^{\alpha'}_{\beta} L^{\beta}_{\gamma'} = \delta^{\alpha'}_{\gamma'}$ that

$$L^{\beta}_{\gamma'} = \delta^{\beta}_{\gamma'} - \partial_{\gamma'} \xi^\beta.$$

Then

$$\begin{aligned} g_{\mu'\nu'} &= L^{\alpha}_{\mu'} L^{\beta}_{\nu'} g_{\alpha\beta} = (\delta^{\alpha}_{\mu} - \partial_{\mu'} \xi^\alpha)(\delta^{\beta}_{\nu} - \partial_{\nu'} \xi^\beta)(\eta_{\alpha\beta} + h_{\alpha\beta}) \\ &= \eta_{\mu\nu} + h_{\mu\nu} - \partial_{\mu'} \xi_\nu - \partial_{\nu'} \xi_\mu + O(h \cdot \partial \xi) \\ &= \eta_{\mu'\nu'} + h_{\mu'\nu'}. \end{aligned}$$

Then we see that $\eta_{\mu'\nu'} = \eta_{\mu\nu}$ and

$$h_{\mu'\nu'} = h_{\mu\nu} - \partial_{\mu'} \xi_\nu - \partial_{\nu'} \xi_\mu. \quad (4.7)$$

We call an infinitesimal coordinate transformation $x^\alpha \rightarrow x^\alpha + \xi^\alpha$ of this kind a ***gauge transformation*** in linearized gravity. We want to find the equations of motion obeyed by perturbations $h_{\mu\nu}$, which come by examining EFE to the first order. The Christoffel symbols are given by

$$\begin{aligned} \Gamma^{\rho}_{\mu\nu} &= \frac{1}{2} g^{\rho\lambda} (\partial_\mu g_{\nu\lambda} + \partial_\nu g_{\lambda\mu} - \partial_\lambda g_{\mu\nu}) \\ &= \frac{1}{2} \eta^{\rho\lambda} (\partial_\mu h_{\nu\lambda} + \partial_\nu h_{\lambda\mu} - \partial_\lambda h_{\mu\nu}). \end{aligned}$$

Since the connection coefficients are first-order quantities, the only contribution to the Riemann tensor will come from the derivatives of Γ 's, not $O(\Gamma^2)$. Lowering an index for convenience, we obtain

$$R_{\mu\alpha\nu\beta} = \frac{1}{2} (\partial_\alpha \partial_\nu h_{\mu\beta} + \partial_\mu \partial_\beta h_{\alpha\nu} - \partial_\alpha \partial_\beta h_{\mu\nu} - \partial_\mu \partial_\nu h_{\alpha\beta}).$$

The Ricci tensor comes from contracting over μ and ρ , giving

$$R_{\alpha\beta} = \eta^{\mu\nu} R_{\mu\alpha\gamma\beta} = \frac{1}{2} (\partial_\alpha \partial^\mu h_{\mu\beta} + \partial_\beta \partial^\mu h_{\mu\alpha} - \partial_\alpha \partial_\beta h - \square h_{\alpha\beta})$$

which is manifestly symmetric in α and β . In this expression we defined $h = \eta^{\mu\nu} h_{\mu\nu} = h^\mu_\mu$, and the d'Alembertian is simply the one from flat space,

$$\square = \eta^{\mu\nu} \partial_\mu \partial_\nu = -\partial_t^2 + \partial_x^2 + \partial_y^2 + \partial_z^2.$$

Contracting again gives

$$R = \partial_\mu \partial_\nu h^{\mu\nu} - \square h.$$

Notice that the gauge transformation leaves the linearized Riemann tensor unchanged:

$$\begin{aligned} \delta R_{\mu\nu\rho\sigma} &= \frac{1}{2} (\partial_\rho \partial_\nu \partial_\mu \xi_\sigma + \partial_\rho \partial_\nu \partial_\sigma \xi_\mu + \partial_\sigma \partial_\mu \partial_\nu \xi_\rho + \partial_\sigma \partial_\mu \partial_\rho \xi_\nu \\ &\quad - \partial_\sigma \partial_\nu \partial_\mu \xi_\rho - \partial_\sigma \partial_\nu \partial_\rho \xi_\mu - \partial_\rho \partial_\mu \partial_\nu \xi_\sigma - \partial_\rho \partial_\mu \partial_\sigma \xi_\nu) \\ &= 0. \end{aligned}$$

Putting it all together we obtain the Einstein tensor:

$$\begin{aligned} G_{\alpha\beta} &= R_{\alpha\beta} - \frac{1}{2} \eta_{\alpha\beta} R \\ &= \frac{1}{2} (\partial_\alpha \partial^\mu h_{\mu\beta} + \partial_\beta \partial^\mu h_{\mu\alpha} - \partial_\alpha \partial_\beta h - \square h_{\alpha\beta} + \eta_{\alpha\beta} \square h - \eta_{\alpha\beta} \partial^\mu \partial^\nu h_{\mu\nu}). \end{aligned}$$

This is already a mess. So let us clean up a little bit. We define

$$\bar{h}_{\alpha\beta} = h_{\alpha\beta} - \frac{1}{2} \eta_{\alpha\beta} h.$$

Then the trace is given by

$$\bar{h} = \eta^{\alpha\beta} \bar{h}_{\alpha\beta} = h - \frac{1}{2} h (\eta^{\alpha\beta} \eta_{\alpha\beta}) = -h.$$

Therefore, we call $\bar{h}_{\alpha\beta}$ the **trace-reversed metric perturbation**. This is useful because if we insert $h_{\alpha\beta} = \bar{h}_{\alpha\beta} + \frac{1}{2} \eta_{\alpha\beta} h$, we see that the Einstein tensor simplifies to

$$G_{\alpha\beta} = \frac{1}{2} (\partial_\alpha \partial^\mu \bar{h}_{\mu\beta} + \partial_\beta \partial^\mu \bar{h}_{\mu\alpha} - \eta_{\alpha\beta} \partial^\mu \partial^\nu \bar{h}_{\mu\nu} - \square \bar{h}_{\alpha\beta}).$$

In the above expression, the first three terms contains the form $\partial^\mu \bar{h}_{\mu\nu}$, and if we can set $\partial^\mu \bar{h}_{\mu\nu} = 0$, which is equivalent to divergence free condition $\partial_\mu \bar{h}^{\mu\nu} = 0$ as you may check, then the Einstein tensor will be simplified into only one term. These are 4 conditions to satisfy, and the gauge generators ξ_ν are 4 free functions. This suggest we can do it. Given (4.7), it is not hard to arrive at

$$\bar{h}_{\mu'\nu'} = \bar{h}_{\mu\nu} - \partial_\mu \xi_\nu - \partial_\nu \xi_\mu + \eta_{\mu\nu} \partial^\alpha \xi_\alpha.$$

This suggests that

$$\partial^{\mu'} \bar{h}_{\mu'\nu'} = \partial^\mu \bar{h}_{\mu\nu} - \square \xi_\nu - \partial^\mu \partial_\nu \xi_\mu + \partial_\nu \partial^\alpha \xi_\alpha = \partial^\mu \bar{h}_{\mu\nu} - \square \xi_\nu$$

where the last two terms cancel. Therefore, if we choose our gauge generator just right, we can adjust the perturbed metric such that $\partial^{\mu'} \bar{h}_{\mu'\nu'} = 0$. In particular, this can be done if we choose

$$\square \xi_\nu = \partial^\mu \bar{h}_{\mu\nu}. \quad (4.8)$$

But (4.8) is a simple wave equation, whose solution is guaranteed to exist. The gauge that satisfies (4.8) is called the Lorentz gauge in linearized gravity. Therefore the Einstein tensor in this frame becomes

$$G_{\alpha\beta} = -\frac{1}{2} \square \bar{h}_{\alpha\beta}.$$

Thus the EFE is simply

$$\square \bar{h}_{\alpha\beta} = -16\pi G T_{\alpha\beta}. \quad (4.9)$$

(4.9) can be solved with radiative Green's function. We will do this in the next section. But now let us look at the solution to (4.9) in Newtonian limit. This means that the source to be a static nonrelativistic perfect fluid. Then $\partial_0 h = 0$, and $\rho \gg p$, therefore we can write

$$T_{\alpha\beta} = \rho U_\alpha U_\beta \sim \begin{pmatrix} \rho & & & \\ & 0 & & \\ & & 0 & \\ & & & 0 \end{pmatrix}.$$

Then our equation now becomes

$$\nabla^2 \bar{h}_{00} = -16\pi G \rho$$

where $\square = \nabla^2$ simply because $\partial_0 h_{00} = 0$. But since the Newtonian potential ϕ solves $\nabla^2 \phi = 4\pi G \rho$, we immediately have

$$\bar{h}_{00} = -4\phi$$

and all the remaining $\bar{h}_{\mu\nu} = 0$. Then since

$$h_{\mu\nu} = \bar{h}_{\mu\nu} - \frac{1}{2}\eta_{\mu\nu}\bar{h},$$

and we have $\bar{h} = \eta^{\mu\nu}\bar{h}_{\mu\nu} = -\bar{h}_{00} = 4\phi$, this leads to

$$h_{\mu\nu} = \begin{pmatrix} -2\phi & & & \\ & -2\phi & & \\ & & -2\phi & \\ & & & -2\phi \end{pmatrix}.$$

Therefore the perturbed metric for static Newtonian sources is given by

$$\boxed{ds^2 = -(1 + 2\phi)dt^2 + (1 - 2\phi)(dx^2 + dy^2 + dz^2)}. \quad (4.10)$$

Part II

Quantum Mechanics

Chapter 5

Quantum Theory and Representations

5.1 Introduction and Overview

5.1.1 Axioms of Quantum Mechanics

- The state of a quantum mechanical system is given by a nonzero vector in a complex vector space $\mathcal{H}$ with Hermitian inner product $\langle \cdot, \cdot \rangle$. States are usually denoted by $|\psi\rangle \in \mathcal{H}$.
- The observables of a quantum mechanical system are given by self-adjoint linear operators on $\mathcal{H}$.
- Time evolution of states $|\psi(t)\rangle \in \mathcal{H}$ is given by the Schrödinger equation

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = H |\psi(t)\rangle$$

where H is a distinguished observable called the Hamiltonian. H has eigenvalues that are bounded below.

- States for which the value of an observable can be characterized by a well-defined number are the eigenstates for the corresponding self-adjoint operator. The value of the observable in such a state will be a real number, the eigenvalue of the operator.
- Given an observable O and two unit-norm states $|\psi_1\rangle, |\psi_2\rangle$ that are eigenstates of O with distinct eigenvalues λ_1, λ_2 , the complex linear combination

$$c_1 |\psi_1\rangle + c_2 |\psi_2\rangle$$

will not have a well-defined value for the observable O . If one attempts to measure this observable, one will get either λ_1 or λ_2 , with probability $\frac{|c_1|^2}{|c_1|^2 + |c_2|^2}$ and $\frac{|c_2|^2}{|c_1|^2 + |c_2|^2}$ respectively.

5.1.2 Group Action and Representation

Given a group G acting on a space M , the space of all functions $F(M)$ on M comes with a natural G -action: Let $f \in F(M), g \in G, x \in M$, this action is

defined by

$$(g \cdot f)(x) = f(g^{-1} \cdot x).$$

One can verify that this is a group action.

Definition 5.1. A **representation** (π, V) if a group G is a homomorphism

$$\pi : G \rightarrow GL(V).$$

This means that $\pi(gh) = \pi(g)\pi(h) \in GL(V)$.

The trivial representation is given by $\pi(g) = I \in GL(V)$ for all $g \in G$.

Definition 5.2. Let (π, V) be a representation on a vector space V with Hermitian inner product. We say π is **unitary** if it preserves inner product, i.e.,

$$\langle \pi(g)v, \pi(g)w \rangle = \langle v, w \rangle$$

for all $g \in G, v, w \in V$.

5.1.3 Representations and Quantum Mechanics

The big idea is: whenever we have a physical quantum system with a group G acting on it, the state space $\mathcal{H}$ will carry a unitary representation of G . To see this intuitively (this will be discussed in greater detail later), for a representation π and $g \in G$ that are close to identity, exponentiation can be used to write

$$\pi(g) = e^A$$

for some A close to the zero matrix (in fact what we have is $\pi(e^{tX}) = e^{t\pi'(X)}$ where π' is the differential map of π). It turns out that if $\pi(g)$ is unitary, then A will be skew-adjoint:

$$A^\dagger = -A.$$

Taking $B = iA$ we find

$$B^\dagger = B.$$

Therefore B is a self-adjoint operator that corresponds to some observable of the quantum system, which we obtain purely from the representation π itself! On the other hand, knowing the observable B will give us A , which we may then exponentiate to reconstruct the representation $\pi(g)$.

5.2 $U(1)$ and its Representations

5.2.1 Irreducible Representations

Definition 5.3. A representation π is called **irreducible** if it has no subrepresentation, i.e., no nonzero proper subspaces $W \subset V$ such that $(\pi|_W, W)$ is a representation.

Definition 5.4. Given $(\pi_1, V), (\pi_2, W)$ representations, we can define the direct sum representation $(\pi_1 \oplus \pi_2, V \oplus W)$, which is given by

$$g \in G \xrightarrow{\pi} \pi(g) = \begin{pmatrix} \pi_1(g) & 0 \\ 0 & \pi_2(g) \end{pmatrix}.$$

Theorem 5.5. *Any unitary representation can be written as a direct sum*

$$\pi = \pi_1 \oplus \cdots \oplus \pi_n$$

where each π_j is irreducible.

Theorem 5.6 (Schur's Lemma). *If (π, V) is irreducible, then the only linear map $T : V \rightarrow V$ commuting with all $\pi(g)$ are of the form $T = \lambda I$, for some $\lambda \in \mathbf{C}$.*

Theorem 5.7. *All irreducible representations of an Abelian group is one-dimensional.*

5.2.2 $U(1)$ and Its Representation

Definition 5.8. The group $U(1)$ is the unit circle in $\mathbf{C}$, i.e.,

$$U(1) = \{e^{i\theta} : \theta \in \mathbf{R}\}.$$

The group multiplication is the usual complex multiplication.

By previous theorem, all irreducible representations of $U(1)$, which is abelian, will be 1-dimensional. Thus they will be given by differentiable maps

$$\pi : U(1) \rightarrow GL(1, \mathbf{C}).$$

Theorem 5.9. *All irreducible representations of the group $U(1)$ are unitary and given by*

$$\pi_k : e^{i\theta} \mapsto e^{ik\theta}$$

where $k \in \mathbf{Z}$.

Proof. See [3] pg 20-21. □

5.2.3 The Charge Operator

By the previous theorems, for a general $U(1)$ representation, we can decompose the state space as

$$\mathcal{H} = \mathcal{H}_{q_1} \oplus \cdots \oplus \mathcal{H}_{q_n}$$

for some integers $q_1, \dots, q_n$, where $n = \dim \mathcal{H}$ and each $\mathcal{H}_{q_j}$ is a copy of $\mathbf{C}$, with $U(1)$ acting by the π_{q_j} representation.

Choosing a basis element for each $\mathcal{H}_{q_j}$, define the charge operator by

$$Q = \begin{pmatrix} q_1 & 0 & \cdots & 0 \\ 0 & q_2 & \cdots & 0 \\ \cdots & & & \cdots \\ 0 & 0 & \cdots & q_n \end{pmatrix}.$$

Since $\pi_N(e^{i\theta}) = e^{iN\theta}$, through some calculation we find its derivative to be $\pi'_N(\theta) = iN\theta$. We mentioned previously that the representation can be recovered by exponentiating the derivatives. Therefore, this motivates us to find that

$$\pi(e^{i\theta}) = e^{iQ\theta} = \begin{pmatrix} e^{iq_1\theta} & 0 & \cdots & 0 \\ 0 & e^{iq_2\theta} & \cdots & 0 \\ \cdots & & & \cdots \\ 0 & 0 & \cdots & e^{iq_n\theta} \end{pmatrix} \in U(n) \subset GL(n, \mathbf{C}).$$

The standard physics terminology is that Q is the generator of the $U(1)$ action by unitary transformations on the state space $\mathcal{H}$. In fact one can further see the relation by calculating the derivative of π , and in this case

$$\begin{aligned}\pi' : T_1 U(1) = i\mathbf{R} &\rightarrow T_{\pi(1)} GL(n, \mathbf{C}) = M(n, \mathbf{C}), \\ i\theta &\mapsto iQ\theta.\end{aligned}$$

5.2.4 Conservation of Charge and $U(1)$ Symmetry

Recall that the evolution of a state $|\psi(t)\rangle$ is determined by the Schrödinger equation

$$\frac{d}{dt} |\psi(t)\rangle = -iH |\psi(0)\rangle$$

where we set $\hbar = 1$. The solution to this equation is given by the propagator operator $U(t)$ where

$$U(t) = e^{-itH}$$

and

$$|\psi(t)\rangle = U(t) |\psi(0)\rangle.$$

Suppose now

$$[Q, H] = 0,$$

then we also have $[Q, U(t)] = 0$ for all t . Therefore, suppose a state has a well-defined charge q at $t = 0$, i.e., $Q |\psi(0)\rangle = q |\psi(0)\rangle$ then it will also have a well-defined charge at any time t , which is again q , for

$$Q |\psi(t)\rangle = QU(t) |\psi(0)\rangle = U(t)Q |\psi(0)\rangle = qU(t) |\psi(0)\rangle = q |\psi(t)\rangle.$$

Therefore, when $[Q, H] = 0$, the charge Q is conserved. Then the group $U(1)$ is said to act as a symmetry group of the system, with $\pi(e^{i\theta})$ the symmetry transformations.

5.3 Lie Groups and Lie Algebras

5.3.1 Lie Groups

Definition 5.10. A group G is a **Lie group** if it is also a differentiable manifold.

Definition 5.11. Here are some important Lie groups:

1. $GL(n, \mathbf{F}) = \{L \in M(n, \mathbf{F}) : \det L \neq 0\}.$
2. $SL(n, \mathbf{F}) = \{L \in GL(n, \mathbf{F}) : \det L = 1\}.$
3. $O(n) = \{L \in GL(n, \mathbf{R}) : L^{-1} = L^T\}.$
4. $SO(n) = \{L \in O(n) : \det L = 1\}.$
5. $U(n) = \{L \in GL(n, \mathbf{C}) : L^{-1} = L^\dagger\}.$
6. $SU(n) = \{L \in U(n) : \det L = 1\}.$

Remark 5.12. Let $L : V \rightarrow V$ be the a linear operator on a complex vector space V that preserves inner product, i.e., for all $v, w \in V$ we have $\langle Lv, Lw \rangle = \langle v, w \rangle$. This is equivalent to $\langle L^\dagger Lv, w \rangle = \langle v, w \rangle$ for all $v, w \in V$, i.e.,

$$L^\dagger = L^{-1}.$$

Therefore, we see that the set of linear maps on V that preserves inner product is given by exactly $U(n)$, up to isomorphism (since we have an isomorphism $V \cong \mathbf{C}^n$). When V is a real vector space, we see that $L^\dagger = L^T$, and the the set of linear maps on V that preserves inner product is given by exactly $O(n)$, again up to isomorphism.

Remark 5.13. With some work one can show that in the case $n = 2$, any matrix element in $SU(2)$ will have the form

$$\begin{pmatrix} \alpha & \beta \\ -\bar{\beta} & \bar{\alpha} \end{pmatrix} \quad |\alpha|^2 + |\beta|^2 = 1, \quad \alpha, \beta \in \mathbf{C}.$$

One can show then that $SU(2)$ is diffeomorphic to the 3-sphere S^3 , by viewing $S^3 \subset \mathbf{R}^4 \cong \mathbf{C}^2$, and use the diffeomorphic map $\varphi : \begin{pmatrix} \alpha & \beta \\ -\bar{\beta} & \bar{\alpha} \end{pmatrix} \mapsto (\alpha, \beta) \in \mathbf{C}^2$.

5.3.2 Lie Algebras

Definition 5.14. Let G be a Lie group and $e \in G$ be the identity element. Then the *Lie algebra* $\mathfrak{g}$ of G is defined to be its tangent space at identity, i.e.,

$$\mathfrak{g} = T_e G.$$

Definition 5.15. When G is a matrix Lie group of $n \times n$ invertible matrices, the Lie algebra $\mathfrak{g}$ of G is, equivalently, the space of $n \times n$ matrices X such that $e^{tX} \in G$ for all $t \in \mathbf{R}$.

Proposition 5.16 (Lie algebras of important Lie groups). The following statements are true:

1. $\mathfrak{gl}(n, \mathbf{R}) = M(n, \mathbf{R})$, the set of all $n \times n$ matrices with entries in $\mathbf{R}$
2. $\mathfrak{sl}(n, \mathbf{F}) = \{X \in M(n, \mathbf{F}) : \text{tr } X = 0\}$.
3. $\mathfrak{o}(n) = \{X \in M(n, \mathbf{R}) : X^T = -X\}$.
4. $\mathfrak{so}(n) = \mathfrak{o}(n)$.
5. $\mathfrak{u}(n) = \{X \in M(n, \mathbf{C}) : X^\dagger = -X\}$.
6. $\mathfrak{su}(n) = \{X \in M(n, \mathbf{C}) : X^\dagger = -X, \text{tr } X = 0\}$.

Proof. To prove (1), notice that for every $X \in M(n, \mathbf{F})$, the matrix $\Omega = e^{tX}$ is invertible, hence in $GL(n, \mathbf{F})$. Thus the claim follows from definition.

To prove (2), notice that e^{tX} will have determinant 1 if and only if $\text{tr}(tX) = t \text{tr}(X) = 0$, which implies $\text{tr } X = 0$. Here we use the formula

$$\det(e^X) = e^{\text{tr } X}.$$

To prove (5), notice that if $\Omega \in U(n)$, then $\Omega\Omega^\dagger = I$. Using $(XY)^T = Y^T X^T$, we can deduce that $(e^{tX})^\dagger = e^{tX^\dagger}$. Now writing $\Omega = e^{tX}$, the condition that $\Omega\Omega^\dagger = I$ becomes

$$e^{tX}(e^{tX})^\dagger = e^{tX}e^{tX^\dagger} = I.$$

Taking the derivative of this equation gives

$$e^{tX}X^\dagger e^{tX^\dagger} + X e^{tX}e^{tX^\dagger} = 0.$$

Evaluating at $t = 0$ gives

$$X + X^\dagger = 0.$$

This proves (5). In the real case, $X^\dagger = X^T$, hence (3) also follows immediately.

(6) is obvious.

To prove (3), note that any skew-symmetric real matrix has vanishing trace. That is, if $X^T = -X$, then clearly $X_{ii} = 0$ for all i , therefore all diagonal elements are zero. This concludes the proof. $\square$

5.3.3 Lie Group and Lie Algebra Representations

Definition 5.17. A Lie algebra representation (ϕ, V) of a Lie algebra on a complex n -dimensional vector space V is given by a real linear map $\phi : \mathfrak{g} \rightarrow \mathfrak{gl}(V) \cong M(n, \mathbb{C})$ such that

$$[\phi(X), \phi(Y)] = \phi([X, Y])$$

for all $X, Y \in \mathfrak{g}$. Such a representation is called unitary if the image of ϕ is in $\mathfrak{u}(n)$, i.e., $\phi(X)^\dagger = -\phi(X)$. We say ϕ is irreducible if there is no proper nonzero subspaces of V that is invariant under ϕ .

We can define the adjoint representation and check that indeed $[X, Y]$ is an well-defined element in $\mathfrak{g}$ for $X, Y \in \mathfrak{g}$. This is done in pg 57-58 in [3]. I will not go through the detail.

Given a Lie group homomorphism $\pi : G \rightarrow GL(V)$, one can consider the differential $\pi' : T_e G \rightarrow T_{\pi(e)} GL(V)$ of this map. This is then a map between Lie algebras: $\pi' : \mathfrak{g} \rightarrow \mathfrak{gl}(V)$. To see explicitly what π' is, we recall the definition of the differential map. For every $X \in T_e G = \mathfrak{g}$, there exists (at least in a small neighborhood of the identity) a unique curve γ_X such that $\gamma_X(0) = e$ and $\gamma'_X(0) = X$. Then the differential map π' takes X to the tangent vector to the curve $\pi(\gamma_X(t))$ at $t = 0$.

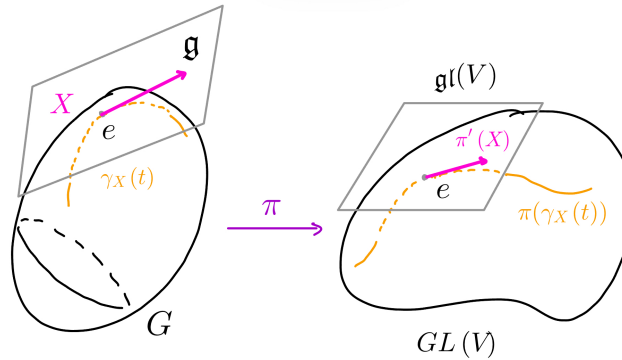


Figure 5.1: An illustration of how π' works.

In the case where G is a matrix group, note that $\gamma(t) = e^{tX}$ is a curve in G with the property that $\gamma(0) = e$ and $\gamma'(0) = X$. Therefore, we may take $\gamma_X(t) = e^{tX}$. Then the differential map can be readily written as

$$\pi'(X) = \frac{d}{dt}\pi(e^{tX})|_{t=0}.$$

Now comes the key result in this section (the proof is given again on page 65 of [3], and I am going to skip it). This theorem tells us that a Lie group representation naturally gives a Lie algebra representation, and one can use this Lie algebra representation to recover the Lie group representation.

Theorem 5.18. *Let $\pi : G \rightarrow GL(n, \mathbf{C})$ be a Lie group homomorphism. Then $\pi' : \mathfrak{g} \rightarrow \mathfrak{gl}(n, \mathbf{C}) = M(n, \mathbf{C})$ given by*

$$\pi'(X) = \frac{d}{dt}(\pi(e^{tX}))|_{t=0}$$

satisfies

1. $\pi(e^{tX}) = e^{t\pi'(X)}.$
2. For $g \in G$, we have

$$\pi'(gXg^{-1}) = \pi(g)\pi'(X)\pi(g)^{-1}.$$

See pg 57 of [3] for a proof that $gXg^{-1} \in \mathfrak{g}$ when $X \in \mathfrak{g}$.

3. $\pi'([X, Y]) = [\pi'(X), \pi'(Y)],$ i.e., π' is a Lie algebra homomorphism.

5.4 Irreducible Representations of $SU(2)$

The goal of this chapter is to classify all irreducible representations of the group $SU(2)$. Recall that this is the group of 2×2 complex matrices with $L^{-1} = L^\dagger$ and determinant 1. We explained that any such matrix can must be of the form

$$\begin{pmatrix} \alpha & \beta \\ -\bar{\beta} & \bar{\alpha} \end{pmatrix} \quad |\alpha|^2 + |\beta|^2 = 1, \quad \alpha, \beta \in \mathbf{C}.$$

In this section, we shall classify all the irreducible representations of $SU(2)$, and construct them explicitly. It will turn out that irreducible representations of $SU(2)$ are classified by a nonnegative integer $n = 0, 1, 2, 3, \dots$, which we will denote by π_n . And each π_n acts on a vector space of dimension $n + 1$.

To classify irreducible representations of $SU(2)$, we first linearize the problem by classifying irreducible representations of its Lie algebra $\mathfrak{su}(2)$, and then we exponentiate these irreducible representations of Lie algebras to get irreducible representation of Lie group, by using the following theorems:

Theorem 5.19. *Let G be a connected matrix Lie group with Lie algebra $\mathfrak{g}$. Let $\pi : G \rightarrow GL(V)$ be a representation of G , and let $\pi' : \mathfrak{g} \rightarrow \mathfrak{gl}(V)$ be its associated representation of Lie algebra. Then π is irreducible if and only if π' is irreducible.*

A rigorous proof can be found in [5], proof of Proposition 16.39. Here we give an intuitive proof that is not very rigorous.

Proof. Suppose π is irreducible, then clearly π' is irreducible. To prove the converse, suppose π' is irreducible. Let $W \subset V$ be a G -invariant subspace. Then W is invariant under $\pi(e^{tX})$, where $t \in \mathbf{R}$ and $X \in \mathfrak{g}$. Then we claim that W is also invariant under

$$\pi'(X) = \frac{d}{dt}\pi(e^{tX})|_{t=0}.$$

To see this, let p be the projection map onto $W^\perp$. Since p is a linear and continuous map, we deduce that

$$p\left(\frac{d}{dt}\pi(e^{tX})|_{t=0}w\right) = \frac{d}{dt}|_{t=0} \circ p(\pi(e^{tX})w) = 0$$

for $w \in W$. But this contradicts the irreducibility of π' . $\square$

Theorem 5.20. *If G is a simply connected matrix Lie group, then every representation ϕ of $\mathfrak{g}$ can be constructed from some Lie group representation Φ of G , via*

$$\phi(X) = \frac{d}{dt}\Phi(e^{tX})|_{t=0}.$$

The proof of Theorem 5.20 can be found in almost any basic Lie theory text, such as [4]. Since we have discussed previously that there is a diffeomorphism $SU(2) \cong S^3$, which is simply connected, the previous theorems imply that it suffices to classify all irreducible representations of $\mathfrak{su}(2)$.

5.4.1 Structure of $\mathfrak{su}(2)$

Consider the 2×2 matrices

$$\sigma_1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \sigma_2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad \sigma_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

These are called Pauli matrices, which satisfies the commutation relations $[\sigma_j, \sigma_k] = 2i \sum_{l=1}^3 \epsilon_{jkl} \sigma_l$. It can be shown that these matrices form a basis of $\mathfrak{su}(2)$. Sometimes, people also define

$$X_j = -\frac{i}{2}\sigma_j$$

and consider them as a basis, for they satisfy a simpler commutation relation $[X_j, X_k] = \sum_{l=1}^3 \epsilon_{jkl} X_l$. Therefore, any group element $\Omega \in SU(2)$ near the identity matrix satisfies

$$\Omega \simeq I + \epsilon_1 X_1 + \epsilon_2 X_2 + \epsilon_3 X_3$$

for small real numbers ϵ_j . However, most physicists prefer another set of self-adjoint operators as their basis, which are

$$S_j = iX_j = \frac{1}{2}\sigma_j.$$

They corresponds to the spin of a particle. One can check that

$$[S_1, S_2] = iS_3, [S_2, S_3] = iS_1, [S_3, S_1] = iS_2.$$

5.4.2 The Highest Weight Theorem

We first provide a motivation of classifying irreducible representations of $\mathfrak{su}(2)$. We start from our original problem: classify irreducible representation of $SU(2)$. Notice that there is a natural inclusion $U(1) \rightarrow SU(2)$, given by

$$e^{i\theta} \mapsto \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix} \in SU(2).$$

We identify $U(1)$ with this isomorphic copy in $SU(2)$. Now let (π, V) be an irreducible representation of $SU(2)$. By restricting π to the subgroup $U(1) \subset SU(2)$, we can decompose our representation as

$$(\pi|_{U(1)}, V) = (\pi_{q_1}, \mathbf{C}) \oplus (\pi_{q_2}, \mathbf{C}) \oplus \cdots \oplus (\pi_{q_m}, \mathbf{C})$$

Therefore, we may choose a basis of V such that

$$\pi \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix} = \begin{pmatrix} e^{i\theta q_1} & 0 & \cdots & 0 \\ 0 & e^{i\theta q_2} & \cdots & 0 \\ \cdots & & \cdots & \cdots \\ 0 & 0 & \cdots & e^{i\theta q_m} \end{pmatrix}.$$

Theorem 5.21. *If q is in the set $\{q_j\}$, then so is $-q$.*

Proof. See pg 105 of [3]. □

Now notice that

$$e^{i2\theta S_3} = \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix}.$$

Therefore,

$$\pi(e^{i2\theta S_3}) = \begin{pmatrix} e^{i\theta q_1} & 0 & \cdots & 0 \\ 0 & e^{i\theta q_2} & \cdots & 0 \\ \cdots & & \cdots & \cdots \\ 0 & 0 & \cdots & e^{i\theta q_m} \end{pmatrix}.$$

Using the formula $\pi'(X) = \frac{d}{d\theta}(e^{\theta X})|_{\theta=0}$, we have that

$$\pi'(i2S_3) = \frac{d}{d\theta} \begin{pmatrix} e^{i\theta q_1} & 0 & \cdots & 0 \\ 0 & e^{i\theta q_2} & \cdots & 0 \\ \cdots & & \cdots & \cdots \\ 0 & 0 & \cdots & e^{i\theta q_m} \end{pmatrix} \Big|_{\theta=0} = \begin{pmatrix} iq_1 & 0 & \cdots & 0 \\ 0 & iq_2 & \cdots & 0 \\ \cdots & & \cdots & \cdots \\ 0 & 0 & \cdots & iq_m \end{pmatrix}.$$

Since π' is a linear map¹, we end up getting

$$\pi'(S_3) = \begin{pmatrix} \frac{q_1}{2} & 0 & \cdots & 0 \\ 0 & \frac{q_2}{2} & \cdots & 0 \\ \cdots & & \cdots & \cdots \\ 0 & 0 & \cdots & \frac{q_m}{2} \end{pmatrix}.$$

Therefore, we are motivated to make the following definition:

¹A rigorous argument would involve complexification of Lie algebra, which is introduced in section 5.5 of [3]. We will not go into such complication, and will just assume that π' is linear even with respect to complex coefficients.

Definition 5.22. If $\pi'(S_3)$ has an eigenvalue $\frac{k}{2}$, we say that k is a weight of the representation (π, V) . The subspace $V_k \subset V$ of the representation V is called the k th weight space of the representation if for all $v \in V_k$ we have $\pi'(S_3)v = \frac{k}{2}v$.

We call $\dim(V_k)$ the multiplicity of the weight k in the representation (π, V) .

Now we make one important definition:

Definition 5.23 (Raising and Lowering Operators). We define

$$S_+ = S_1 + iS_2 = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, S_- = S_1 - iS_2 = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}.$$

$S_{\pm}$ are elements in $\mathfrak{sl}(2, \mathbf{C})$. One can easily check that $(S_{\pm})^{\dagger} = S_{\mp}$, and similarly $\pi'(S_{\pm})^{\dagger} = \pi'(S_{\mp})$. To see why these two operators are called raising and lowering operators, notice that

$$[S_3, S_+] = [S_3, S_1 + iS_2] = iS_2 + i(-iS_1) = S_1 + iS_2 = S_+.$$

Since π' is a Lie algebra homomorphism, we have that

$$\pi'(S_3)\pi'(S_+) - \pi'(S_+)\pi'(S_3) = \pi'([S_3, S_+]) = \pi'(S_+).$$

So for any $v \in V_k$, we have

$$\pi'(S_3)\pi'(S_+)v = \pi'(S_+)\pi'(S_3)v + \pi'(S_+)v = \left(\frac{k}{2} + 1\right)\pi'(S_+)v.$$

Therefore, $v \in V_k$ implies that $\pi'(S_+)v \in V_{k+2}$. A similar calculation shows that $\pi'(S_-)$ takes V_k to V_{k-2} . Therefore $\pi'(S_+)$ takes a vector with a well-defined weight k to a vector with weight $k+2$, i.e., raising the weight, and similarly $\pi'(S_-)$ lowers the weight by 2.

Since V is finite dimensional, and eigenvectors of $\pi'(S_3)$ with distinct eigenvalues are linearly independent, there must exist an integer n , and a nonzero vector $v \in V_n \subset V$, such that

$$\pi'(S_+)v = 0.$$

Such a vector v is called a highest weight vector, with highest weight n .

Now we prove the key results in this section:

Theorem 5.24 (Highest weight theorem). *Finite dimensional irreducible representations of $SU(2)$ have weights of the form*

$$-n, -n+2, \dots, n-2, n$$

for n a nonnegative integer, each with multiplicity 1, with n a highest weight.

Proof. Suppose (π, V) is an irreducible representation of $SU(2)$. Then (π', V) is an irreducible representation of $\mathfrak{su}(2)$ by Theorem 5.19. Choose a highest weight vector $v_n \in V_n$ and a highest weight n (existence of n and v_n are guaranteed by finite-dimensionality of V). Define

$$v_{n-2j} = \pi'(S_-)^j v_n \in V_{n-2j}.$$

Consider the subspace of V defined by

$$W = \text{span}\{v_{n-2j} : j \geq 0\}.$$

W is $n + 1$ -dimensional with v_{-n} being a nonzero lowest weight vector (i.e. $\pi'(S_-)v_{-n} = 0$). Otherwise by Theorem 5.21, n will not be a highest weight. We claim that $W = V$, thus proving π is $n + 1$ -dimensional with weight $-n, -n + 2, \dots, n - 2, n$, each with multiplicity 1.

If we could show that W is an invariant subspace under π' , then irreducibility of π' would imply $W = V$, and we are done. Note that since

$$\mathfrak{su}(2) = \text{span}_{\mathbf{C}}\{S_3, S_-, S_+\},$$

it suffices to check that W is invariant under $\pi'(S_3), \pi'(S_-), \pi'(S_+)$. Clearly W is invariant under $\pi'(S_-)$ by definition, and W is obviously invariant under $\pi'(S_3)$ since W is the span of eigenvectors of $\pi'(S_3)$. One can show inductively that

$$\pi'(S_+)v_{n-2j} = j(n - j + 1)v_{n-2(j-1)}.$$

To see this, for $j = 0$ we have the highest weight condition on v_n . Assuming validity for j , validity for $j + 1$ can be checked by

$$\begin{aligned} \pi'(S_+)v_{n-2(j+1)} &= \pi'(S_+)\pi'(S_-)v_{n-2j} \\ &= (\pi'([S_+, S_-]) + \pi'(S_-)\pi'(S_+))v_{n-2j} \\ &= (\pi'(2S_3) + \pi'(S_-)\pi'(S_+))v_{n-2j} \\ &= ((n - 2j)v_{n-2j} + \pi'(S_-)j(n - j + 1)v_{n-2(j-1)}) \\ &= ((n - 2j) + j(n - j + 1))v_{n-2j} \\ &= (j + 1)(n - (j + 1) + 1)v_{n-2((j+1)-1)} \end{aligned}$$

where we have used the commutation relation $[S_+, S_-] = 2S_3$. This completes the proof of the theorem. $\square$

Therefore, we have shown that every irreducible representation (π_n, V^{n+1}) of $SU(2)$ can be characterized by a nonnegative interger n , namely the highest weight, and is $n + 1$ -dimensional.

5.4.3 Irreducible Representations of $SU(2)$

Although the highest weight theorem gives a characterization of irreducible representations of $SU(2)$, there is no clue how to construct such representations. This section will give an explicit construction of (π_n, V^{n+1}) for every n , following mostly pg 82-84 of [4]. Now, let n be fixed, and let V^{n+1} be the vector space of homogeneous polynomials with two variables of total degree n , i.e., the space of polynomials of the form

$$f(z_1, z_2) = a_0 z_1^n + a_1 z_1^{n-1} z_2 + \dots + a_{n-1} z_1 z_2^{n-1} + a_n z_2^n. \quad (5.1)$$

A basis for such a space is given by $z_1^n, z_1^{n-1} z_2, \dots, z_1 z_2^{n-1}, z_2^n$. Hence this is indeed an $n + 1$ -dimensional vector space.

For simplicity we denote π_n by π . Now viewing $(z_1, z_2) \in \mathbf{C}^2$, the representation $\pi : SU(2) \rightarrow V^{n+1}$ is defined by

$$\pi(g)f(z) = f(g^{-1}z).$$

Explicitly, if $f(z_1, z_2)$ are given by (5.1), then the action is explicitly given by

$$\pi(g)f(z_1, z_2) = \sum_{k=0}^n a_k (g_{11}^{-1}z_1 + g_{12}^{-1}z_2)^{m-k} (g_{21}^{-1}z_1 + g_{22}^{-1}z_2)^k$$

where g_{ij} are the entries of the element $g \in SU(2)$. By expanding the right-hand side of this formula, we see that $\pi(g)f$ is again a homogeneous polynomial of degree n , hence $\pi(g)$ maps V^{n+1} to V^{n+1} . It is easy to check that π is a homomorphism. Therefore, π is a well-defined representation.

Now we compute π' . We have

$$\pi'(X)f(z) = \frac{d}{dt}\pi(e^{tX})f(z)|_{t=0} = \frac{d}{dt}f(e^{-tX}z)|_{t=0}.$$

Now let $z(t) = (z_1(t), z_2(t))$ be the curve in $\mathbf{C}^2$ defined as $z(t) = e^{-tX}z$. By chain rule, we have that

$$\pi'(X)f = \frac{\partial f}{\partial z_1} \frac{dz_1}{dt} \Big|_{t=0} + \frac{\partial f}{\partial z_2} \frac{dz_2}{dt} \Big|_{t=0}.$$

Since $dz/dt|_{t=0} = -Xz$, we obtain that

$$\pi'(X)f = -\frac{\partial f}{\partial z_1}(X_{11}z_1 + X_{12}z_2) - \frac{\partial f}{\partial z_2}(X_{21}z_1 + X_{22}z_2).$$

Now plugging in the entries of $S_3, S_{\pm}$, we find that

$$\begin{aligned} \pi'(S_3) &= \frac{1}{2} \left(-z_1 \frac{\partial}{\partial z_1} + z_2 \frac{\partial}{\partial z_2} \right) \\ \pi'(S_+) &= -z_2 \frac{\partial}{\partial z_1} \\ \pi'(S_-) &= -z_1 \frac{\partial}{\partial z_2}. \end{aligned}$$

Applying these operators to a basis element $z_1^{n-k}z_2^k$ for V^{n+1} gives

$$\begin{aligned} \pi'(S_3)(z_1^{n-k}z_2^k) &= \frac{-n+2k}{2} z_1^{n-k}z_2^k \\ \pi'(S_+)(z_1^{n-k}z_2^k) &= -(n-k)z_1^{n-k-1}z_2^{k+1} \\ \pi'(S_-)(z_1^{n-k}z_2^k) &= -kz_1^{n-k+1}z_2^{k-1}. \end{aligned}$$

Thus, $z_1^{n-k}z_2^k$ is an eigenvector of $\pi'(S_3)$ with weight $-n+2k$. Therefore, z_2^n will be an explicit highest weight vector for the representation (π_n, V^n) . $\pi'(S_+)$ shift the weight of an eigenvector by 2, and $\pi'(S_-)$ shift the weight of an eigenvector down by 2.

Proposition 5.25. For each $n \geq 0$, the representation π_n is irreducible.

Proof. See Proposition 4.11 of [4]. □

5.4.4 Unitarity of the Representation

We conclude this section with a discussion of unitarity of the irreducible representations (π_n, V^{n+1}) of $SU(2)$. To discuss unitarity of representation, we need a notion of inner product on V^{n+1} . Usually we can define the inner product on $L^2(\mathbf{C}^2)$ by setting $\langle f, g \rangle = \int_{\mathbf{C}^2} \bar{f}g$. However, this is a useless inner product for f, g homogeneous polynomials, since they diverge. The trick is to define

$$\langle f, g \rangle = \frac{1}{\pi^2} \int_{\mathbf{C}^2} \bar{f}(z_1, z_2) g(z_1, z_2) e^{-(|z_1|^2 + |z_2|^2)} dx_1 dy_1 dx_2 dy_2.$$

Then one can check that the polynomials

$$\frac{z_1^j z_2^k}{\sqrt{j!k!}}$$

will be an orthonormal basis of the space of polynomial functions with respect to this inner product, and the operators $\pi'(X)$ will be skew-adjoint for any $X \in \mathfrak{su}(2)$. This shows that π_n is unitary with respect to this inner product.

5.5 Irreducible Representations of $SO(3)$

In this chapter, we follow mostly Chapter 16-17 of [5] and classify irreducible representations of the group $SO(3)$. It turns out that our previous work in classifying irreducible representations of $SU(2)$ is of great help: We will develop an analogous result for the highest weight theorem of irreducible representations of $SO(3)$, using the exact same idea of defining raising and lowering operators, and we will connect representations of $SO(3)$ with representations of $SU(2)$. But first, let us study some properties of $SO(3)$ group itself.

5.5.1 Properties of $SO(3)$ and $\mathfrak{so}(3)$

Recall that $SO(3)$ consists of the set of real matrices R with

$$R^T R = R R^T = I.$$

Let $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$ denote the standard basis in $\mathbf{R}^3$. Then rotations about coordinate axis by an angle θ are given by

$$R(\theta, \mathbf{e}_1) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix}$$

$$R(\theta, \mathbf{e}_2) = \begin{pmatrix} \cos \theta & 0 & \sin \theta \\ 0 & 1 & 0 \\ -\sin \theta & 0 & \cos \theta \end{pmatrix}$$

$$R(\theta, \mathbf{e}_3) = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

It turns out that any element $R \in SO(3)$ can be written as a composition of rotations along the axis parallel to these basis vectors. Therefore, we can choose 3 real numbers ϕ, θ, ψ (called the Euler angles) and write

$$R(\phi, \theta, \psi) = R(\psi, \mathbf{e}_3) R(\theta, \mathbf{e}_1) R(\phi, \mathbf{e}_3).$$

The infinitesimal picture on $\mathfrak{so}(3)$ is in fact much easier to understand. Recall that $\mathfrak{so}(3)$ is the set of all real 3×3 anti-symmetric matrices. A basis of $\mathfrak{so}(3)$ can be given by

$$l_1 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix} \quad l_2 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ -1 & 0 & 0 \end{pmatrix} \quad l_3 = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

which satisfies the commutation relation

$$[l_1, l_2] = l_3, [l_2, l_3] = l_1, [l_3, l_1] = l_2.$$

Definition 5.26. Let $\mathfrak{g}, \mathfrak{h}$ be two Lie algebras. A map $\phi : \mathfrak{g} \rightarrow \mathfrak{h}$ is called a Lie algebra homomorphism if it is linear and

$$\phi([X, Y]) = [\phi(X), \phi(Y)]$$

for all $X, Y \in \mathfrak{g}$. A Lie algebra homomorphism is called a Lie algebra isomorphism if it is bijective.

With the definition above, it can be shown that $\mathfrak{su}(2)$ is isomorphic to $\mathfrak{so}(3)$. The map ϕ is in fact very easy to construct. Consider the following basis for $\mathfrak{su}(2)$:

$$E_1 = \frac{i}{2}\sigma_3 \quad E_2 = \frac{i}{2}\sigma_2 \quad E_3 = \frac{i}{2}\sigma_1.$$

It can be shown that $[E_1, E_2] = E_3$ and relations obtained from this by cyclic permutation of the indices holds. Therefore, E_1, E_2, E_3 satisfies the exact same commutation relation as l_1, l_2, l_3 . So set $\phi(E_j) = l_j$ for $j = 1, 2, 3$, and extend by linearity.

Unfortunately, just because $\mathfrak{g}, \mathfrak{h}$ are isomorphic does not imply that their Lie groups G, H are isomorphic (It is true though if both groups are connected and simply connected). For a counterexample, $SU(2)$ is not isomorphic to $SO(3)$ despite their isomorphic Lie algebras. The reason is because $SO(3)$ is not simply connected. In fact $\pi_1(SO(3)) = \mathbf{Z}/2$.

Definition 5.27. Suppose G is a connected matrix Lie group with Lie algebra $\mathfrak{g}$. A **universal cover** of G is an ordered pair $(\tilde{G}, \Phi)$ consisting of a simply connected matrix Lie group $\tilde{G}$ and a Lie group homomorphism $\Phi : \tilde{G} \rightarrow G$ such that the associated Lie algebra homomorphism $\phi : \tilde{\mathfrak{g}} \rightarrow \mathfrak{g}$ is an isomorphism.

Now we cite the following theorem without proving it:

Theorem 5.28. Suppose G, H are matrix Lie groups with Lie algebras $\mathfrak{g}, \mathfrak{h}$. And suppose $\phi : \mathfrak{g} \rightarrow \mathfrak{h}$ is a Lie algebra homomorphism. If G is connected and simply connected, then there exists a unique Lie group homomorphism $\Phi : G \rightarrow H$ such that $\Phi' = \phi$.

This theorem implies that if G is connected and simply connected, then every representation of G can be obtained by exponentiating some representation of the Lie algebra $\mathfrak{g}$. The proof requires the application of the Baker-Campbell-Hausdorff formula. For a detailed proof, consult Section 3.6 of [4]. If we set $G = SU(2)$, then since $SU(2) \cong S^3$ which is connected and simply connected, the Lie algebra isomorphism ϕ defined by $\phi(E_j) = l_j$ above can be extended to a unique Lie group homomorphism Φ . Then one can show

Theorem 5.29. *$(SU(2), \Phi)$ is a universal cover of $SO(3)$ with $\ker \Phi = \pm I$, where Φ is the unique homomorphism such that $\Phi' = \phi$, and $\phi(E_j) = l_j$.*

Proof. This is shown on pg 349, example 16.34 of [5]. $\square$

We will not go into the detail of the proof. It is sufficient for us to realize what $\Phi : SU(2) \rightarrow SO(3)$ looks like. Since $\Phi' = \phi$, we have that

$$\Phi(e^{tX}) = e^{t\phi(X)}.$$

And since $SU(2)$ is simply connected, any element in $SU(2)$ can be written as e^{tX} for some real t and $X \in \mathfrak{su}(2)$. Therefore, we know everything about the map Φ already. One can also realize Φ through quaternions, which is explained in Chapter 6 of [3]. However, all we really need here is to know that there is a covering map $\Phi : SU(2) \rightarrow SO(3)$.

5.5.2 Irreducible Representation of $\mathfrak{so}(3)$

Recall that in the classification of irreps of $SU(2)$, we first look at its Lie algebra and deduce the highest weight theorem. We will use the same strategy and look at irreps of $\mathfrak{so}(3)$ first. We continue to use the basis l_1, l_2, l_3 for $\mathfrak{so}(3)$.

Theorem 5.30 (Highest Weight Theorem). *Let $\pi : \mathfrak{so}(3) \rightarrow \mathfrak{gl}(V)$ be a finite-dimensional irreducible representation of $\mathfrak{so}(3)$ on V . We define operators $L_{\pm}, L_3$ on V by*

$$\begin{aligned} L_+ &= i\pi(l_1) - \pi(l_2) \\ L_- &= i\pi(l_1) + \pi(l_2) \\ L_3 &= i\pi(l_3). \end{aligned}$$

Let $l = \frac{1}{2}(\dim V - 1)$, so that $\dim V = 2l + 1$. Then there exists a basis $v_0, v_1, \dots, v_{2l}$ of V such that

$$\begin{aligned} L_3 v_j &= (l - j)v_j \\ L_- v_j &= \begin{cases} v_{j+1} & j < 2l \\ 0 & j = 2l \end{cases} \\ L_+ v_j &= \begin{cases} j(2l + 1 - j)v_{j-1} & j > 0 \\ 0 & j = 0. \end{cases} \end{aligned}$$

Thus, the number l completely determines the structure of an irrep of $\mathfrak{so}(3)$. Since $\dim V$ must be a nonnegative integer, l can only be one of the following values: $l = 0, \frac{1}{2}, 1, \frac{3}{2}, \dots$

Proof. Since π is a Lie algebra homomorphism, then $\pi(l_j)$ s satisfy the same commutation relations as l_j s. So we have

$$[L_3, L_+] = L_+ \quad [L_3, L_-] = -L_- \quad [L_+, L_-] = 2L_3.$$

Since over $\mathbf{C}$, the operator L_3 has at least one eigenvector v with eigenvalue λ , consider L_+v . Using the commutation relation, we have

$$L_3L_+v = (L_+L_3 + L_+)v = (\lambda + 1)L_+v.$$

So either $L_+v = 0$ or L_+v is an eigenvector of L_3 with eigenvalue $\lambda + 1$. We call L_+ raising operator for this reason. Since L_3 only has finitely many eigenvalues, there exists $k \geq 0$ such that $(L_+)^k v \neq 0$ but $(L_+)^{k+1} v = 0$. Clearly $(L_+)^k v$ is an eigenvector of L_3 with eigenvalue $\lambda + k$.

Now let us introduce new Notation $v_0 = (L_+)^k v$ and $\mu = \lambda + k$. Then v_0 is a nonzero eigenvector of L_3 with eigenvalue μ and $L_+v_0 = 0$. We now forget about v, λ and consider only v_0, μ . We define

$$v_j = (L_-)^j v_0 \quad j = 0, 1, 2, \dots$$

By similar arguments using commutation relation, we see that L_- has the effect of lowering eigenvalue of L_3 by 1 or giving zero vector. Thus

$$L_3 v_j = (\mu - j) v_j.$$

Next, one can check by induction that

$$L_+ v_j = j(2\mu + 1 - j) v_{j-1} \quad j = 0, 1, 2, 3, \dots$$

When $j = 0$ this is the highest weight condition. Since L_3 has only finitely many eigenvectors, v_j must eventually be zero. Thus there exists some $N \geq 0$ such that $v_N \neq 0$ but $v_{N+1} = 0$. Since $v_{N+1} = 0$, apply the previous equation with $j = N$ gives

$$0 = L_+ v_{N+1} = (N + 1)(2\mu - N) v_N.$$

Now $v_N \neq 0$ and $N + 1 > 0$. So we must have $2\mu - N = 0$, i.e., $\mu = N/2$.

Now we set $l = N/2$ and putting $\mu = N/2 = l$, we have all the formulas in the statement of the theorem. Since v_j s are eigenvectors of L_3 with distinct eigenvalues, they are linearly independent. Furthermore, the span of v_j s are invariant under $L_{\pm}, L_3$, which is easy to check. Since V is assumed to be irreducible, v_j s must be a basis of V . Then

$$\dim V = \#\{v_j\} = N + 1 = 2l + 1.$$

This completes the proof. □

Theorem 5.31. *For any $l = 0, \frac{1}{2}, 1, \frac{3}{2}, \dots$ there exists an irreducible representation of $\mathfrak{so}(3)$ of dimension $2l + 1$, and any two irreducible representations of $\mathfrak{so}(3)$ of dimension $2l + 1$ are isomorphic.*

Proof. We construct V simply by defining a space V with basis $v_0, v_1, \dots, v_{2l}$ and defining action of $\mathfrak{so}(3)$ by the formulas in the previous theorem. It is easy to check that $L_{\pm}, L_3$ defined in this way, have the correct commutation relation. Hence (π, V) defined this way is indeed a representation of $\mathfrak{so}(3)$.

Now we show that V is irreducible. Suppose that W is a nonzero invariant subspace of V . Take a nonzero vector $w = \sum_{j=0}^{2l} a_j v_j \in W$. Let j_0 be the largest index for which a_j is nonzero. According to the formula in the previous theorem, $(L_+)^{j_0} w$ will be a nonzero multiple of v_0 . This means $v_0 \in W$. Repeatedly applying L_- to v_0 , we see that $v_j \in W$ for all j . Hence $W = V$.

Finally, the previous theorem tells us that any two irreducible representations of $\mathfrak{so}(3)$ has a basis as in the formulas of the theorem. Then we can construct an isomorphism between any two irreducible representations by mapping this basis in one space to the corresponding basis in the other space. $\square$

5.5.3 Irreducible Representations of $SO(3)$

Having classified irreducible representation of the Lie algebra $\mathfrak{so}(3)$, we now classify irreducible representation of $SO(3)$. It can be shown that $SO(3)$ is connected. Therefore, by Theorem 5.19, a representation Π of $SO(3)$ is irreducible if and only if its associated Lie algebra representation $\Pi' = \pi$ is irreducible. It can be shown that Π_1, Π_2 of $SO(3)$ are isomorphic if and only if π_1, π_2 are isomorphic (see [5], Proposition 16.39). Thus, to classify irreducible representations of $SO(3)$ up to isomorphism, one only need to determine which irreducible representations of $\mathfrak{so}(3)$ come from a representation of the group $SO(3)$.

Theorem 5.32. *Let $\pi_l : \mathfrak{so}(3) \rightarrow \mathfrak{gl}(V)$ be an irreducible representation of $\mathfrak{so}(3)$, with $l = \frac{1}{2}(\dim V - 1)$. If l is an integer, i.e., if $\dim V$ is odd, then there is a representation $\Pi_l : SO(3) \rightarrow GL(V)$ such that $\Pi'_l = \pi_l$. If l is a half-integer, i.e., if $\dim V$ is even, then no such representation Π_l exists.*

Proof. If l is a half-integer, then L_3 is diagonal in the basis $\{v_j\}$ with eigenvalue being half-integers. Therefore $e^{-2\pi i L_3} = -I$. Then

$$e^{2\pi i \pi_l(L_3)} = e^{-2\pi i L_3} = -I.$$

Simple calculation shows that $e^{2\pi i l_3} = I$. Thus if a corresponding representation Π_l of $SO(3)$ existed, then we would have

$$\Pi_l(I) = \Pi_l(e^{2\pi i l_3}) = e^{2\pi i \pi_l(L_3)} = -I$$

which is a contradiction.

If l is an integer, let $\phi : \mathfrak{su}(2) \rightarrow \mathfrak{so}(3)$ be the isomorphism such that $\phi(E_j) = l_j, j = 1, 2, 3$. We obtain a representation π'_l of $\mathfrak{su}(2)$ by $\pi'_l(X) = \pi_l(\phi(X))$. Since $SU(2)$ is simply connected, there is a representation Π'_l whose derivative is π'_l . Then we compute

$$\Pi'_l(-I) = \Pi'_l(e^{2\pi i E_1}) = e^{2\pi i \pi'_l(E_1)} = e^{2\pi i \pi_l(l_1)} = e^{2\pi i L_3} = I$$

since the eigenvalues of L_3 are now integers.

Now, as we have discussed, there is a surjective homomorphism $\Phi : SU(2) \rightarrow SO(3)$ for which $\Phi' = \phi$ and $\ker \Phi = \pm I$. Since $\ker \Pi'_l$ contains $\pm I$, the map Π'_l factors through $SO(3)$, i.e., for every $U \in SU(2)$, define $\Pi_l : SO(3) \rightarrow GL(V)$, where $GL(V)$ is the codomain of Π'_l , by

$$\Pi_l(\Phi(U)) = \Pi'_l(U)$$

for every $U \in SU(2)$. Since Φ is surjective and $\ker \Phi \subset \ker \Pi'_l$, the map Π_l is well-defined. One can check that Π_l is the desired irreducible representation of $SO(3)$, for its Lie algebra is exactly π_l , an irreducible representation of $\mathfrak{so}(3)$. This is because the relation $\Pi'_l = \Pi_l \circ \Phi$, hence this induces the differential relation $\pi'_l = \sigma_l \circ \phi$, where σ_l is the derivative of Π_l . Therefore, $\sigma_l = \pi'_l \circ \phi^{-1} = \pi_l$, since we defined $\pi'_l = \pi_l \circ \phi$.

Consult Proposition 17.10 of [5] for details. $\square$

5.5.4 Spherical Harmonics

We will now construct irreducible representations of $SO(3)$ explicitly. First, recall that in the case of $SO(3)$, there is an irreducible representation Π_l of dimension $2l + 1$ for every $l = 0, 1, 2, \dots$. We will now provide an explicit construction of Π_l and the corresponding vector space V^{2l+1} .

Now, fix an integer l , and denote the irreducible representation of $SO(3)$ simply by (Π, V) . Since $SO(3)$ acts on $\mathbf{R}^3$, then as discussed previously, we have an induced representation on $F(\mathbf{R}^3)$, all complex functions on $\mathbf{R}^3$, by

$$(\Pi(g)f)(\mathbf{x}) = f(g^{-1}\mathbf{x}).$$

However, the space of all complex functions on $\mathbf{R}^3$ is far too large for our vector space V , which is only $2l + 1$ -dimensional. We need to find a proper subspace $V \subset F(\mathbf{R}^3)$, on which the representation Π defined above is irreducible. Since $SO(3)$ is connected, Π is irreducible if and only if Π' is irreducible. Therefore, a natural strategy of finding V is to compute L_+ , and use the highest weight condition to find the highest weight vector (which, as you will see, is equivalent to solving a PDE), and apply L_- to this highest weight vector repeatedly and get the whole list of $2l + 1$ vectors, which will then be the basis for V .

Now, we remember that

$$\Pi'(X)f = \frac{d}{dt}\Pi(e^{tX})f|_{t=0} = \frac{d}{dt}f(e^{-tX}\mathbf{x})|_{t=0}$$

for $X \in \mathfrak{so}(3)$. Recall that a basis of $\mathfrak{so}(3)$ is given by

$$l_1 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix} \quad l_2 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ -1 & 0 & 0 \end{pmatrix} \quad l_3 = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

They satisfy the commutation relation

$$[l_1, l_2] = l_3, [l_2, l_3] = l_1, [l_3, l_1] = l_2.$$

To be consistent with physics literatures, we define

$$L_j = i\Pi'(l_j) \quad j = 1, 2, 3.$$

Then one can check that

$$[L_1, L_2] = iL_3, [L_2, L_3] = iL_1, [L_3, L_1] = iL_2.$$

We also define

$$l_{\pm} = l_1 \pm il_2 \quad L_{\pm} = i\Pi'(l_{\pm}).$$

It turns out that L_j will correspond to angular momentum of a particle in 3 directions, respectively. Now using the formula for calculating $\Pi'(X)$ above (see [3], section 8.3 for detailed calculation), we get that

$$\begin{aligned} \Pi'(l_1) &= x_3 \frac{\partial}{\partial x_2} - x_2 \frac{\partial}{\partial x_3} \\ \Pi'(l_2) &= x_1 \frac{\partial}{\partial x_3} - x_3 \frac{\partial}{\partial x_1} \\ \Pi'(l_3) &= x_2 \frac{\partial}{\partial x_1} - x_1 \frac{\partial}{\partial x_2}. \end{aligned}$$

Now, we realize that when $SO(3)$ acts on $\mathbf{R}^3$, it keeps the distance r to the origin of a vector $\mathbf{x}$ fixed. The innate rotational symmetry of $SO(3)$ action motivates us to restrict our attention to functions on the unit sphere S^2 , i.e., function of the form $f(\theta, \phi)$, and change our coordinate system from x_1, x_2, x_3 to spherical coordinates (r, θ, ϕ) .

We have to figure out the form of $\Pi'(l_j)$ with respect to this coordinate system, instead of the coordinate system x_1, x_2, x_3 . The coordinate transformation is given by

$$\begin{aligned} x_1 &= r \sin \theta \cos \phi \\ x_2 &= r \sin \theta \sin \phi \\ x_3 &= r \cos \theta. \end{aligned}$$

The harder part is to transform the derivative operators into spherical coordinates. But essentially this is just a straightforward change of basis of tangent vectors. See, for example, p.63 of [1] or section 8.3 of [3]. After some calculation, we get

$$\begin{aligned} L_1 &= i\Pi'(l_1) = i \left(x_3 \frac{\partial}{\partial x_2} - x_2 \frac{\partial}{\partial x_3} \right) = i \left(\sin \phi \frac{\partial}{\partial \theta} + \cot \theta \cos \phi \frac{\partial}{\partial \phi} \right) \\ L_2 &= i\Pi'(l_2) = i \left(x_1 \frac{\partial}{\partial x_3} - x_3 \frac{\partial}{\partial x_1} \right) = i \left(-\cos \phi \frac{\partial}{\partial \theta} + \cot \theta \sin \phi \frac{\partial}{\partial \phi} \right) \\ L_3 &= i\Pi'(l_3) = i \left(x_2 \frac{\partial}{\partial x_1} - x_1 \frac{\partial}{\partial x_2} \right) = -i \frac{\partial}{\partial \phi} \end{aligned}$$

which then gives

$$\begin{aligned} L_+ &= i\Pi'(l_+) = e^{i\phi} \left(\frac{\partial}{\partial \theta} + i \cot \theta \frac{\partial}{\partial \phi} \right) \\ L_- &= i\Pi'(l_-) = e^{-i\phi} \left(-\frac{\partial}{\partial \theta} + i \cot \theta \frac{\partial}{\partial \phi} \right). \end{aligned}$$

Now, as mentioned before, we first find the highest weight vector v_0 in the previous section. Then we will apply L_- to v_0 repeatedly to obtain a basis of our vector space. We will call our highest weight vector v_0 a function $Y_l^l(\theta, \phi)$,

which is a function defined on a unit sphere. It satisfies $L_3 Y_l^l(\theta, \phi) = l Y_l^l(\theta, \phi)$, and $L_+ Y_l^l(\theta, \phi) = 0$. Therefore, we get two PDEs:

$$\begin{aligned} L_3 Y_l^l(\theta, \phi) &= -i \frac{\partial}{\partial \phi} Y_l^l(\theta, \phi) = l Y_l^l(\theta, \phi) \\ L_+ Y_l^l(\theta, \phi) &= e^{i\phi} \left(\frac{\partial}{\partial \theta} + i \cot \theta \frac{\partial}{\partial \phi} \right) Y_l^l(\theta, \phi) = 0. \end{aligned}$$

The first equation suggests that

$$Y_l^l(\theta, \phi) = e^{il\phi} F_l(\theta)$$

for some function $F_l(\theta)$ that depends only on θ , and using the second equation we get

$$\left(\frac{\partial}{\partial \theta} - l \cot \theta \right) F_l(\theta) = 0.$$

This has the solution

$$F_l(\theta) = C_l \sin^l \theta$$

for some constant C_l . So finally we get our highest weight vector

$$Y_l^l(\theta, \phi) = C_l e^{il\phi} \sin^l \theta.$$

Repeatedly applying the lowering operator L_- gives the rest of the basis:

$$\begin{aligned} Y_l^m(\theta, \phi) &= C_{lm} (L_-)^{l-m} Y_l^l(\theta, \phi) \\ &= C_{lm} \left(e^{-i\phi} \left(-\frac{\partial}{\partial \theta} + i \cot \theta \frac{\partial}{\partial \phi} \right) \right)^{l-m} e^{il\phi} \sin^l \theta \end{aligned}$$

for $m = l, l-1, \dots, -l+1, -l$. Therefore, our $2l+1$ dimensional vector space V is given by the span of the functions $Y_l^m(\theta, \phi)$ on the unit sphere S^2 , for $m = l, l-1, \dots, -l+1, -l$.

We still need to show that representation (Π, V) is irreducible. But this is very easy. We know that Π is irreducible if and only if Π' is irreducible. Since we know that $l_\pm, l_3$ form a basis of $\mathfrak{so}(3)$, we only need to check that $V = \text{span}\{Y_l^m(\theta, \phi)\}_{m=l, \dots, -l}$ is closed under the action of $L_\pm, L_3$. Using similar argument as we did in the proof of Theorem 5.31, one can show that V is indeed invariant under $L_\pm, L_3$. Therefore, (Π, V) is indeed irreducible.

The functions $Y_l^m(\theta, \phi)$ are called **spherical harmonics**. They belong to the set of square integrable functions on the unit sphere, denoted by $L^2(S^2)$. This space can be given an inner product by

$$\langle f, g \rangle = \int_{S^2} \bar{f} g \sin \theta \, d\theta \, d\phi = \int_{\phi=0}^{2\pi} \int_{\theta=0}^{\pi} \bar{f}(\theta, \phi) g(\theta, \phi) \sin \theta \, d\theta \, d\phi.$$

Then one can define the L^2 -norm in this space by

$$\|f\|_{L^2(S^2)} = \langle f, f \rangle^{1/2}.$$

Then, by some nontrivial functional analysis technique, one can show that $Y_l^m(\theta, \phi)$ form a basis for $L^2(S^2)$, where $l = 0, 1, 2, \dots$ and $-l \leq m \leq l$. In other words, any $f(\theta, \phi) \in L^2(S^2)$ can be written as a series expansion

$$f(\theta, \phi) = \sum_{l=0}^{\infty} \sum_{m=-l}^l a_l^m Y_l^m(\theta, \phi)$$

in the sense that the infinite series on the right hand side converges to f in L^2 -norm. In fact, this is an orthonormal basis, in the sense that

$$\langle Y_l^m(\theta, \phi), Y_{l'}^{m'}(\theta, \phi) \rangle = \delta_{mm'} \delta_{ll'}.$$

Therefore, we see that

$$a_l^m = \langle Y_l^m(\theta, \phi), f(\theta, \phi) \rangle.$$

Conventionally, the constant C_{ml} in the function $Y_l^m(\theta, \phi)$ is chosen such that this basis is orthonormal, i.e., such that $\|Y_l^m(\theta, \phi)\|_{L^2(S^2)} = 1$.

5.5.5 The Casimir operator

One can define the Casimir operator on $SO(3)$ to be the second order differential operator

$$L^2 \equiv L_1^2 + L_2^2 + L_3^2.$$

One can show that for any $X \in \mathfrak{so}(3)$, we have

$$[L^2, \Pi'_l(X)] = 0.$$

Knowing this, a version of Schur's lemma says that L^2 will act on an irreducible representation as a scalar (i.e., all vectors in the representation are eigenvectors of L^2 , with the same eigenvalue). This eigenvalue can be used to characterize the irreducible representation. To compute its eigenvalue, note that

$$\begin{aligned} L_- L_+ &= (L_1 - iL_2)(L_1 + iL_2) \\ &= L_1^2 + L_2^2 + i[L_1, L_2] \\ &= L_1^2 + L_2^2 - L_3. \end{aligned}$$

So

$$L^2 = L_1^2 + L_2^2 + L_3^2 = L_- L_+ + L_3 + L_3^2.$$

Let f be the highest weight vector of the irreducible representation Π_l . Then

$$L_+ f = 0, \quad L_3 f = l f$$

with solution $f \propto Y_l^l(\theta, \phi)$. Therefore,

$$L^2 f = L_- L_+ f + (L_3 + L_3^2) f = (0 + l + l^2) f = l(l+1) f.$$

We have thus shown that our irreducible representation Π_l can be characterized as the representation on which L^2 acts by the scalar $l(l+1)$. We will reencounter this operator in as the angular part of the Laplace operator on $\mathbf{R}^3$. Finally, an explicit calculation shows that

$$\begin{aligned} L^2 &= L_1^2 + L_2^2 + L_3^2 \\ &= \left(i \left(\sin \phi \frac{\partial}{\partial \theta} + \cot \theta \cos \phi \frac{\partial}{\partial \phi} \right) \right)^2 \\ &\quad + \left(i \left(-\cos \phi \frac{\partial}{\partial \theta} + \cot \theta \sin \phi \frac{\partial}{\partial \phi} \right) \right)^2 + \left(-i \frac{\partial}{\partial \phi} \right)^2 \\ &= - \left(\frac{1}{\sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial}{\partial \theta} \right) + \frac{1}{\sin^2 \theta} \frac{\partial^2}{\partial \phi^2} \right). \end{aligned}$$

5.6 Tensor Products

Let V, W be two vector spaces over $\mathbf{C}$, with basis $\{e_1, \dots, e_m\}, \{f_1, \dots, f_n\}$ respectively. Then the direct sum $V \oplus W$ gives a vector space of dimension $m + n$ with basis $\{e_1, \dots, e_m, f_1, \dots, f_n\}$. The tensor product $V \otimes W$ gives a vector space with dimension mn with basis

$$e_i \otimes f_j \quad i = 1, \dots, m \quad j = 1, \dots, n.$$

We require $\otimes$ to be a bilinear operation. This is, of course, not a rigorous definition of the tensor product, since we haven't check that this definition is independent of the choice of basis. However, we will not do a formal definition of tensor products, which is quite abstract and can be found in standard abstract algebra textbooks.

There are natural isomorphisms

$$V \otimes W \cong W \otimes V \quad U \otimes (W \otimes W) \cong (U \otimes V) \otimes W.$$

Also, given linear operators A on V and B on W , naturally we can define a linear operator $A \otimes B$ on $V \otimes W$ by

$$(A \otimes B)(v \otimes w) = Av \otimes Bw.$$

If one writes out the matrix representation carefully, $A \otimes B$ will be an $mn \times mn$ matrix.

One more thing we shall say about tensor product is the following. Let $\mathcal{L}(V, W)$ denote linear maps from V to W . Then there is an isomorphism

$$V^* \otimes W \cong \mathcal{L}(V, W)$$

which is given by the correspondence

$$l \otimes w \leftrightarrow (v \mapsto l(v)w).$$

5.6.1 Tensor Products in Quantum Mechanics: Pauli Exclusion and Entanglement

The application of tensor products in quantum mechanics is the following: Let $\mathcal{H}_1, \mathcal{H}_2$ be two quantum systems. One can consider the composite quantum system corresponding to considering the two systems as a single one, with no interaction between them, by taking a new state space

$$\mathcal{H} = \mathcal{H}_1 \otimes \mathcal{H}_2$$

with operators of the form

$$A \otimes I_B + I_A \otimes B \quad A \in \mathcal{L}(\mathcal{H}_1, \mathcal{H}_1), B \in \mathcal{L}(\mathcal{H}_2, \mathcal{H}_2).$$

If one consider a state space $\mathcal{H}$ as a system describing a single particle, then a system of n such particles is described by the state space $\mathcal{H}^{\otimes n}$.

One has a representation $(\pi, \mathcal{H}^{\otimes n})$ of the symmetric group S_n given by

$$\pi(\sigma)(v_1 \otimes v_2 \otimes \dots \otimes v_n) = v_{\sigma(1)} \otimes v_{\sigma(2)} \otimes \dots \otimes v_{\sigma(n)}.$$

A fundamental axiom of quantum mechanics is that if $\mathcal{H}^{\otimes n}$ describes n identical particles, then all physical states are either symmetric (bosons) or anti-symmetric (fermions).

Definition 5.33. A state $v \in \mathcal{H}^{\otimes n}$ is called

- symmetric, or ***bosonic*** if for all $\sigma \in S_n$, one have $\pi(\sigma)v = v$. The space of such states is denoted $S^n(\mathcal{H})$.
- antisymmetric, or ***fermionic*** if for all $\sigma \in S_n$, one have $\pi(\sigma)v = (-1)^{|\sigma|}v$, where $|\sigma|$ is the minimal number of transpositions that by composition give σ . The space of such states is denoted $\Lambda^n(\mathcal{H})$.

Note that in the fermionic case, let $\sigma \in S_2$ be the nontrivial transposition. Then $\pi(\sigma)$ acts on $\mathcal{H} \otimes \mathcal{H}$ by interchanging two vectors. Thus it takes $w \otimes w \in \mathcal{H} \otimes \mathcal{H}$ to itself. Antisymmetry requires $\pi(\sigma)$ takes this state to its negative. So the state cannot be nonzero. As a result, we have proven

Pauli Exclusion Principle. One cannot have nonzero states in $\mathcal{H}^{\otimes n}$ describing two identical fermions in the same state $w \in \mathcal{H}$.

Definition 5.34. A vector in $V \otimes W$ is called ***decomposable*** if it is of the form $v \otimes w$ for some $v \in V, w \in W$. If it cannot be put in this form, it is called ***indecomposable***.

Clearly, basis $e_i \otimes f_j$ are all decomposable. However, their linear combinations are in general indecomposable. In the physics literature, the language used is

Definition 5.35. An indecomposable state in the tensor product state space $\mathcal{H}_1 \otimes \mathcal{H}_2$ is called an entangled state.

5.6.2 Tensor Products of Representations

Definition 5.36. Let $(\pi_V, V), (\pi_W, W)$ be representations of a group G . There is a tensor product representation $(\pi_{V \otimes W}, V \otimes W)$ of G , defined by

$$(\pi_{V \otimes W}(g))(v \otimes w) = \pi_V(g)v \otimes \pi_W(g)w.$$

One can easily check that $\pi_{V \otimes W}(gh) = \pi_{V \otimes W}(g)\pi_{V \otimes W}(h)$, hence this is indeed a representation.

To see what happens to the Lie algebra representation, we first have to figure out what $v \otimes w$ looks like in the coordinate form. Suppose with respect to the ordered basis $e_1, \dots, e_m$ of V and $f_1, \dots, f_n$ of W , one has

$$v = \begin{pmatrix} v_1 \\ \vdots \\ v_m \end{pmatrix} \quad w = \begin{pmatrix} w_1 \\ \vdots \\ w_n \end{pmatrix}$$

where m, n are not necessarily equal. Then using the linearity of $\otimes$, one can see that $v \otimes w$ with respect to the ordered basis

$$e_1 \otimes f_1, \dots, e_1 \otimes f_n, e_2 \otimes f_1, \dots, e_2 \otimes f_n, \dots, e_m \otimes f_1, \dots, e_m \otimes f_n$$

is given by

$$v \otimes w = \begin{pmatrix} v_1 w_1 \\ \vdots \\ v_1 w_n \\ v_2 w_1 \\ \vdots \\ v_2 w_n \\ \vdots \\ v_m w_1 \\ \vdots \\ v_m w_n \end{pmatrix} \in V \otimes W.$$

Then by the chain rule, if v, w are vector-valued functions of t , one can see that

$$\frac{d}{dt} v \otimes w = \frac{dv}{dt} \otimes w + v \otimes \frac{dw}{dt}.$$

Therefore, we have that

$$\begin{aligned} \pi'_{V \otimes W}(X)(v \otimes w) &= \frac{d}{dt} \pi_{V \otimes W}(e^{tX})(v \otimes w)_{t=0} \\ &= \frac{d}{dt} (\pi_V(e^{tX})v \otimes \pi_W(e^{tX})w)_{t=0} \\ &= \left(\left(\frac{d}{dt} \pi_V(e^{tX})v \right) \otimes \pi_W(e^{tX})w \right)_{t=0} \\ &\quad + \left(\pi_V(e^{tX})v \otimes \left(\frac{d}{dt} \pi_W(e^{tX})w \right) \right)_{t=0} \\ &= (\pi'_V(X)v) \otimes w + v \otimes (\pi'_W(X)w). \end{aligned}$$

Therefore, one has

$$\pi'_{V \otimes W}(X) = (\pi'_V(X) \otimes \mathbf{1}_W) + (\mathbf{1}_V \otimes \pi'_W(X)).$$

5.6.3 Clebsch-Gordan Decomposition

Theorem 5.37 (Clebsch-Gordan Decomposition Theorem). *Let (Π_k, V^k) denote the $k+1$ -dimensional irreducible representations of $SU(2)$. Then tensor product $(\Pi_{V^{n_1} \otimes V^{n_2}}, V^{n_1} \otimes V^{n_2})$ decomposes into irreducible representations as*

$$(\Pi_{n_1+n_2}, V^{n_1+n_2}) \oplus (\Pi_{n_1+n_2-2}, V^{n_1+n_2-2}) \oplus \cdots \oplus (\Pi_{|n_1-n_2|}, V^{|n_1-n_2|}).$$

To prove this theorem, we need to first make a

Definition 5.38. The **character** of a representation (π, V) of a group G is a function $\chi_V : G \rightarrow \mathbf{C}$, defined by

$$\chi_V(g) = \text{tr}(\pi(g)).$$

One can easily check the following properties of the character:

Proposition 5.39. We have:

- $\chi_{V \oplus W} = \chi_V + \chi_W$
- $\chi_{V \otimes W} = \chi_V \cdot \chi_W$
- $\chi_V(ghg^{-1}) = \chi_h$.

To prove the last claim, use $\text{tr}(AB) = \text{tr}(BA)$. Therefore, we see that the character is a function on conjugacy classes. Namely, if two elements $g, h \in G$ are conjugates of each other, then $\chi_V(g) = \chi_V(h)$. In the case where $G = SU(2)$, we have the following result:

Lemma 5.40. *Let $A \in SU(2)$, then there exists $Q \in SU(2)$ and $\Lambda \in U(1)$ such that $A = Q\Lambda Q^{-1}$. That is, every element in $SU(2)$ is conjugate to some element in $U(1) \subset SU(2)$.*

Proof. The spectral theorem for normal operators (i.e., operators such that $[A, A^\dagger] = 0$) says that for normal operator A we have $A = Q\Lambda Q^{-1}$ where Q is unitary and Λ is diagonal. Since if $A \in SU(2)$, then $AA^\dagger = A^\dagger A = I$, so A is clearly normal, we want to show that Q can be chosen to be in $SU(2)$, i.e., we can choose Q with determinant 1. Since Q is unitary, we have $QQ^\dagger = Q^\dagger Q = I$, hence $(\det Q)(\det Q^\dagger) = 1$. Let $d = \det Q$. Then $\det(Q^\dagger) = \det(\overline{Q}) = \overline{d}$ since transpose does not change determinant. Thus $d\overline{d} = 1$. Hence d is on the unit circle. Let s be the square root of d . Consider $Q' = \overline{s}Q$. We have $\det(Q') = (\overline{s})^2 \det Q = \overline{d}d = 1$. So $Q' \in SU(2)$. Now $A' = Q'\Lambda(Q')^\dagger$ is diagonal, and it differs with A by a scalar multiple. This proves our claim.

Now suppose $A = Q\Lambda Q^{-1}$ with $A, Q \in SU(2)$ and Λ diagonal. Then $\Lambda = Q^{-1}AQ \in SU(2)$ as well. Since $\Lambda \in SU(2)$, we can write

$$\Lambda = \begin{pmatrix} \alpha & \beta \\ -\overline{\beta} & \overline{\alpha} \end{pmatrix} \quad |\alpha|^2 + |\beta|^2 = 1.$$

Since Λ is diagonal, $\beta = 0$ and $|\alpha|^2 = 1$. Therefore $\Lambda \in U(1)$. □

Corollary 5.41. *Let (Π_V, V) and (Π_W, W) be two representations of $SU(2)$. Then if $\chi_V = \chi_W$ on $U(1)$, we must have $\chi_V = \chi_W$ on the entire group $SU(2)$.*

Proof. Immediate by the previous lemma, and the fact that character is invariant under conjugation. □

Proof of the Clebsch-Gordan Decomposition Theorem. We know that the irreducible representation (Π_n, V^n) acts on $U(1)$ by

$$\Pi_n \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix} = \begin{pmatrix} e^{in\theta} & & & \\ & e^{i(n-2)\theta} & & \\ & & \ddots & \\ & & & e^{i(-n+2)\theta} & \\ & & & & e^{i(-n)\theta} \end{pmatrix}.$$

Therefore, we have that

$$\chi_{V^n} \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix} = e^{in\theta} + e^{i(n-2)\theta} + \cdots + e^{-in\theta} = \frac{e^{i(n+1)\theta} - e^{-i(n+1)\theta}}{e^{i\theta} - e^{-i\theta}}.$$

To see this, one only need to use the identity

$$\left(e^{in\theta} + e^{i(n-2)\theta} + \dots + e^{-i(n-2)\theta} + e^{-in\theta}\right) \left(e^{i\theta} - e^{-i\theta}\right) = e^{i(n+1)\theta} - e^{-i(n+1)\theta}.$$

Suppose now that $n_2 > n_1$, then we have that on $U(1)$,

$$\begin{aligned} \chi_{V^{n_1} \otimes V^{n_2}} &= \chi_{V^{n_1}} \chi_{V^{n_2}} \\ &= (e^{in_1\theta} + e^{i(n_1-2)\theta} + \dots + e^{-in_1\theta}) \frac{e^{i(n_2+1)\theta} - e^{-i(n_2+1)\theta}}{e^{i\theta} - e^{-i\theta}} \\ &= \frac{(e^{i(n_1+n_2+1)\theta} - e^{-i(n_1+n_2+1)\theta}) + \dots + (e^{i(n_2-n_1+1)\theta} - e^{-i(n_2-n_1+1)\theta})}{e^{i\theta} - e^{-i\theta}} \\ &= \chi_{V^{n_1+n_2}} + \chi_{V^{n_1+n_2-2}} + \dots + \chi_{V^{n_2-n_1}}. \end{aligned}$$

But since character is invariant under conjugation, the equality must hold on all of $SU(2)$. Finally, we use the fact that $SU(2) \cong S^3$ is compact. For a compact Lie group, one can always decompose any representation Π into the direct sum of irreducible ones (up to equivalence), with $\Pi \sim \bigoplus_n c_n \Pi_n$, where $c_n \neq 0$ for only finitely many n , and $c_n = \langle \chi_{\Pi_n}, \chi_{\Pi} \rangle$. Therefore, any representation of $SU(2)$ is completely determined by its character. Since we have the equality of the character above, equality of representation must also follow. $\square$

Remark 5.42. The inner product of character mentioned above can be defined as

$$\langle \chi_{\Pi_1}, \chi_{\Pi_2} \rangle = \int_{SU(2)} \overline{\chi_{\Pi_1}}(g) \chi_{\Pi_2}(g) d\mu(g)$$

where $d\mu$ is some invariant volume-measure. Then it can be shown that under this inner product, irreducible representations are mutually orthogonal:

$$\langle \chi_{\Pi_n}, \chi_{\Pi_m} \rangle = \delta_{nm},$$

and every representation on a compact Lie group is completely reducible. This is almost exactly the same as in the finite group representation theory case. For a detailed treatment of representation theory of compact Lie group, consult [4].

5.6.4 Complexification of Vector Space

We conclude with a discussion of complexifying vector spaces. Let V be an $\mathbf{R}$ -vector space. Then intuitively by complexification, we mean that we allow complex coefficients in scalar multiplication, making V a $\mathbf{C}$ -vector space. Formally, we can identify the complexification of V with

$$V \otimes_{\mathbf{R}} \mathbf{C}.$$

Suppose that as a real vector space, V has basis $e_1, \dots, e_n$. Then intuitively, complexifying V will make it a $2n$ -dimensional $\mathbf{R}$ -vector space, with basis elements $e_1, \dots, e_n, ie_1, \dots, ie_n$. But this is an n -dimensional $\mathbf{C}$ -vector space with basis $e_1, \dots, e_n$, but over $\mathbf{C}$. Now consider the identification

$$V_{\mathbf{C}} \ni e_j \leftrightarrow e_j \otimes 1 \in V \otimes \mathbf{C}.$$

One can show that this is an isomorphism. Then the scalar multiplication by a complex number $\alpha \in \mathbf{C}$ is now given by the identification

$$\alpha v = \alpha \left(\sum_{j=1}^n v_j e_j \right) \leftrightarrow \sum_{j=1}^n e_j \otimes (\alpha v_j).$$

5.7 Representation of the Group $(\mathbf{R}, +)$

Theorem 5.43. *Irreducible representations of $(\mathbf{R}, +)$ are all given by*

$$\pi_c(a) = e^{ca} \in GL(\mathbf{C}), \quad c \in \mathbf{C}.$$

Proof. This follows exactly from the proof of irreducible representation of $U(1)$ in pg 20-21 of [3], except we forget about the periodicity condition. $\square$

The Lie group $(\mathbf{R}, +)$ gives some of the most fundamental symmetries in nature. For example, for a function $\psi : \mathbf{R} \rightarrow \mathbf{C}$, it is said to be translational symmetric if $\psi(q + a) = \psi(q)$ for all $a \in \mathbf{R}$. In fact, if we consider the Lie group $(\mathbf{R}, +)$ acting on itself by translation

$$q \rightarrow a \cdot q = a + q,$$

this induces a representations on functions taking real values, given by the usual

$$\pi(a)\psi(q) = \psi(a^{-1} \cdot q) = \psi(q - a).$$

On L^2 functions, this representation is clearly unitary, for

$$\langle \pi(a)f, \pi(a)g \rangle = \int_{-\infty}^{\infty} \bar{f}(q - a)g(q - a)dq = \int_{-\infty}^{\infty} \bar{f}(q)g(q) = \langle f, g \rangle.$$

And we have seen, many times, that given a unitary representation, its induced Lie algebra representation is skew-symmetric, from which we can obtain a self-adjoint operator, which corresponds to some physical observable. Therefore, we want to find π' of the representation above.

What we get is:

$$\begin{aligned} \pi'(a)f(q) &= \frac{d}{dt}\pi(e^{ta})f(q)|_{t=0} = \frac{d}{dt}f(q - e^{ta})|_{t=0} = \frac{d}{dt}f(q - ta)|_{t=0} \\ &= -a \frac{d}{dq}f(q). \end{aligned}$$

Note that in this case we have $e^{ta} = ta$. This is because the exponential function here is quite different from the exponential function in the real analysis case. A detailed discussion can be found [here](#). But to be concise, this is because an element a in the Lie group $G = \mathbf{R}$ can be identified with the matrix group H element

$$a \leftrightarrow \begin{pmatrix} 1 & a \\ 0 & 1 \end{pmatrix} \in H.$$

By definition, the Lie algebra $\mathfrak{h}$ of this matrix group H is all the matrices whose exponential is in the group, and we find that

$$\exp \begin{pmatrix} 0 & a \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 1 & a \\ 0 & 1 \end{pmatrix}. \quad (5.2)$$

Therefore, the Lie algebra of this matrix Lie group H which we identify as $\mathbf{R}$ is given by the set of matrices of the form

$$\mathfrak{h} = \left\{ \begin{pmatrix} 0 & a \\ 0 & 0 \end{pmatrix} : a \in \mathbf{R} \right\} \cong \mathbf{R} \quad \mathfrak{h} \ni \begin{pmatrix} 0 & a \\ 0 & 0 \end{pmatrix} \leftrightarrow a \in \mathbf{R}.$$

Therefore (5.2) can be written, counter-intuitively, as

$$e^a = a \quad a \in \mathbf{R}.$$

In conclusion, we get that

$$\pi'(a) = -a \frac{d}{dq}$$

which is a skew-adjoint operator. Multiplying by i we get a self-adjoint operator

$$-ia \frac{d}{dq}.$$

5.7.1 Hamiltonian and Momentum Operator

Hamiltonian and Time Translation

One of the axiom of QM, the Schrödinger equation, says that for all state $|\psi\rangle$, we have

$$i\hbar \frac{d}{dt} |\psi\rangle = H |\psi\rangle.$$

Therefore

$$H = i\hbar \frac{d}{dt}.$$

Indeed, we see that the self-adjoint operator $\pi'(a)$ is proportional to H up to some scaling by some real constant (which doesn't matter) if we consider the time coordinate $q = t$. This suggests that the Hamiltonian is the generator of the time-translation operator. Indeed, one has

$$\Pi(a)\psi(t) = \psi(t-a) = \sum_{n=0}^{\infty} \frac{1}{n!} (-a)^n \frac{d^n}{dt^n} \psi(t) = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\frac{-a}{i\hbar} H \right)^n \psi(t),$$

which suggests that

$$\Pi(a) = \exp \left(\frac{i}{\hbar} a H \right).$$

But by changing a to t , one see that we covered exactly the propagator

$$\Pi(-t) = U(t) = e^{-\frac{i}{\hbar} t H}$$

that solves the Schrödinger equation with $|\psi(t)\rangle = U(t) |\psi(0)\rangle$.

Momentum and Space Translation

Similarly, if now we choose our coordinate $q = x$ to be spatial, then the operator $\pi(a) = -a \frac{d}{dx}$ will correspond to spatial translation. One can define the momentum operator

$$P = -i\hbar \frac{\partial}{\partial x}$$

which is just the self-adjoint $i\pi(a)$ up to some real constant multiple. Similar to the derivative above, we find that

$$\pi(a) = \exp\left(-\frac{ia}{\hbar}P\right).$$

We say that momentum is the generator of translations. In more than 1 dimensions, the momentum operator can be defined componentwise by $P_j = -i\hbar\frac{\partial}{\partial x_j}$, hence the total momentum operator can be given by $-i\hbar\nabla$.

The Schrödinger Equation for Free Particles

For a free particle described by a state $\psi(\mathbf{x}, t)$, the energy momentum relation reads

$$H = \frac{|P|^2}{2m}$$

with no potentials. Therefore, the Schrödinger equation can be written as follows by plugging in $H = i\hbar\frac{\partial}{\partial t}$ and $P = i\hbar\nabla$:

$$i\hbar\frac{\partial}{\partial t}\psi(\mathbf{x}, t) = -\frac{\hbar^2}{2m}\nabla^2\psi(\mathbf{x}, t). \quad (5.3)$$

This is the Schrödinger equation of a free particle. By separation of variable, one can find that the solution to this equation is given by

$$\psi(\mathbf{x}, t) = \psi_E(\mathbf{x})e^{-\frac{i}{\hbar}tE}$$

where $\psi_E(\mathbf{x})$ satisfies

$$H\psi_E(\mathbf{x}) = -\frac{\hbar^2}{2m}\nabla^2\psi_E(\mathbf{x}) = E\psi_E(\mathbf{x}).$$

The solution is given by

$$\psi_E(\mathbf{x}) = e^{i(k_1x_1+k_2x_2+k_3x_3)} = e^{i\mathbf{k}\cdot\mathbf{x}}$$

where

$$\frac{\hbar^2|\mathbf{k}|^2}{2m} = E.$$

Therefore, we have found the stationary solutions to (5.3), which are labeled by a vector $|\mathbf{k}\rangle = e^{i\mathbf{k}\cdot\mathbf{x}}$, satisfying

$$P_j|\mathbf{k}\rangle = \hbar k_j|\mathbf{k}\rangle \quad H|\mathbf{k}\rangle = \frac{\hbar^2|\mathbf{k}|^2}{2m}|\mathbf{k}\rangle.$$

Then a general time-dependent state $\Psi(\mathbf{x}, t)$ satisfying (5.3) will therefore be a linear combination of stationary states:

$$\Psi(\mathbf{x}, t) = \int d^3\mathbf{k} \phi(\mathbf{k})e^{i\mathbf{k}\cdot\mathbf{x}}e^{-it\hbar\frac{|\mathbf{k}|^2}{2m}}. \quad (5.4)$$

where $\phi(k)$ is some function depending smoothly on $\mathbf{k}$. Actually, (5.4) poses a much more severe problem than you think, so severe that we have to redefine the spaces in which our quantum theory takes place. To find a suitable explanation of (5.4), we have to invoke a lot of nontrivial functional analysis, which will be discussed in the next section.

5.8 Free Particles

In the previous chapter, we have solved the Schrödinger equation for a free particle (5.3). To simplify our discussion, we will consider only 1-dimensional case, but the 3-dimensional case can be generalized quite straightforwardly. In 1-dimension (5.3) reads

$$i\hbar \frac{\partial}{\partial t} \psi(x, t) = -\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} \psi(x, t).$$

The stationary solutions are given by plane waves

$$|k\rangle = e^{ikx}.$$

To be a physically realizable state, it has to be normalizable, since the amplitude $|\psi(x, t)|^2$ of a physically realizable state $\psi(x, t)$ is the probability density function of particle. However,

$$\langle k|k\rangle = \int dx e^{-ikx} e^{ikx} = \int dx = \infty.$$

The non-normalizable nature of the eigenfunctions $|k\rangle$ of P, H suggests the following problems:

- Physically, this means that a free particle cannot have stationary momentum and energy.
- Mathematically, this means the eigenvectors of the operator P, H are not in $L^2(\mathbf{R})$. Therefore, it suggests that we have to define quantum theory on a broader space than $L^2(\mathbf{R})$.

To solve this problem, there are two ways to go. First is to pose periodicity condition $\psi(x + 2\pi) = \psi(x)$ for all x , thus reducing problem to continuous functions on S^1 , which will have convergent $L^2(S^1)$ norm. Detailed treatment using Fourier analysis on S^1 can be found in Chapter 11 of [3]. A somewhat equivalent treatment of this topic is given in section 2.2 of [6], which is more commonly known as the infinite square well problem. We will not be discussing this case despite its important physical interests. The main issue we will focus on is resolving the difficulties that present on the entire $\mathbf{R}$.

5.8.1 Schwartz Space and Tempered Distributions

The solution to the difficult problem that the momentum operator P does not have eigenvalue in $L^2(\mathbf{R})$ is to consider a slightly more general class than just L^2 -functions, called the space of tempered distributions, in which P will have eigenvectors.

Definition 5.44. The Schwartz space $\mathcal{S}(\mathbf{R})$ consists of smooth functions $f : \mathbf{R} \rightarrow \mathbf{C}$ such that the function f itself and its β -th derivative $D^\beta f$ fall off faster than any polynomial functions. Formally speaking, $f \in \mathcal{S}(\mathbf{R})$ if and only if

$$\sup_{x \in \mathbf{R}} |x^\alpha D^\beta f| < \infty \quad \forall \alpha, \beta.$$

Proposition 5.45. We have $\mathcal{S}(\mathbf{R}) \subset L^p(\mathbf{R})$ for $1 \leq p \leq \infty$.

Proof. Take $f \in \mathcal{S}(\mathbf{R})$. Denote $\|f\|_{\alpha\beta} = \sup_{x \in \mathbf{R}} |x^\alpha D^\beta f|$. Then we have

$$\int |f(x)|^p = \int ((1+x^2)|f(x)|)^p \frac{1}{(1+x^2)^p}.$$

Since f is a Schwartz function, $(1+x^2)f$ is bounded, and $\frac{1}{(1+x^2)^p} \leq \frac{1}{1+x^2} \in L^1$ for $p \geq 1$. Hence we have that

$$\int |f(x)|^p \leq \|(1+x^2)f\|_\infty^p \int \frac{1}{1+x^2} \leq \pi (\|f\|_{0,0} + \|f\|_{2,0})^p < \infty.$$

This concludes the proof. In particular, we have $\mathcal{S}(\mathbf{R}) \subset L^2(\mathbf{R})$. $\square$

Now, let V^* denote the dual space of a vector space V , which are all linear functionals on V . Then $L^2(\mathbf{R}) \cong (L^2(\mathbf{R}))^*$ since $L^2(\mathbf{R})$ is a Hilbert space. To see more detailed treatment one can refer to [7]. But the basic idea is, there is a conjugate-linear isomorphism from L^2 to its dual given by

$$f \mapsto \int \bar{f}(\cdot) \in (L^2)^*.$$

In Dirac notation, this is simply written as

$$|f\rangle \mapsto \langle f|.$$

Then identifying L^2 with its dual, there is a natural inclusion

$$\mathcal{S}(\mathbf{R}) \subset L^2(\mathbf{R}) \subset \mathcal{S}'(\mathbf{R}) \quad (5.5)$$

where $\mathcal{S}'(\mathbf{R}) = \mathcal{S}(\mathbf{R})^*$, and the inclusion $L^2(\mathbf{R}) \subset \mathcal{S}'(\mathbf{R})$ are given precisely by $|f\rangle \mapsto \langle f|$. And this inclusion is essentially a consequence of the Hölder's inequality, which is explained in [7]:

Proposition 5.46. We have a natural inclusion $L^p(\mathbf{R}) \subset \mathcal{S}'(\mathbf{R})$, for $p \geq 1$.

Proof. Identifying $f \in L^p$ with $\int \bar{f}\phi d\mu$ where $\phi \in \mathcal{S}(\mathbf{R})$, we have that

$$\left| \int \bar{f}(x)\phi(x)d\mu \right| \leq \|f\|_{L^p} \cdot \|(1+|x|^2)^{-1}\|_{L^{p'}} \cdot \|(1+|x|^2)\phi\|_\infty < \infty$$

since $\phi \in \mathcal{S}(\mathbf{R})$. In fact, this is even true in the $\mathbf{R}^n$ case, where we have

$$\left| \int_{\mathbf{R}^n} f(x)\phi(x)d\mu \right| \leq \|f\|_{L^p} \cdot \|(1+|x|^2)^{-n}\|_{L^{p'}} \cdot \|(1+|x|^2)^n\phi\|_\infty < \infty.$$

This proves the claim. $\square$

Now that we have verified (5.5), we consider a quantum theory defined not on $L^2(\mathbf{R})$, but on $\mathcal{S}'(\mathbf{R})$. We make the following definition:

Definition 5.47. Let $F \in \mathcal{S}'(\mathbf{R})$ and A be an operator on $\mathcal{S}(\mathbf{R})$. We say that F is a generalized eigenvector of A with eigenvalue λ if and only if

$$(AF)(\varphi) := F(A^\dagger \varphi) = (\lambda F)(\phi) \quad \forall \varphi \in \mathcal{S}(\mathbf{R}).$$

Now we consider the plane wave solution $|k\rangle = e^{ikx}$. It sits naturally in $\mathcal{S}'(\mathbf{R})$, which is given by $\langle k|(|\phi\rangle) = \int e^{-ikx}\phi(x) dx$. It can be shown that $P(\mathcal{S}(\mathbf{R})) \subset \mathcal{S}(\mathbf{R})$. Hence the restriction of the momentum operator P on $\mathcal{S}(\mathbf{R})$ gives an operator on $\mathcal{S}(\mathbf{R})$. Now we verify that $|k\rangle$ (which is really $\langle k|$ identified) is a generalized eigenvector of P .

$$\begin{aligned} (P\langle k|)|\varphi\rangle &= \langle k|P^\dagger|\varphi\rangle = \int_{-\infty}^{\infty} e^{-ikx} \left(-i\hbar \frac{\partial}{\partial x} \varphi\right) dx \\ &= -i\hbar \left(\int_{-\infty}^{\infty} \frac{\partial}{\partial x} [e^{-ikx} \varphi(x)] dx - \int_{-\infty}^{\infty} \varphi \frac{\partial}{\partial x} e^{-ikx} dx \right) = \hbar k \langle k|\varphi\rangle. \end{aligned}$$

Note that we have used $P = P^\dagger$. This shows that we can indeed think of $|k\rangle$ as eigenvectors of P , sitting inside $\mathcal{S}'(\mathbf{R})$. Sometimes, it is more convinient to write

$$|k\rangle = \frac{1}{\sqrt{2\pi}} e^{ikx}$$

and write

$$\widehat{\phi}(k) = \langle k|\phi\rangle = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \phi(x) dx.$$

One sees that this is nothing but the Fourier transform of the Schwarz function ϕ .

There is another important linear functional in $\mathcal{S}'(\mathbf{R})$, which we write δ_y . It is defined by

$$\delta_y(f) = f(y).$$

Motivated by the beauty of the Dirac notation, we want to write δ_y as some function that we can integrate against any Schwarz function, i.e., some function g such that $\delta_y(f) = f(y) = \langle g|f\rangle$. Unfortunately, there is no measurable function g that will make this work. However, we can still write (notationally)

$$f(y) = \int_{-\infty}^{\infty} \delta(x-y) f(x) dx.$$

But it is important to note that $\delta(x-y)$ is not a function. Notationally, one can write $\langle \delta_y|$. We claim that this is a generalized eigenvector with eigenvalue y of the position operator X defined by

$$X\varphi(x) = x\varphi(x)$$

which is multiplication by the function x . The proof is straightforward:

$$(X\langle \delta_y|)(|\phi\rangle) = \langle \delta_y|X^\dagger|\phi\rangle = \langle \delta_y|x\phi(x)\rangle = y\phi(y) = y\langle \delta_y|\phi\rangle.$$

Therefore, we see that although P, X does not have eigenvectors in L^2 space, we can extend the space to generalized distributions in which they do have eigenvalues, and in which L^2 naturally sits.

In general, given a Hilbert space $\mathcal{H}$ and an algebra of operators $\mathcal{A}$ on $\mathcal{H}$, there exists a maximal subspace $\mathcal{S} \subset \mathcal{H}$ on which $\mathcal{A}$ acts. Then we have the natural inclusion $\mathcal{S} \subset \mathcal{H} \subset \mathcal{S}^*$. Although $A \in \mathcal{A}$ won't have eigenvectors in $\mathcal{H}$ if it has continuous spectra, it will have eigenvectors in $\mathcal{S}^*$. A triple

$$(\mathcal{S}, \mathcal{H}, \mathcal{S}^*)$$

is called a rigged Hilbert space. It turns out that rigged Hilbert space is a natural place in which quantum mechanics takes place. A detailed survey can be found in [8].

5.8.2 Fourier Transform

Definition 5.48. The Fourier transform of a function $\psi \in \mathcal{S}(\mathbf{R})$ is given by a function $\widehat{\psi}$, where

$$\widehat{\psi}(k) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \psi(x) dx.$$

One can show that $\widehat{\psi} \in \mathcal{S}(\mathbf{R})$, and define the inverse Fourier transform by

$$\psi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \widehat{\psi}(k) e^{ikx} dk.$$

Remark 5.49. One can consider the Fourier transform and its inverse transform using the following analogy: In a finite dimensional vector space V , any vector can be written as $|v\rangle = \sum_j v(j) |e_j\rangle$ where $|e_j\rangle$ is an orthonormal basis: $\langle e_i | e_j \rangle = \delta_{ij}$. The coefficients $v(j)$ can be obtained by $v(j) = \langle e_j | v \rangle$. In the continuous case where V is infinite dimensional, the set of functions $|k\rangle = e^{ikx}$ forms an orthonormal basis (we'll explain later) with $\langle k | k' \rangle = \delta(k - k')$. Then we can view $\widehat{\psi}(k) = \langle k | \psi \rangle$ just as $v(j)$ in the discrete case. And the inverse Fourier transform is just the continuous version of reconstructing $|v\rangle = \sum_j v(j) |e_j\rangle$ with the knowledge of the coefficients $v(j)$.

Proposition 5.50. The Fourier transform convert differential operator D into multiplication. More specifically,

$$\widehat{D\psi}(k) = ik\widehat{\psi}(k).$$

Proof. This is simply proven by

$$\begin{aligned} \widehat{D\psi}(k) &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \frac{d}{dx} \psi dx \\ &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \left(\frac{d}{dx} (e^{-ikx} \psi) - \left(\frac{d}{dx} e^{-ikx} \right) \psi \right) dx \\ &= ik \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \psi dx = ik\widehat{\psi}(k). \end{aligned}$$

We used the fact that $\psi(\pm\infty) = 0$ and $|e^{ikx}| = 1$. □

Proposition 5.51. Let T_a be the translation operator defined by $T_a\psi(x) = \psi(x + a)$. Then

$$\widehat{T_a\psi}(k) = e^{ika}\widehat{\psi}(k).$$

Proof. We have

$$\begin{aligned} \widehat{T_a\psi}(k) &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \psi(x + a) dx = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ik(x'-a)} \psi(x') dx' \\ &= e^{ika}\widehat{\psi}(k). \end{aligned}$$

This concludes the proof. □

Theorem 5.52 (Plancherel). *The Fourier transform and inverse Fourier transform extends to unitarity isomorphisms of $L^2(\mathbf{R})$ with itself. In particular,*

$$\|\psi(x)\|_{L^2} = \|\widehat{\psi}(k)\|_{L^2}.$$

Proof. Refer to any standard textbook in Fourier analysis. $\square$

Definition 5.53 (Fourier transform on $\mathcal{S}'(\mathbf{R})$). Since the Fourier transform is defined on Schwartz function space $\mathcal{S}(\mathbf{R})$, we can also define Fourier transform on the space of tempered distributions $\mathcal{S}'(\mathbf{R})$. Let $T \in \mathcal{S}'(\mathbf{R})$ and $\phi \in \mathcal{S}(\mathbf{R})$. We define

$$(\widehat{T})(\phi) = T(\widehat{\phi}).$$

If we write $T(\phi) = \langle T, \phi \rangle$, then the Fourier transform $\widehat{T}$ of T is defined by

$$\langle \widehat{T}, \phi \rangle = \langle T, \widehat{\phi} \rangle.$$

Example 5.54. One of the most important examples of Fourier transformation of distributions is that of δ_0 . We have that

$$\langle \widehat{\delta_0}, \phi \rangle = \langle \phi, \widehat{\delta_0} \rangle = \phi(0).$$

From the formula of the Fourier transform, we have that

$$\widehat{\phi}(0) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \phi(x) dx = \frac{1}{\sqrt{2\pi}} \langle 1, \phi \rangle$$

where 1 is understood as the functional $\phi \mapsto \int \phi dx$. Therefore, the Fourier transform of δ_0 is the constant function

$$\widehat{\delta_0} = \frac{1}{\sqrt{2\pi}}.$$

Using the formula for the inverse transform, one gets

$$\delta(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{ikx} dk.$$

Replacing $x \rightarrow x - x'$, we have

$$\delta(x - x') = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{ik(x-x')} dk.$$

Then we consider the eigenvectors of the momentum operation P :

$$|k\rangle = \frac{1}{\sqrt{2\pi}} e^{ikx}.$$

We have that

$$\langle k' | k \rangle = \int_{-\infty}^{\infty} \left(\frac{1}{\sqrt{2\pi}} e^{-ik'x} \right) \left(\frac{1}{\sqrt{2\pi}} e^{ikx} \right) dx = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{i(k-k')x} dx = \delta(k - k').$$

5.8.3 Dirac Notation, Position and Momentum Space

We have been seeing Dirac notation a lot. But there is still something worth mentioning about it. Suppose in a finite dimensional vector space with an orthonormal basis $|e_j\rangle$, one can write $|v\rangle = \sum_j v(j) |e_j\rangle$. Since $v(j) = \langle e_j | v \rangle$, one can actually write

$$|v\rangle = \sum_j \langle e_j | v \rangle |e_j\rangle = \left(\sum_j |e_j\rangle \langle e_j| \right) |v\rangle.$$

Hence, we get that

$$\sum_j |e_j\rangle\langle e_j| = \mathbf{1}$$

if $|e_j\rangle$ is an orthonormal basis. In the infinite dimensional case, things are more subtle. If a self-adjoint operator Q has continuous spectra, then one can derive an analogous version of the finite-dimensional case through some hard work. For a rigorous treatment, see [5], [8], and [7]. But the idea is straightforward, define eigenvector $|q\rangle$ of the operator Q such that $Q|q\rangle = q|q\rangle$. Then spectral theorem for unbounded operator says that one can find a set of eigenvectors $|q\rangle$ that form an orthonormal basis of our state space. Then replacing the sum by an integral, we get

$$\int dq |q\rangle\langle q| = \mathbf{1}.$$

Let us study the case where our state space is $\mathcal{S}'(\mathbf{R})$, and the self-adjoint operators are position and momentum operators. Then a state $|\psi\rangle$ can be written as²

$$|\psi\rangle = \left(\int dx |x\rangle\langle x| \right) |\psi\rangle = \int dx \langle x|\psi\rangle |x\rangle.$$

Defining $\langle x|\psi\rangle = \psi(x)$, this boils down to the equality

$$\psi(y) = \int dx \psi(x) \delta(y - x).$$

Therefore, what we refer to as wave function $\psi(x)$ is really just the component of the abstract state $|\psi\rangle$ with respect to the basis $|x\rangle$, which are eigenvectors of the operator X . Similarly, one can consider the case of operator P and write

$$|\psi\rangle = \left(\int dk |k\rangle\langle k| \right) |\psi\rangle = \int dk \langle k|\psi\rangle |k\rangle. \quad (5.6)$$

Let us look at (5.6) more carefully. The coefficients $\langle k|\psi\rangle$ are exactly the Fourier coefficients, since

$$\hat{\psi}(k) := \langle k|\psi\rangle = \frac{1}{\sqrt{2\pi}} \int e^{-ikx} \psi(x) dx.$$

Therefore, the Fourier transform formula is written compactly inside $\langle k|\psi\rangle$! Now expanding (5.6) into a more familiar form, we find that (5.6) is nothing but

$$\psi(x) = \frac{1}{\sqrt{2\pi}} \int \hat{\psi}(k) e^{ikx} dk$$

which is the inverse Fourier transform! Therefore, given a state $|\psi\rangle$ in state space, its components $\psi(x)$ with respect to $|x\rangle$ basis is different from its component $\hat{\psi}(k)$ with respect to $|k\rangle$ basis. But $\psi(x), \hat{\psi}(k)$ refer to the same vector $|\psi\rangle$. This is just a consequence of a change of basis. Therefore, we have the elegant correspondence between position space and momentum space of wave functions as the Fourier transform from signal space to frequency space.

²The reason we can conclude $\int dx |x\rangle\langle x| = \mathbf{1}$ is not because $|x\rangle$ is an orthonormal basis. In fact, there is no inner product on $\mathcal{S}'(\mathbf{R})$. The reason we can do this is simply because of the definition of the Dirac delta function, and its corresponding integral representation in the next line.

Therefore, we call functions $\widehat{\psi}(k)$ to be in momentum space, and $\psi(x)$ to be in position space.

In the position space, operators X, P are given by

$$X\psi(x) = x\psi(x) \quad P\psi(x) = -i\hbar \frac{\partial}{\partial x} \psi(x).$$

But in the momentum space, we are now taking the basis to be $|k\rangle$, so X, P will also change. Assuming our state space $\mathcal{H}$ is nice enough, the Fourier transform provides an isomorphism $\mathcal{F} : \mathcal{H} \rightarrow \widehat{\mathcal{H}} = \mathcal{H}$, where we label the hat to suggest that we are working in the momentum space. Then we have the commutative diagram

$$\begin{array}{ccc} \widehat{\mathcal{H}} & \xrightarrow{\widehat{P}} & \widehat{\mathcal{H}} \\ \mathcal{F}^{-1} \downarrow & & \uparrow \mathcal{F} \\ \mathcal{H} & \xrightarrow{P} & \mathcal{H} \end{array}$$

Hence we have $\widehat{P} = \mathcal{F} \circ P \circ \mathcal{F}^{-1}$. Therefore, in the momentum space, the action of momentum operator is given by

$$\begin{aligned} \widehat{P}\widehat{\psi}(k) &= (\mathcal{F} \circ P \circ \mathcal{F}^{-1})(\widehat{\psi}(k)) \\ &= \mathcal{F} \circ P(\psi(x)) \\ &= \mathcal{F}(-i\hbar D\psi(x)) \\ &= (-i\hbar)(ik)\widehat{\psi}(k) \\ &= \hbar k\widehat{\psi}(k) \end{aligned}$$

where we have used $\widehat{D}\widehat{\psi}(k) = ik\widehat{\psi}(k)$. Therefore, we conclude that in the momentum space

$$\widehat{P} = \hbar k.$$

Similarly we have a commutative diagram

$$\begin{array}{ccc} \widehat{\mathcal{H}} & \xrightarrow{\widehat{X}} & \widehat{\mathcal{H}} \\ \mathcal{F}^{-1} \downarrow & & \uparrow \mathcal{F} \\ \mathcal{H} & \xrightarrow{X} & \mathcal{H} \end{array}$$

Therefore we have that

$$\begin{aligned} \widehat{X}\widehat{\psi}(k) &= (\mathcal{F} \circ X \circ \mathcal{F}^{-1})(\widehat{\psi}(k)) \\ &= \mathcal{F} \circ X(\psi(x)) \\ &= \mathcal{F}(x\psi(x)) \\ &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} x\psi(x) dx \\ &= \left(i \frac{d}{dk}\right) \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-ikx} \psi(x) dx \\ &= \left(i \frac{d}{dk}\right) \widehat{\psi}(k). \end{aligned}$$

Therefore, we conclude that in the momentum space, the position operator acts as

$$\hat{X} = i \frac{d}{dk}.$$

There is a derivation of the position operator in momentum space in [6] using the Dirac notation.

5.9 Representations of Euclidean Group

In this chapter, we will first explore the solution to the Schrödinger equation of a free particle using Fourier transform, and then we will digress a little bit to study Euclidean groups $E(2)$. After that, we will show that the time independent solutions to the SDE of free particles are the elements of irreducible representations of the Euclidean group $E(2)$. We will finally generalize this situation to the 3-dimensional case.

5.9.1 Solution to the Free Particle SDE in Momentum Space

Recall that the Schrödinger equation of the free particle in 2-dimension is given by

$$i\hbar \frac{\partial}{\partial t} \psi(\mathbf{q}, t) = -\frac{\hbar^2}{2m} \nabla^2 \psi(\mathbf{q}, t). \quad (5.7)$$

We have seen that Fourier transform takes the derivative operator to multiplication by k in one dimension. Similarly, in three dimension, after Fourier transforming (5.7), one gets (setting $\hbar = 1$ and replacing coordinate k_1, k_2, k_3 with p_1, p_2, p_3)

$$i \frac{\partial}{\partial t} \hat{\psi}(\mathbf{p}, t) = \frac{1}{2m} (p_1^2 + p_2^2) \hat{\psi}(\mathbf{p}, t).$$

This equation can then be solved using separation of variable, giving solutions

$$\hat{\psi}(\mathbf{p}, t) = e^{-i \frac{1}{2m} |\mathbf{p}|^2 t} \hat{\psi}(\mathbf{p}, 0).$$

We want to find states that not only obeys the Schrödinger equation, but is also the momentum eigenstate. The momentum eigenstate $\hat{\psi}(\mathbf{p}, t)$ at $t = 0$ obeys $p_j \hat{\psi}(\mathbf{p}, 0) = p'_j \hat{\psi}(\mathbf{p}, 0)$ acted by each momentum component P_j . Therefore, the solution is given by a delta distribution

$$\hat{\psi}(\mathbf{p}, 0) = \delta(\mathbf{p} - \mathbf{p}') = \delta(p_1 - p'_1) \delta(p_2 - p'_2).$$

The position space wave function can be recovered from the Fourier inversion formula

$$\psi(\mathbf{q}, t) = \frac{1}{2\pi} \int d^2 \mathbf{p} e^{i \mathbf{p} \cdot \mathbf{q}} \hat{\psi}(\mathbf{p}, t).$$

5.9.2 Semidirect Products and The Euclidean Group

Definition 5.55 (Semidirect Products). Let N, K be two Lie groups. Suppose for every $k \in K$, there is an automorphism (isomorphism from a group to itself) $\Phi(k) = \Phi_k : N \rightarrow N$. Then one can define the semidirect product

$$N \rtimes_{\Phi} K$$

of N and K to be the group with elements $(n, k) \in N \times K$ with the group law

$$(n_1, k_1)(n_2, k_2) = (n_1\Phi_{k_1}(n_2), k_1k_2).$$

One can check that this is indeed a group, with inverse element

$$(n, k)^{-1} = (\Phi_{k^{-1}}(n^{-1}), k^{-1}).$$

Definition 5.56 (Euclidean Group). The Euclidean group $E(d)$ is the semidirect product

$$E(d) = \mathbf{R}^d \rtimes_{\Phi} SO(d)$$

where the automorphism is given by the identity map

$$\Phi(R) = R : \mathbf{R}^d \rightarrow \mathbf{R}^d.$$

More specifically, the multiplication is given by

$$(\mathbf{a}_1, R_1)(\mathbf{a}_2, R_2) = (\mathbf{a}_1 + R_1\mathbf{a}_2, R_1R_2).$$

There is a natural action of $E(d)$ on $\mathbf{R}^d$, given by

$$(\mathbf{a}, R) \cdot \mathbf{v} = \mathbf{a} + R\mathbf{v}.$$

One can check that this is indeed a group action. Note that we can identify an element in $E(d)$ with a $(d+1) \times (d+1)$ matrix, making $E(d)$ a matrix Lie group by

$$(\mathbf{a}, R) \leftrightarrow \begin{pmatrix} R & \mathbf{a} \\ \mathbf{0} & 1 \end{pmatrix}.$$

The matrix multiplication is consistent with group multiplication of $E(d)$, since

$$\begin{pmatrix} R_1 & \mathbf{a}_1 \\ \mathbf{0} & 1 \end{pmatrix} \begin{pmatrix} R_2 & \mathbf{a}_2 \\ \mathbf{0} & 1 \end{pmatrix} = \begin{pmatrix} R_1R_2 & \mathbf{a}_1 + R_1\mathbf{a}_2 \\ \mathbf{0} & 1 \end{pmatrix}.$$

Given a semidirect product of two Lie groups, a natural question is to find its corresponding Lie algebra. This turns out to be quite a difficult question in general. But in the case of $E(d)$, we can do it. Recall that the Lie algebra consists of matrices $Y \in \mathfrak{e}(d)$, such that when we exponentiate, we have $e^{tY} \in E(d)$. In this case, we have that

$$\mathfrak{e}(d) = \left\{ \begin{pmatrix} X & \mathbf{a} \\ \mathbf{0} & 0 \end{pmatrix} : X \in \mathfrak{so}(d), \mathbf{a} \in \mathbf{R}^d \right\}.$$

And one can compute the commutator and find that

$$\left[\begin{pmatrix} X_1 & \mathbf{a}_1 \\ \mathbf{0} & 0 \end{pmatrix}, \begin{pmatrix} X_2 & \mathbf{a}_2 \\ \mathbf{0} & 0 \end{pmatrix} \right] = \begin{pmatrix} [X_1, X_2] & X_1\mathbf{a}_2 - X_2\mathbf{a}_1 \\ \mathbf{0} & 0 \end{pmatrix}.$$

Hence, in simple Lie algebra pair, one has that $\mathfrak{e}(d) = \mathbf{R}^d \oplus \mathfrak{so}(d)$, with Lie bracket

$$[(\mathbf{a}_1, X_1), (\mathbf{a}_2, X_2)] = (X_1\mathbf{a}_2 - X_2\mathbf{a}_1, [X_1, X_2]).$$

In fact, this relation can be generalized, as we will see in later sections.

Now we turn to the case where $d = 2$ and study the Lie algebra $\mathfrak{e}(2)$ of $E(2)$. Since any element $R \in SO(2)$ is of the form

$$R = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix},$$

and

$$\begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix} = \exp \begin{pmatrix} 0 & \theta \\ -\theta & 0 \end{pmatrix},$$

we see that the matrix $\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$ is a basis for $\mathfrak{so}(2)$. Therefore, we see that

$$l = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \quad p_1 = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \quad p_2 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

form a basis of $\mathfrak{e}(2)$. Notice that we call these matrices l, p_1, p_2 not by coincidence: They satisfy the commutation relation

$$[l, p_1] = p_2 \quad [l, p_2] = -p_1 \quad [p_1, p_2] = 0.$$

If we identify the matrix l with the angular momentum function $l = q_1 p_2 - q_2 p_1$ on the phase space of (q_1, q_2, p_1, p_2) , using the Poisson bracket, we see that exactly the same relation is satisfied:

$$\{l, p_1\} = p_2 \quad \{l, p_2\} = -p_1 \quad \{p_1, p_2\} = 0.$$

Therefore, $\mathfrak{e}(2) \cong \text{span}\{l, p_1, p_2\}$, which is a sub-Lie algebra of the homogeneous polynomials of degree at most 2 in phase space. Therefore, the Schrödinger representation will provide us a representation of Lie algebra $\mathfrak{e}(2)$ on the state space $\mathcal{H}$ of functions of the position variable q_1, q_2 .

5.9.3 Irreducible Representations of $E(2)$

As we have discussed in the previous section, $\mathfrak{e}(2)$ is isomorphic to the Lie algebra with basis $l = q_1 p_2 - q_2 p_1$ and p_1, p_2 as a Lie sub-algebra of functions on phase space $\mathbf{R}^4 = (q_1, q_2, p_1, p_2)$. Using the canonical quantization, the Schrödinger representation Π'_S gives us a Lie algebra representation of $\mathfrak{e}(2)$ on the state space $\mathcal{H}$ of functions of the position variable q_1, q_2 given by

$$\begin{aligned} \Pi'_S(p_1) &= -iP_1 = -\frac{\partial}{\partial q_1} \\ \Pi'_S(p_2) &= -iP_2 = -\frac{\partial}{\partial q_2} \\ \Pi'_S(l) &= -iL = -i(Q_1 P_2 - Q_2 P_1) = -\left(q_1 \frac{\partial}{\partial q_2} - q_2 \frac{\partial}{\partial q_1}\right). \end{aligned}$$

Now, we want to look for subspaces of $L^2(\mathbf{R}^2)$ such that they are invariant under $\Pi'_S(p_1), \Pi'_S(p_2), \Pi'_S(l)$, and no proper subspace of them are invariant under these operators, hence irreducible. Since $E(2) = \mathbf{R}^2 \rtimes SO(2)$ is topologically just $\mathbf{R}^2 \times SO(2)$, which are connected since both components are, exponentiating these irreducible representations of Lie algebra $\mathfrak{e}(2)$ gives irreducible representations of $E(2)$.

The trick is to define the Casimir operator and raising/lowering operator again. Here we follow mostly [15] and define

$$P^2 = P_1^2 + P_2^2 \quad P_{\pm} = P_1 \pm iP_2.$$

Then one can verify that

$$\begin{aligned} [L, P_{\pm}] &= P_{\pm} \\ P^2 &= P_1^2 + P_2^2 = P_- P_+ = P_+ P_- \\ [P^2, L] &= [P^2, P_{\pm}] = 0. \end{aligned}$$

Since $[P^2, L] = 0$, we can diagonalize them simultaneously by finding a set of normal eigenvectors $|p, m\rangle$ of both of them, say

$$P^2 |p, m\rangle = p^2 |p, m\rangle \quad L |p, m\rangle = m |p, m\rangle.$$

Since P^2, L are self-adjoint (as you can check), we can assume $|p, m\rangle$ to be orthonormal. Note that $p^2 \geq 0$, since we have

$$\langle p, m | P^2 | p, m \rangle = p^2 = \langle p, m | P_- P_+ | p, m \rangle = \langle p, m | P_+^{\dagger} P_+ | p, m \rangle$$

so p^2 is the square of the norm of the vector $P_+ |p, m\rangle$. Now consider $P_+ |p, m\rangle$. Using the relation $[L, P_{\pm}] = \pm P_{\pm}$, we have that

$$\begin{aligned} J(P_+ |p, m\rangle) &= (LP_+ - P_+L + P_+L) |p, m\rangle \\ &= P_+ |p, m\rangle + P_+(m |p, m\rangle) \\ &= (m+1)(P_+ |p, m\rangle). \end{aligned}$$

Using $[P^2, P_{\pm}] = 0$, we have

$$P^2(P_+ |p, m\rangle) = P_+ P^2 |p, m\rangle = p^2(P_+ |p, m\rangle).$$

Using $P^2 = P_- P_+$ and $P_{\pm}^{\dagger} = P_{\mp}$, we find that

$$(\langle p, m | P_+^{\dagger})(P_+ |p, m\rangle) = \langle p, m | P_- P_+ |p, m\rangle = \langle p, m | P^2 |p, m\rangle = p^2$$

where we have used the orthonormality condition $\langle p, m | p, m \rangle = 1$. Therefore, we found that $P_+ |p, m\rangle$ has norm p , and is an eigenvector of L with eigenvalue $m+1$, and is an eigenvector of P^2 with eigenvalue p^2 . Hence we conclude that

$$P_+ |p, m\rangle \propto p |p, m+1\rangle$$

up to a phase factor. Similarly, one can show that

$$P_{\pm} |p, m\rangle \propto p |p, m \pm 1\rangle$$

up to a phase factor.

Theorem 5.57. *Let $p > 0$ be given. The vector space*

$$V^p = \text{span}_{\mathbf{C}}\{|p, m\rangle\}_{m=0, \pm 1, \pm 2, \dots}$$

is invariant under $\Pi'_S(p_1), \Pi'_S(p_2), \Pi'_S(l)$, hence is an irreducible Lie algebra representation of $\mathfrak{e}(2)$.

Proof. Clearly V^p is invariant under $P_{\pm}, L$. Hence it is also invariant under P_1, P_2, L . This proves the claim. Consequently, (Π'_S, V^p) is an irreducible representation of $\mathfrak{e}(2)$. Hence (Π_S, V^p) is an irreducible representation of $E(2)$. $\square$

Now, it remains to find explicit expressions of the eigenvectors $|p, m\rangle$ for a given p . This amounts to solving a system of PDEs

$$P^2 |p, m\rangle = p^2 |p, m\rangle \quad L |p, m\rangle = m |p, m\rangle.$$

Or explicitly, witting $|p, m\rangle = \psi_{p,m}(\mathbf{q})$, and introducing $p^2 = 2ME$, the PDEs above reads

$$\begin{aligned} -\frac{1}{2M} \left(\frac{\partial^2}{\partial q_1^2} + \frac{\partial^2}{\partial q_2^2} \right) \psi_{p,m}(\mathbf{q}) &= E \psi(\mathbf{q}) \\ -i \left(q_1 \frac{\partial}{\partial q_2} - q_2 \frac{\partial}{\partial q_1} \right) \psi_{p,m}(\mathbf{q}) &= m \psi(\mathbf{q}). \end{aligned}$$

Notice that the first equation is exactly the solution to the time independent SDE of a free particle with mass M and energy E , whose solution in the momentum space after Fourier transform is given by

$$\widehat{\psi}_{p,m}(\mathbf{p}, t) = e^{-i \frac{1}{2M} |\mathbf{p}|^2 t} \widehat{\psi}(\mathbf{p}, 0)$$

where we define $\widehat{\psi}(\mathbf{p}) = \widehat{\psi}(\mathbf{p}, 0)$. These will have well-defined momentum $\mathbf{p}_0$ when

$$\widehat{\psi}(\mathbf{p}) = \delta(\mathbf{p} - \mathbf{p}_0).$$

The position space wave functions can be recovered from the Fourier inversion formula

$$\psi_{p,m}(\mathbf{q}, t) = \frac{1}{2\pi} \int d^2 \mathbf{p} e^{i \mathbf{p} \cdot \mathbf{q}} \widehat{\psi}_{p,m}(\mathbf{p}, t).$$

Since in the momentum space, the momentum operator becomes the multiplication operator

$$\mathbf{P} \widehat{\psi}(\mathbf{p}) = \mathbf{p} \widehat{\psi}(\mathbf{p})$$

an eigenfunction for the Hamiltonian with eigenvalue E will satisfy

$$\left(\frac{|\mathbf{p}|^2}{2M} - E \right) \widehat{\psi}_{p,m}(\mathbf{p}) = 0.$$

Therefore, $\widehat{\psi}_{p,m}(\mathbf{p})$ can only be nonzero on $E = |\mathbf{p}|^2/2M$. That is, to say the free particle solutions with energy E will thus be parametrized by the distributions that are supported on the circle

$$|\mathbf{p}|^2 = 2ME.$$

Going to polar coordinates $\mathbf{p} = (p \cos \theta, p \sin \theta)$, such solutions are given by distributions of the form

$$\widehat{\psi}_{p,m}(\mathbf{p}) = \widehat{\psi}_E(\theta) \delta(p^2 - 2ME) \quad (5.8)$$

depending on both θ, p . Since one has equality of distribution³

$$\delta(p^2 - 2ME) = \frac{1}{2\sqrt{2ME}}\delta(p - \sqrt{2ME})$$

One has that (setting $t = 0$)

$$\begin{aligned}\psi_{p,m}(\mathbf{q}) &= \frac{1}{2\pi} \int d^2\mathbf{p} e^{i\mathbf{p}\cdot\mathbf{q}} \hat{\psi}_{p,m}(\mathbf{p}) \\ &= \frac{1}{2\pi} \iint e^{i\mathbf{p}\cdot\mathbf{q}} \hat{\psi}_E(\theta) \delta(p^2 - 2ME) p \, dp \, d\theta \\ &= \frac{1}{2\pi} \iint e^{i\mathbf{p}\cdot\mathbf{q}} \hat{\psi}_E(\theta) \frac{1}{2\sqrt{2ME}} \delta(p - \sqrt{2ME}) p \, dp \, d\theta \\ &= \frac{1}{4\pi} \int_0^{2\pi} e^{i\sqrt{2ME}(q_1 \cos \theta + q_2 \sin \theta)} \hat{\psi}_E(\theta) \, d\theta.\end{aligned}$$

To determine what $\hat{\psi}_E(\theta)$ is, we need to use the second PDE, which is $L\psi_{p,m}(\mathbf{q}) = m\psi_{p,m}(\mathbf{q})$. One can show that in the momentum space, the angular momentum operator becomes

$$L = -i \left(p_1 \frac{\partial}{\partial p_2} - p_2 \frac{\partial}{\partial p_1} \right) \xrightarrow{\mathcal{F}} \hat{L} = -i \frac{\partial}{\partial \theta}.$$

Therefore, after Fourier transforming,

$$\hat{L}\hat{\psi}_{p,m} = m\hat{\psi}_{p,m}$$

together with (5.8) gives that

$$\hat{\psi}_E(\theta) = e^{im\theta}.$$

Plugging this back to the integral equation of $\psi_{p,m}(\mathbf{q})$, we find that the wave function along the q_2 direction is given by

$$\psi_{p,m}(0, q) = \frac{1}{4\pi} \int_0^{2\pi} e^{i\sqrt{2ME}(q \sin \theta)} e^{im\theta} \, d\theta = \frac{1}{2} J_{-m}(\sqrt{2ME}q) = \frac{1}{2} J_{-m}(pq)$$

where J_n is the n 'th Bessel function.

Therefore, we have studied the solution to (5.7) in two dimensional case, and found that time independent solutions $\psi_{p,m}(\mathbf{q}) = |p, m\rangle$ with energy $2ME = p^2$ correspond to elements in the irreducible representations (Π_S, V^p) of the Euclidean group $E(2)$.

³See [3] pg 249 for detailed discussion. Refer to (11.10) of [3] and the proof given by this [Stackexchange post](#).

Chapter 6

The Schrödinger Equations

In this chapter, we study some of the most common and important quantum mechanical systems by explicitly solving Schrödinger equations with various potentials.

6.1 Schrödinger Equation in Position Space

The Schrödinger equation in position space reads

$$i\hbar \frac{\partial \Psi}{\partial t} = -\frac{\hbar^2}{2m} \frac{\partial^2 \Psi}{\partial x^2} + V\Psi. \quad (6.1)$$

Now our task is to solve $\Psi(x, t)$ with various potential function $V(x)$. The first attempt is to guess a solution of the form

$$\Psi(x, t) = \psi(x)\varphi(t). \quad (6.2)$$

This gives

$$\frac{\partial \Psi}{\partial t} = \psi \frac{d\varphi}{dt}, \quad \frac{\partial^2 \Psi}{\partial x^2} = \frac{d^2 \psi}{dx^2} \varphi.$$

Substitute above into the Shronddinger Equation we get

$$i\hbar \psi \frac{d\varphi}{dt} = -\frac{\hbar^2}{2m} \frac{d^2 \psi}{dx^2} \varphi + V\psi \varphi.$$

Divide both sides by $\psi\varphi$ gives

$$i\hbar \frac{1}{\varphi} \frac{d\varphi}{dt} = -\frac{\hbar^2}{2m} \frac{1}{\psi} \frac{d^2 \psi}{dx^2} + V. \quad (6.3)$$

Now, the left hand side is a function of t alone whereas the right hand side is a function of x alone. We conclude that *both sides must be equal to a constant E* in order for this equation to hold. Otherwise by varying t we can change the left hand side without changing the right hand side, and the equation will no longer be true.

This gives

$$i\hbar \frac{1}{\varphi} \frac{d\varphi}{dt} = E,$$
$$-\frac{\hbar^2}{2m} \frac{1}{\psi} \frac{d^2 \psi}{dx^2} + V = E.$$

The first equation is easy to solve:

$$\varphi(t) = e^{-iEt/\hbar}. \quad (6.4)$$

The second equation is called the ***time independent Schrödinger equation*** (TISE):

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} + V\psi = E\psi. \quad (6.5)$$

Now recall we defined the energy or hamiltonian operator:

$$\hat{H} = \frac{\hat{p}^2}{2m} + V = -\frac{\hbar^2}{2m} \frac{\partial^2}{\partial x^2} + V$$

Hence the TISE can also be written as

$$\hat{H}\psi = E\psi. \quad (6.6)$$

Note that $\hat{H}$ is an operator on a vector ψ (the set of all continuous functions defined on the real line forms a vector space). On the right hand side we have a constant E times that vector ψ . This should remind you of something: if $\hat{H}$ were a matrix, then E would be an eigenvalue of $\hat{H}$. Solution functions $\psi(x)$ (if there exists any) like this are called *stationary states* or eigenstates. Although the wave function itself is dependent on time:

$$\Psi(x, t) = \psi(x)e^{-iEt/\hbar},$$

the probability density does not:

$$|\Psi(x, t)|^2 = \bar{\Psi}\Psi = \bar{\psi}e^{+iEt/\hbar}\psi e^{-iEt/\hbar} = |\psi(x)|^2.$$

Hence it is stationary. Also the expectation value of any quantity $\langle Q \rangle$ reduces to

$$\langle Q(x, p) \rangle = \int \bar{\psi} \left[Q \left(x, -i\hbar \frac{d}{dx} \right) \right] \psi dx,$$

which is independent of time. Hence every expectation value stays constant. In particular, since $\hat{H}\psi = E\psi$, we have

$$\langle H \rangle = \int \bar{\psi} \hat{H} \psi dx = E \int |\psi|^2 dx = E \int |\Psi|^2 dx = E.$$

$$\langle H^2 \rangle = \int \bar{\psi} \hat{H}^2 \psi dx = E^2 \int |\psi|^2 dx = E^2.$$

$$\sigma_H^2 = \langle H^2 \rangle - \langle H \rangle^2 = E^2 - E^2 = 0.$$

Hence the expectation value of the Hamiltonian is the constant E , and it has zero variance (spread). But zero spread implies every member of the sample has the same value. Thus a separable solution (stationary states) has the property that every measurement of the total energy is certain to return to the value E .

Finally, depending on our choice of the constant E , we have different solutions to the SE:

$$\Psi_1(x, t) = \psi_1(x)e^{-iE_1t/\hbar}, \quad \Psi_2(x, t) = \psi_2(x)e^{-iE_2t/\hbar}, \dots$$

But stationary states form a complete basis in Hilbert space, hence any wave function $\Psi(x, t)$ can be expressed as a linear combination of the stationary states:

$$\Psi(x, t) = \sum_{n=1}^{\infty} c_n \psi_n(x) e^{-iE_n t/\hbar}. \quad (6.7)$$

In particular, setting $t = 0$ gives

$$\Psi(x, 0) = \sum_{n=1}^{\infty} c_n \psi_n(x).$$

We know that each coefficient c_n is the projection of $\Psi(x, 0)$ onto the corresponding stationary state ψ_n . or

$$c_n = \langle \Psi(x, 0), \psi_n \rangle = \int \Psi(x, 0) \psi_n dx.$$

This inner product is a measurement of how close two vectors are in vector space since it is related to the angle between two vectors. It is a measurement of *how possible it is for a state $\Psi(x, 0)$ to collapse into another state ψ_n* . Since we defined the probability to be the square of wave function:

$$|\Psi(x, t)|^2 = \langle \Psi(x, t), \Psi(x, t) \rangle = \bar{\Psi}(x, t) \Psi(x, t),$$

we also expect c_n to be related to the probability of a state $\Psi(x, 0)$ to collapse into another state ψ_n . But since we are doing inner product of two different vectors and the probability always needs to be positive, we define $|c_n|^2$ to be the probability of a state $\Psi(x, 0)$ to collapse into another state ψ_n , or that of a measurement of the energy would return the value E_n . This implies that

$$\sum_{n=1}^{\infty} |c_n|^2 = 1 \quad \langle H \rangle = \sum_{n=1}^{\infty} |c_n|^2 E_n.$$

6.2 Infinite Square Well

Suppose now we have a square well with width $0 < x < a$ (dimension a). The potential outside the well is infinite and the potential inside the well is zero.

$$V(x) = \begin{cases} 0, & 0 \leq x \leq a \\ \infty, & \text{otherwise} \end{cases}.$$

Outside the well, $\psi(x) = 0$ since an infinite force prevent the particle from escaping. Inside the well where $V = 0$, the TISE reads

$$-\frac{\hbar^2}{2m} \frac{d^2 \psi}{dx^2} = E \psi.$$

Now, assuming $E > 0$ (if $E < 0$, we have an exercise showing that it will not work), we can write

$$\frac{d^2 \psi}{dx^2} = -k^2 \psi, \quad \text{where } k \equiv \frac{\sqrt{2mE}}{\hbar}.$$

Then the solution is of the form

$$\psi(x) = A \sin kx + B \cos kx.$$

The continuity of $\psi(x)$ requires that

$$\psi(0) = \psi(a) = 0.$$

Hence

$$\psi(0) = A \sin 0 + B \cos 0 = B$$

implies $B = 0$ and we have $\psi(x) = A \sin kx$. since $\psi(a) = 0$ and $A \neq 0$ (otherwise we have the trivial solution), we have

$$ka = 0, \pm\pi, \pm2\pi, \pm3\pi, \dots$$

Throwing away $k = 0$ (otherwise $\psi(x) = 0$ again) gives

$$k_n = \frac{n\pi}{a}, \quad \text{with } n = 1, 2, 3, \dots$$

Thus by the definition of k we can arrive at the allowed energy

$$E_n = \frac{\hbar^2 k_n^2}{2m} = \frac{n^2 \pi^2 \hbar^2}{2ma^2}. \quad (6.8)$$

We now determine A by normalizing the wave function:

$$\int_0^a |A|^2 \sin^2(kx) dx = |A|^2 \frac{a}{2} = 1.$$

This gives $A = \sqrt{2/a}$. The inside the well, the solutions are

$$\psi_n(x) = \sqrt{\frac{2}{a}} \sin\left(\frac{n\pi}{a}x\right). \quad (6.9)$$

We immediately notice that ψ forms an orthonormal basis:

$$\int \bar{\psi}_m(x) \psi_n(x) dx = \delta_{mn}.$$

Since this orthonormal basis is complete, any function f can be written as a linear combination of them:

$$f(x) = \sum_{n=1}^{\infty} c_n \psi_n(x) = \sqrt{\frac{2}{a}} \sum_{n=1}^{\infty} c_n \sin\left(\frac{n\pi}{a}x\right).$$

In order to determine the coefficients c_n , we use fourier's trick:

$$\int \bar{\psi}_m(x) f(x) dx = \sum_{n=1}^{\infty} c_n \int \bar{\psi}_m(x) \psi_n(x) dx = \sum_{n=1}^{\infty} c_n \delta_{mn} = c_m.$$

This gives

$$c_n = \int \bar{\psi}_n(x) f(x) dx.$$

Thus, the stationary states of the infinite square well are

$$\Psi_n(x, t) = \sqrt{\frac{2}{a}} \sin\left(\frac{n\pi}{a}x\right) e^{-i(n^2\pi^2\hbar/2ma^2)t}. \quad (6.10)$$

We said before that any state can be written as a linear combination of stationary states:

$$\Psi(x, t) = \sum_{n=1}^{\infty} c_n \sqrt{\frac{2}{a}} \sin\left(\frac{n\pi}{a}x\right) e^{-i(n^2\pi^2\hbar/2ma^2)t}. \quad (6.11)$$

Now our task is to determine the coefficients c_n . We utilize the initial condition (which is usually given)

$$\Psi(x, 0) = \sum_{n=1}^{\infty} c_n \psi_n(x).$$

Now by fourier's trick

$$c_n = \sqrt{\frac{2}{a}} \int_0^a \sin\left(\frac{n\pi}{a}x\right) \Psi(x, 0) dx. \quad (6.12)$$

Now it is a good time to check whether our previous claim that $|c_n|^2$ represents probability and their sum is 1 holds:

$$\begin{aligned} 1 &= \int |\Psi(x, 0)|^2 dx = \sum_{m=1}^{\infty} \sum_{n=1}^{\infty} c_m^* c_n \int \bar{\psi}_m(x) \psi_n(x) dx \\ &= \sum_{n=1}^{\infty} \sum_{m=1}^{\infty} c_m^* c_n \delta_{mn} = \sum_{n=1}^{\infty} |c_n|^2. \end{aligned} \quad (6.13)$$

So this is indeed true. Also, since for stationary states we have $\hat{H}\psi_n = E_n\psi_n$, we can calculate the expectation value of the Hamiltonian.

$$\langle H \rangle = \int \bar{\Psi} \hat{H} \Psi dx = \sum_m \sum_n c_m^* c_n E_n \int \bar{\psi}_m \psi_n dx = \sum_n |c_n|^2 E_n. \quad (6.14)$$

Thus our proposition holds in the infinite square well problem, I hope this convince you that $|c_n|^2$ are indeed probabilities.

6.3 The Harmonic Oscillator

In this chapter we study the single most important quantum model: the ***simple harmonic oscillator***. In this senario our quantum harmonic oscillator has the potential

$$V(x) = \frac{1}{2}m\omega^2 x^2.$$

So the TISE reads

$$-\frac{\hbar^2}{2m} \frac{d^2\psi}{dx^2} + \frac{1}{2}m\omega^2 x^2 \psi = E\psi. \quad (6.15)$$

There are two ways of solving (6.15). One via algebraic method and the other is analytic, which we will not get into and refer the curious readers to [6]. Let

us look at the more elegant algebraic method, first proposed by Dirac. We write (6.15) in operator form:

$$\hat{H}\psi = \frac{1}{2m} [\hat{p}^2 + (m\omega x)^2] \psi = E\psi. \quad (6.16)$$

Recall the factorization $u^2 + v^2 = (iu + v)(-iu + v)$. Hence we define

$$\hat{a}_{\pm} \equiv \frac{1}{\sqrt{2\hbar m\omega}} (\mp i\hat{p} + m\omega x). \quad (6.17)$$

Then this gives

$$\begin{aligned} \hat{a}_- \hat{a}_+ &= \frac{1}{2\hbar m\omega} (i\hat{p} + m\omega x)(-i\hat{p} + m\omega x) \\ &= \frac{1}{2\hbar m\omega} [\hat{p}^2 + (m\omega x)^2 - im\omega(x\hat{p} - \hat{p}x)]. \end{aligned} \quad (6.18)$$

Hence

$$\hat{a}_- \hat{a}_+ = \frac{1}{2\hbar m\omega} [\hat{p}^2 + (m\omega x)^2] - \frac{i}{2\hbar} [x, \hat{p}].$$

Simple calculations give

$$[x, \hat{p}] = i\hbar.$$

Hence we conclude that

$$\hat{a}_- \hat{a}_+ = \frac{1}{\hbar\omega} \hat{H} + \frac{1}{2}.$$

Or

$$\hat{H} = \hbar\omega \left(\hat{a}_- \hat{a}_+ - \frac{1}{2} \right). \quad (6.19)$$

Similarly we have

$$\hat{a}_+ \hat{a}_- = \frac{1}{\hbar\omega} \hat{H} - \frac{1}{2}.$$

In particular we have this wonderful relationship

$$[\hat{a}_-, \hat{a}_+] = 1. \quad (6.20)$$

Hence the Hamiltonian can also be written as

$$\hat{H} = \hbar\omega \left(\hat{a}_+ \hat{a}_- + \frac{1}{2} \right).$$

Now the TISE (6.16) becomes

$$\hbar\omega \left(\hat{a}_{\pm} \hat{a}_{\mp} \pm \frac{1}{2} \right) \psi = E\psi. \quad (6.21)$$

The key observation is that if ψ solves the TISE $\hat{H}\psi = E\psi$ with energy E , then $(\hat{a}_+ \psi)$ solves the TISE with energy $E + \hbar\omega$. This is because

$$\begin{aligned} \hat{H}(\hat{a}_+ \psi) &= \hbar\omega \left(\hat{a}_+ \hat{a}_- + \frac{1}{2} \right) (\hat{a}_+ \psi) = \hbar\omega \left(\hat{a}_+ \hat{a}_- \hat{a}_+ + \frac{1}{2} \hat{a}_+ \right) \psi \\ &= \hbar\omega \hat{a}_+ \left(\hat{a}_- \hat{a}_+ + \frac{1}{2} \right) \psi = \hat{a}_+ \left[\hbar\omega \left(\hat{a}_+ \hat{a}_- + 1 + \frac{1}{2} \right) \psi \right] \\ &= \hat{a}_+ (\hat{H} + \hbar\omega) \psi = \hat{a}_+ (E + \hbar\omega) \psi = (E + \hbar\omega) (\hat{a}_+ \psi). \end{aligned} \quad (6.22)$$

This concludes the proof. Similarly $(\hat{a}_-)\psi$ solves the TISE with energy $E - \hbar\omega$. Because

$$\begin{aligned}\hat{H}(\hat{a}_-\psi) &= \hbar\omega\left(\hat{a}_-\hat{a}_+ - \frac{1}{2}\right)(\hat{a}_-\psi) = \hbar\omega\hat{a}_-\left(\hat{a}_+\hat{a}_- - \frac{1}{2}\right)\psi \\ &= \hat{a}_-\left[\hbar\omega\left(\hat{a}_-\hat{a}_+ - 1 - \frac{1}{2}\right)\psi\right] = \hat{a}_-(\hat{H} - \hbar\omega)\psi = \hat{a}_-(E - \hbar\omega)\psi \\ &= (E - \hbar\omega)(\hat{a}_-\psi).\end{aligned}\tag{6.23}$$

We call $\hat{a}_+$ the rising operator, which allows us to climb up the energy by generating a new solution with energy differ by $+\hbar\omega$. Similarly we call $\hat{a}_-$ the lowering operator, which allows us to climb down the energy by generating a new solution with energy differ by $-\hbar\omega$.

But if we apply the lowering operator repeatedly, we will eventually reach a state with energy less than zero, which cannot happen since then the wave function will no longer be normalizable (by Griffiths Problem 2.2, E must exceed the minimum value of $V(x)$, which is zero in the infinite square well case. But this does not mean that E cannot be smaller than zero. It's just that in this case it cannot. In fact, in section 2.6 we will see such situation where $E < 0$, called the bound state). Thus we conclude that there must exist a ground state ψ_0 such that

$$\hat{a}_-\psi_0 = 0.$$

This gives

$$\frac{1}{\sqrt{2\hbar m\omega}}\left(\hbar\frac{d}{dx} + m\omega x\right)\psi_0 = 0.$$

Or,

$$\frac{d\psi_0}{dx} = -\frac{m\omega}{\hbar}x\psi_0.$$

The solution to this equation is

$$\psi_0(x) = Ae^{-\frac{m\omega}{2\hbar}x^2}.$$

We normalize it:

$$1 = |A|^2 \int_{-\infty}^{\infty} e^{-m\omega x^2/\hbar} dx = |A|^2 \sqrt{\frac{\pi\hbar}{m\omega}}.$$

Hence we have our ground state

$$\psi_0(x) = \left(\frac{m\omega}{\pi\hbar}\right)^{1/4} e^{-\frac{m\omega}{2\hbar}x^2}.\tag{6.24}$$

The TISE for this ground state is $\hbar\omega(\hat{a}_+\hat{a}_- + 1/2)\psi_0 = E_0\psi_0$. Knowing that $\hat{a}_-\psi_0 = 0$ we immediately arrive at

$$E_0 = \frac{1}{2}\hbar\omega.\tag{6.25}$$

Now that we have the ground state, we can repeatedly apply the raising operator to achieve all possible states.

$$\psi_n(x) = A_n(\hat{a}_+)^n\psi_0(x), \quad \text{with} \quad E_n = \left(n + \frac{1}{2}\right)\hbar\omega.\tag{6.26}$$

But the difficult things to do now is to determine the normalization constant A_n . Now we turn to this problem. From (6.26) we see that $\hat{a}_\pm \psi_n$ differs from $\psi_{n\pm 1}$ by a constant. Hence we may assume that

$$\hat{a}_+ \psi_n = c_n \psi_{n+1}, \quad \hat{a}_- \psi_n = d_n \psi_{n-1}. \quad (6.27)$$

Now we need the following result:

Lemma 6.1. *For any functions f, g we have*

$$\int_{-\infty}^{\infty} \bar{f}(\hat{a}_\pm g) dx = \int_{-\infty}^{\infty} \overline{(\hat{a}_\mp f)} g dx. \quad (6.28)$$

Proof. This is because

$$\int_{-\infty}^{\infty} \bar{f}(\hat{a}_\pm g) dx = \frac{1}{\sqrt{2\hbar m\omega}} \int_{-\infty}^{\infty} \bar{f} \left(\mp \hbar \frac{d}{dx} + m\omega x \right) g dx.$$

Suppose the integral exists, then $f, g \rightarrow 0$ at $x = \pm\infty$. Integrate by part will throw away the boundary terms, which gives

$$\begin{aligned} \int_{-\infty}^{\infty} \bar{f}(\hat{a}_\pm g) dx &= \frac{1}{\sqrt{2\hbar m\omega}} \int_{-\infty}^{\infty} \overline{(\pm \hbar f' + m\omega x f)} g dx \\ &= \int_{-\infty}^{\infty} \overline{(\hat{a}_\pm f)} g dx. \end{aligned}$$

This proves the claim. $\square$

In particular, now we have

$$\int_{-\infty}^{\infty} \overline{(\hat{a}_\pm \psi_n)} (\hat{a}_\pm \psi_n) dx = \int_{-\infty}^{\infty} \overline{(\hat{a}_\mp \hat{a}_\pm \psi_n)} \psi_n dx.$$

Now, (6.21) says

$$\hbar\omega \left(\hat{a}_\pm \hat{a}_\mp \pm \frac{1}{2} \right) \psi_n = E_n \psi_n$$

and we have (6.26). Hence

$$\hat{a}_+ \hat{a}_- \psi_n = n \psi_n, \quad \hat{a}_- \hat{a}_+ \psi_n = (n+1) \psi_n. \quad (6.29)$$

Hence now we have

$$\begin{aligned} \int_{-\infty}^{\infty} \overline{(\hat{a}_+ \psi_n)} (\hat{a}_+ \psi_n) dx &= |c_n|^2 \int_{-\infty}^{\infty} |\psi_{n+1}|^2 dx = (n+1) \int_{-\infty}^{\infty} |\psi_n|^2 dx, \\ \int_{-\infty}^{\infty} \overline{(\hat{a}_- \psi_n)} (\hat{a}_- \psi_n) dx &= |d_n|^2 \int_{-\infty}^{\infty} |\psi_{n-1}|^2 dx = n \int_{-\infty}^{\infty} |\psi_n|^2 dx. \end{aligned}$$

Since ψ_n and $\psi_{n\pm 1}$ are normalized, we have $|c_n|^2 = n+1$ and $|d_n|^2 = n$. This implies

$$\hat{a}_+ \psi_n = \sqrt{n+1} \psi_{n+1}, \quad \hat{a}_- \psi_n = \sqrt{n} \psi_{n-1}.$$

This means that

$$\psi_1 = \hat{a}_+ \psi_0,$$

$$\begin{aligned}
\psi_2 &= \frac{1}{\sqrt{2}} \hat{a}_+ \psi_1 = \frac{1}{\sqrt{2}} (\hat{a}_+)^2 \psi_0, \\
\psi_3 &= \frac{1}{\sqrt{3}} \hat{a}_+ \psi_2 = \frac{1}{\sqrt{3 \cdot 2}} (\hat{a}_+)^3 \psi_0, \\
\psi_4 &= \frac{1}{\sqrt{4}} \hat{a}_+ \psi_3 = \frac{1}{\sqrt{4 \cdot 3 \cdot 2}} (\hat{a}_+)^4 \psi_0 \dots
\end{aligned}$$

Then clearly

$$\psi_n = \frac{1}{\sqrt{n!}} (\hat{a}_+)^n \psi_0. \quad (6.30)$$

Hence we have figured out the normalization coefficient

$$A_n = \frac{1}{\sqrt{n!}}.$$

It can be checked that the stationary states for harmonic oscillator is also orthonormal:

$$\int_{-\infty}^{\infty} \bar{\psi}_m \psi_n dx = \delta_{mn}.$$

This is because

$$\begin{aligned}
\int_{-\infty}^{\infty} \bar{\psi}_m (\hat{a}_+ \hat{a}_-) \psi_n dx &= n \int_{-\infty}^{\infty} \bar{\psi}_m \psi_n dx \\
&= \int_{-\infty}^{\infty} \overline{(\hat{a}_- \psi_m)} (\hat{a}_- \psi_n) dx = \int_{-\infty}^{\infty} \overline{(\hat{a}_+ \hat{a}_- \psi_m)} \psi_n dx \\
&= m \int_{-\infty}^{\infty} \bar{\psi}_m \psi_n dx.
\end{aligned}$$

Unless $m = n$, the integral must be zero. Thus we can again use Fourier's trick to determine c_n .

6.4 The Uncertainty Principle

In this section we prove what is arguably one of the most important results in quantum mechanics, the ***Heisenberg's uncertainty principle***. The statement is as follows:

Theorem 6.2. (*Generalized Uncertainty Principle*) For two observables $\hat{A}, \hat{B}$, we always have

$$\sigma_A^2 \sigma_B^2 \geq \left(\frac{1}{2i} \langle [\hat{A}, \hat{B}] \rangle \right)^2.$$

Since the commutator $[\hat{A}, \hat{B}]$ carries its own factor of i , the right hand side is positive.

Proof. For any observable $\hat{A}$ we have

$$\sigma_A^2 = \langle (\hat{A} - \langle A \rangle) \Psi | (\hat{A} - \langle A \rangle) \Psi \rangle = \langle f | f \rangle,$$

where $f \equiv (\hat{A} - \langle A \rangle) \Psi$. Likewise, for an observable $\hat{B}$, we have

$$\sigma_B^2 = \langle g | g \rangle, \quad \text{where } g \equiv (\hat{B} - \langle B \rangle) \Psi.$$

Schwarz inequality says that

$$\sigma_A^2 \sigma_B^2 = \langle f|f \rangle \langle g|g \rangle \geq |\langle f|g \rangle|^2.$$

Now, for any complex number z , we have

$$|z|^2 = [\text{Re}(z)]^2 + [\text{Im}(z)]^2 \geq [\text{Im}(z)]^2 = \left(\frac{1}{2i}\right)^2 (z - \bar{z})^2.$$

Therefore, let $z = \langle f|g \rangle$,

$$\sigma_A^2 \sigma_B^2 \geq \left(\frac{1}{2i}\right)^2 (\langle f|g \rangle - \langle g|f \rangle)^2.$$

Now, since $\hat{A} - \langle A \rangle$ is Hermitian,

$$\begin{aligned} \langle f|g \rangle &= \langle (\hat{A} - \langle A \rangle) \Psi | (\hat{B} - \langle B \rangle) \Psi \rangle \\ &= \langle \Psi | (\hat{A} - \langle A \rangle) (\hat{B} - \langle B \rangle) | \Psi \rangle \\ &= \langle \Psi | \hat{A} \hat{B} - \hat{A} \langle B \rangle - \hat{B} \langle A \rangle + \langle A \rangle \langle B \rangle | \Psi \rangle \\ &= \langle \hat{A} \hat{B} \rangle - \langle B \rangle \langle A \rangle - \langle A \rangle \langle B \rangle + \langle A \rangle \langle B \rangle \\ &= \langle \hat{A} \hat{B} \rangle - \langle A \rangle \langle B \rangle. \end{aligned}$$

Note that $\langle A \rangle, \langle B \rangle$ are numbers. So they commute. Similarly we have

$$\langle g|f \rangle = \langle \hat{B} \hat{A} \rangle - \langle A \rangle \langle B \rangle.$$

Hence

$$\langle f|g \rangle - \langle g|f \rangle = \langle \hat{A} \hat{B} - \hat{B} \hat{A} \rangle = \langle [\hat{A}, \hat{B}] \rangle.$$

Therefore,

$$\sigma_A^2 \sigma_B^2 \geq \left(\frac{1}{2i} \langle [\hat{A}, \hat{B}] \rangle\right)^2.$$

This concludes the proof. $\square$

As an easy corollary, let $\hat{A} = \hat{X}$ and $\hat{B} = \hat{P}$. And since $[\hat{X}, \hat{P}] = i\hbar$, we have

$$\sigma_X^2 \sigma_P^2 \geq \left(\frac{1}{2i} i\hbar\right)^2 = \left(\frac{\hbar}{2}\right)^2.$$

Then the uncertainty principle immediately follows. This can also be explained in experimental physics. Suppose you first measure the position of the particle. Then once you measure the position, the wave function collapses into a spike. But then since $p = h/\lambda = 2\pi\hbar/\lambda$, with wavefunction collapsing into a spike, there is no clear notion of the wavelength λ .

There is, in fact, an uncertainty principle for every pair of observables whose operators do not commute. We call them incompatible observables.

6.5 Quantum Mechanics in Three Dimensions

Before going into 3D, let us consider a simple case, where the particle lives in 2D. Then the generalization of Hamiltonian is simple:

$$\hat{H} = \frac{\hat{p}_x^2}{2m} + \frac{\hat{p}_y^2}{2m} + V(x, y) = -\frac{\hbar}{2m} \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} \right) + V(x, y)$$

Further, let us assume the potential is separable. This means that

$$V(x, y) = V_1(x) + V_2(y).$$

This is rarely the case in actual situations. But let's just consider this for simplicity. Then the Hamiltonian becomes

$$\hat{H} = \left(\frac{\hat{p}_x^2}{2m} + V_1(x) \right) + \left(\frac{\hat{p}_y^2}{2m} + V_2(y) \right) = \hat{H}_x + \hat{H}_y.$$

Now, suppose that $\psi_x(x)$ solves the 1D T.I.S.E. $\hat{H}_x\psi(x) = E_x\psi(x)$ and $\psi_y(y)$ solves the 1D time independent Schrödinger equation (T.I.S.E.) $\hat{H}_y\psi(y) = E_y\psi(y)$, then I claim the solution to this 2D problem $\hat{H}\psi(x, y) = E\psi(x, y)$ is just

$$\psi(x, y) = \psi_x(x)\psi_y(y) \text{ with } E = E_x + E_y.$$

The proof is a simple computation, which is left to the reader as exercise. Thus in this case, probability multiplies, and energy adds. You will also notice that in this case, we can have degeneracy, which does not appear in 1D situation.

The schrodinger equation $i\hbar\partial_t\Psi = \hat{H}\Psi$ still holds. But now the 3D momentum becomes $\mathbf{p} = -i\hbar\nabla$. And also potential $V(x, y, z)$ is a function of three coordinates. Now, the Hamiltonian becomes

$$\hat{H} = \frac{1}{2m} (\hat{p}_x^2 + \hat{p}_y^2 + \hat{p}_z^2) + V = -\frac{\hbar^2}{2m} \nabla^2 + V,$$

where

$$\nabla^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} = \nabla \cdot \nabla$$

And the Shrodinger equation becomes

$$i\hbar\frac{\partial\Psi}{\partial t} = -\frac{\hbar^2}{2m}\nabla^2\Psi + V\Psi.$$

And now the normalization condition becomes

$$\int d^3\mathbf{r} |\Psi|^2 = \int_{\Omega} dx dy dz |\Psi(x, y, z)|^2 = 1.$$

Here, Ω is the region where the particle is allowed, or where the wave function is defined. We write the differential $dx dy dz = d^3\mathbf{r}$ and the triple integral as just one integral for convinience. The probability of finding the particle in a small region with volume $dx dy dz = d^3\mathbf{r}$ centered at $\mathbf{r} = (x, y, z)$ now becomes $|\Psi|^2 d^3\mathbf{r}$. Still, the stationary state that evolve with time will be of the form $\Psi_n(\mathbf{r}, t) = \psi_n(\mathbf{r})e^{-iE_n t/\hbar}$, where ψ_n solves the T.I.S.E

$$-\frac{\hbar^2}{2m}\nabla^2\psi + V\psi = E\psi.$$

And then we can express any general state by a sum:

$$\Psi(\mathbf{r}, t) = \sum_{n=0}^{\infty} c_n \psi_n(\mathbf{r}) e^{-iE_n t/\hbar}.$$

The generalization seems straightforward for now. Also notice that now if we label $\mathbf{r} = (r_1, r_2, r_3) = (x, y, z)$ and $\mathbf{p} = (p_1, p_2, p_3) = (p_x, p_y, p_z)$, we have the commutator relations

$$[r_i, p_j] = -[p_i, r_j] = i\hbar\delta_{ij}, \quad [r_i, r_j] = [p_i, p_j] = 0.$$

Also, the familiar relation between momentum and position and the uncertainty principle still hold:

$$m \frac{d\langle \mathbf{r} \rangle}{dt} = \langle \mathbf{p} \rangle, \quad \sigma_{r_i} \sigma_{p_i} \geq \frac{\hbar}{2}.$$

Sometimes it is easier to do things in spherical coordinates. In spherical coordinates, the Laplacian becomes

$$\nabla^2 = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \left(\frac{\partial^2}{\partial \phi^2} \right).$$

The TISE in spherical coordinate reads

$$-\frac{\hbar^2}{2m} \left[\frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial \psi}{\partial r} \right) + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial \psi}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \left(\frac{\partial^2 \psi}{\partial \phi^2} \right) \right] + V\psi = E\psi.$$

We try to separate the wave function into products

$$\psi(r, \theta, \phi) = R(r)Y(\theta, \phi).$$

We substitute this back into the TISE and get

$$-\frac{\hbar^2}{2m} \left[\frac{Y}{r^2} \frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) + \frac{R}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial Y}{\partial \theta} \right) + \frac{R}{r^2 \sin^2 \theta} \frac{\partial^2 Y}{\partial \phi^2} \right] + VRY = ERY.$$

Dividing by $\psi = RY$ and multiply by $-2mr^2/\hbar^2$ gives

$$\left\{ \frac{1}{R} \frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) - \frac{2mr^2}{\hbar^2} [V(r) - E] \right\} + \frac{1}{Y} \left\{ \frac{1}{\sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial Y}{\partial \theta} \right) + \frac{1}{\sin^2 \theta} \frac{\partial^2 Y}{\partial \phi^2} \right\} = 0.$$

The first term is a function of only r and the second term is a function of only ϕ and θ . This suggests us to use a separation constant $\ell(\ell + 1)$:

$$\frac{1}{R} \frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) - \frac{2mr^2}{\hbar^2} [V(r) - E] = \ell(\ell + 1),$$

$$\frac{1}{Y} \left\{ \frac{1}{\sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial Y}{\partial \theta} \right) + \frac{1}{\sin^2 \theta} \frac{\partial^2 Y}{\partial \phi^2} \right\} = -\ell(\ell + 1).$$

We will justify the choice of the separation constant in due course. The first equation is called the radial equation and the second one is called the angular equation.

However, one deceptively simple question to ask is: Is this really legal? How are we allowed to divide ψ on both sides for ψ can be zero at some point? Then aren't we dividing by zero? This question should have been asked way before when we derived TISE in 1D. To justify this we need the following lemma:

Lemma 6.3. *Let I, J be intervals on x, y axis respectively. Assume that*

$$a(x)b(y) = c(x)d(y) \quad \forall x \in I, \forall y \in J \quad (6.31)$$

with both sides not equal to 0 for at least one point (x_0, y_0) . Then there exists a constant $\mu \neq 0$ with

$$\begin{aligned} c(x) &= \mu a(x) & x \in I, \\ b(y) &= \mu d(y) & y \in J. \end{aligned}$$

This means that we can sort of "cheat" and pretend $\frac{a(x)}{c(x)} = \frac{b(y)}{d(y)} = \mu$.

Proof. Put $x = x_0$ in (6.31) to obtain

$$b(y) = \frac{c(x_0)}{a(x_0)} d(y) \quad y \in J.$$

And therefore $b(y) = \mu d(y)$ with $\mu = c(x_0)/a(x_0) \neq 0$. Now put $y = y_0$ in (6.31) and get

$$a(x)\mu d(y_0) = c(x)d(y_0) \quad x \in I.$$

This implies $c(x) = \mu a(x)$, as desired. $\square$

One can see without too much work that this lemma justifies the dividing by ψ step. This is cheating, of course, but it actually is very intuitive, and even better: it is correct!

6.5.1 The Angular Equation

We now try to solve the angular equation:

$$\sin \theta \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial Y}{\partial \theta} \right) + \frac{\partial^2 Y}{\partial \phi^2} = -\ell(\ell + 1) \sin^2 \theta Y.$$

We use separation of variables again:

$$Y(\theta, \phi) = \Theta(\theta)\Phi(\phi).$$

We plug this in and divide by $Y = \Theta\Phi$ again (justified!) and get

$$\left\{ \frac{1}{\Theta} \left[\sin \theta \frac{d}{d\theta} \left(\sin \theta \frac{d\Theta}{d\theta} \right) \right] + \ell(\ell + 1) \sin^2 \theta \right\} + \frac{1}{\Phi} \frac{d^2 \Phi}{d\phi^2} = 0.$$

This time let the separation constant be m^2 :

$$\frac{1}{\Theta} \left[\sin \theta \frac{d}{d\theta} \left(\sin \theta \frac{d\Theta}{d\theta} \right) \right] + \ell(\ell + 1) \sin^2 \theta = m^2,$$

$$\frac{1}{\Phi} \frac{d^2 \Phi}{d\phi^2} = -m^2.$$

Instantly we can find the solution to the second equation be $\Phi(\phi) = e^{im\phi}$. This time we write it as a complex exponential instead of a combination of trig functions for convinience. As ϕ advances by 2π we get to the same point.

Thus the boundary condition requires $\Phi(\phi + 2\pi) = \Phi(\phi)$. This suggests that $m = 0, \pm 1, \pm 2, \dots$. The other equation is harder to solve:

$$\sin \theta \frac{d}{d\theta} \left(\sin \theta \frac{d\Theta}{d\theta} \right) + [\ell(\ell + 1) \sin^2 \theta - m^2] \Theta = 0.$$

The solution is

$$\Theta(\theta) = AP_\ell^m(\cos \theta),$$

where P_ℓ^m is the associated Legendre function defined by

$$P_\ell^m(x) \equiv (-1)^m (1 - x^2)^{m/2} \left(\frac{d}{dx} \right)^m P_\ell(x)$$

and $P_\ell(x)$ is the ℓ -th Legendre polynomial, defined by

$$P_\ell(x) \equiv \frac{1}{2^\ell \ell!} \left(\frac{d}{dx} \right)^\ell (x^2 - 1)^\ell.$$

Notice that $\ell \geq 0$ so this definition makes sense. If $m > \ell$, then the definition says that $P_\ell^m = 0$. Hence for any given ℓ , we have $2\ell + 1$ choices of m :

$$m = -\ell, -\ell - 1, \dots, -1, 0, 1, \dots, \ell - 1, \ell.$$

The volume element in spherical coordinate is

$$d^3\mathbf{r} = r^2 \sin \theta dr d\theta d\phi = r^2 dr d\Omega, \quad \text{where} \quad d\Omega \equiv \sin \theta d\theta d\phi.$$

The normalization condition becomes

$$\int |\psi|^2 r^2 \sin \theta dr d\theta d\phi = \int |R|^2 r^2 dr \int |Y|^2 d\Omega = 1.$$

We normalize R and Y separately:

$$\int_0^\infty |R|^2 r^2 dr = 1 \quad \text{and} \quad \int_0^{2\pi} \int_0^\pi |Y|^2 \sin \theta d\theta d\phi = 1$$

Notice that Griffiths actually made a mistake in his book. θ goes from 0 to π and ϕ goes from 0 to 2π , not the other way around!

The normalized angular wave function is

$$Y_\ell^m(\theta, \phi) = \sqrt{\frac{(2\ell + 1)(\ell - m)!}{4\pi(\ell + m)!}} e^{im\phi} P_\ell^m(\cos \theta).$$

And they are mutually orthogonal:

$$\int_0^{2\pi} \int_0^\pi \bar{Y}_\ell^m(\theta, \phi) Y_{\ell'}^{m'}(\theta, \phi) \sin \theta d\theta d\phi = \delta_{\ell\ell'} \delta_{mm'}.$$

6.5.2 The Radial Equation

If we use variable $u(r) \equiv rR(r)$, the radial equation simplifies to

$$-\frac{\hbar^2}{2m} \frac{d^2 u}{dr^2} + \left[V + \frac{\hbar^2}{2m} \frac{\ell(\ell+1)}{r^2} \right] u = Eu.$$

We can write

$$-\frac{\hbar^2}{2m} \frac{d^2 u}{dr^2} + V_{\text{eff}} u = Eu,$$

where the effective potential is defined to be

$$V_{\text{eff}} = V + \frac{\hbar^2}{2m} \frac{\ell(\ell+1)}{r^2}.$$

Now the normalization condition reads

$$\int_0^\infty |u|^2 dr = 1.$$

6.6 Hydrogen Atom

Our task is to solve the TISE:

$$-\frac{\hbar^2}{2m} \nabla^2 \psi(\mathbf{r}) + V(\mathbf{r}) \psi(\mathbf{r}) = E \psi(\mathbf{r}).$$

In the problem we are interested in, we only consider radial potential:

$$V(\mathbf{r}) = V(r) \text{ where } r = \|\mathbf{r}\|.$$

Thus it makes sense to consider solving TISE in spherical coordinate. Here we denote the coordinate system as (r, ϕ, θ) , where

$$0 \leq \phi \leq 2\pi \quad 0 \leq \theta \leq \pi.$$

Then the Laplacian $\nabla^2 \psi$ becomes

$$\nabla^2 \psi = \underbrace{\left(\frac{\partial^2 \psi}{\partial r^2} + \frac{2}{r} \frac{\partial \psi}{\partial r} \right)}_{\frac{1}{r^2} \frac{\partial}{\partial r} (r^2 \frac{\partial \psi}{\partial r})} + \frac{1}{r^2 \sin \theta} \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial \psi}{\partial \theta} \right) + \frac{1}{r^2 \sin^2 \theta} \frac{\partial^2 \psi}{\partial \phi^2}.$$

We may consider a separation of variable: $\psi(r, \phi, \theta) = R(r)Y(\phi, \theta)$. Further, we let the separation constant be $\ell(\ell+1)$. Then the DE involving $Y(\phi, \theta)$ becomes the well-known angular equation:

$$\sin \theta \frac{\partial}{\partial \theta} \left(\sin \theta \frac{\partial Y}{\partial \theta} \right) + \frac{\partial^2 Y}{\partial \phi^2} = -\ell(\ell+1) \sin^2 \theta Y.$$

The solution to the angular equation is

$$Y_\ell^m(\theta, \phi) = \sqrt{\frac{(2\ell+1)}{4\pi} \frac{(\ell-m)!}{(\ell+m)!}} e^{im\phi} P_\ell^m(\cos \theta).$$

They are called the spherical harmonics. What we are interested now is to solve the radial part:

$$\frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) - \frac{2mr^2}{\hbar^2} [V(r) - E] R = \ell(\ell + 1) R$$

with a given radial potential $V(r)$.

For the Hydrogen atom, we plug in the Columb potential:

$$V(r) = -\frac{e^2}{4\pi\epsilon_0 r}.$$

Simplifying the radial equation, we get

$$-\frac{\hbar^2}{2mr^2} \frac{d}{dr} \left(r^2 \frac{dR}{dr} \right) + \left[-\frac{e^2}{4\pi\epsilon_0} \frac{1}{r} + \frac{\hbar^2}{2m} \frac{\ell(\ell + 1)}{r^2} \right] R = ER.$$

And this is the equation that we are aiming to solve. There are two methods to solve this equation: one algebraic (involving defining rising and lowering operators), and one analytic (involving asymptotic analysis and infinite series). We will only present the algebraic method here since it is more interesting, and refer interested readers to check [6] for the analytical method.

We wish to simplify the radial equation to the form $\hat{H}_r R = ER$. Thus we need to define a radial Hamiltonian

$$\hat{H}_r = \frac{p_r^2}{2m} + V_{\text{eff}}.$$

p_r is the radial momentum and V_{eff} is the effective potential:

$$V_{\text{eff}} = -\frac{e^2}{4\pi\epsilon_0} \frac{1}{r} + \frac{\hbar^2}{2m} \frac{\ell(\ell + 1)}{r^2}.$$

Our first attempt to define p_r is by

$$\hat{\mathbf{r}} \cdot \mathbf{p} = \hat{\mathbf{r}} \cdot -i\hbar \nabla = \hat{\mathbf{r}} \cdot -i\hbar \left(\hat{r} \frac{\partial}{\partial r} + \hat{\theta} \frac{1}{r} \frac{\partial}{\partial \theta} + \hat{\phi} \frac{1}{r \sin \theta} \frac{\partial}{\partial \phi} \right) = -i\hbar \frac{\partial}{\partial r}.$$

But unfortunately, this expression is not Hermitian. Why? Because in spherical coordinates

$$\left(\frac{\partial}{\partial r} \right)^\dagger \neq -\frac{\partial}{\partial r}.$$

However, we can make it Hermitian by adding a symmetric component:

$$p_r \equiv \frac{1}{2} (\hat{\mathbf{r}} \cdot \mathbf{p} + \mathbf{p} \cdot \hat{\mathbf{r}}) = -\frac{i\hbar}{2} \left(\frac{1}{r} \mathbf{r} \cdot \nabla + \nabla \cdot \left(\frac{\mathbf{r}}{r} \right) \right).$$

The first term is just $1/r(\mathbf{r} \cdot \nabla) = r\partial_r$ by the definition of ∇ in spherical coordinate (last page). We need to be careful with the second term:

$$\nabla \cdot \left(\frac{\mathbf{r}}{r} \right) f = f(\mathbf{r}) \nabla \cdot \frac{\mathbf{r}}{r} + \left(\frac{\mathbf{r}}{r} \cdot \nabla \right) f(\mathbf{r}) = \left(\frac{2}{r} + \frac{\partial}{\partial r} \right) f.$$

Plugging in and we find that

$$p_r = -i\hbar \left(\frac{\partial}{\partial r} + \frac{1}{r} \right).$$

We now prove that

Proposition 6.4. p_r is Hermitian.

Proof. This is just a simple calculation:

$$\begin{aligned}
\langle \psi | p_r \psi \rangle &= \int_0^\infty \int_0^{2\pi} \int_0^\pi \bar{\psi} \left[-i\hbar \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] \psi r^2 \sin \theta d\theta d\phi dr \\
&= \left(\int_0^\infty R^* \left[-i\hbar \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] R r^2 dr \right) \underbrace{\left(\int_0^{2\pi} \int_0^\pi |Y|^2 d\Omega \right)}_{=1 \text{ by normalization}} \\
&= i\hbar \int_0^\infty \left(\frac{2}{r} + \frac{\partial}{\partial r} - \frac{1}{r} \right) R^* R r^2 dr \\
&= \int_0^\infty \int_0^{2\pi} \int_0^\pi \left(\left[-i\hbar \left(\frac{\partial}{\partial r} + \frac{1}{r} \right) \right] \psi \right)^* \psi r^2 \sin \theta d\theta d\phi dr \\
&= \langle p_r \psi | \psi \rangle.
\end{aligned}$$

In the third step we expand two terms, integrate by part on the first term, and merge them together. This concludes the proof. $\square$

Here are two useful commutation relations that are easy to check:

$$[r, p_r] = i\hbar \quad [p_r, 1/r] = \frac{i\hbar}{r^2}.$$

Next we verify

$$p_r^2 = -\frac{\hbar^2}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial}{\partial r} \right).$$

Thus

$$\hat{H}_r R = \left[\frac{p_r^2}{2m} - \frac{e^2}{4\pi\epsilon_0} \frac{1}{r} + \frac{\hbar^2}{2m} \frac{\ell(\ell+1)}{r^2} \right] R = ER.$$

Then for each ℓ we can define the corresponding radial Hamiltonian:

$$H_\ell = \frac{p_r^2}{2m} - \frac{\hbar^2}{ma_0 r} + \frac{\hbar^2}{2m} \frac{\ell(\ell+1)}{r^2}.$$

We introduce a constant

$$a_0 \equiv \frac{4\pi\epsilon_0\hbar^2}{me^2} = 0.529 \times 10^{-10} \text{ m}.$$

Define the raising operator

$$a_\ell \equiv \frac{a_0}{\hbar\sqrt{2}} \left(ip_r - \frac{(\ell+1)\hbar}{r} + \frac{\hbar}{(\ell+1)a_0} \right).$$

The lowering operator is just $a^\dagger$. You can check that

$$a_\ell \equiv \frac{a_0}{\hbar\sqrt{2}} \left(-ip_r - \frac{(\ell+1)\hbar}{r} + \frac{\hbar}{(\ell+1)a_0} \right).$$

Then the Hamiltonian can be written as

$$H_\ell = \frac{\hbar^2}{ma_0^2} \left(a_\ell^\dagger a_\ell - \frac{1}{2(\ell+1)^2} \right).$$

The algebra is messy. So we omit the verification. One very important commutation relation is

$$[a_\ell, a_\ell^\dagger] = \frac{(\ell+1)a_0^2}{r^2} = \frac{ma_0^2}{\hbar^2}(H_{\ell+1} - H_\ell).$$

With the result $[a_\ell, a_\ell^\dagger]$, we can now actually calculate

$$[a_\ell, H_\ell] = (H_{\ell+1} - H_\ell) a_\ell.$$

This can also be written as $a_\ell H_\ell = H_{\ell+1} a_\ell$. Suppose now R_ℓ is a solution to $H_\ell R_\ell = E R_\ell$. Then

$$a_\ell H_\ell R_\ell = H_{\ell+1} a_\ell R_\ell = E a_\ell R_\ell.$$

which means that the state $a_\ell R_\ell$ is an eigenstate of $H_{\ell+1}$, lets call it $R_{\ell+1}$, with the same eigenenergy E .

Can we do this infinitely many times, getting eigenstates $R_{\ell+2}, R_{\ell+3}, \dots$? Lets think of the consequences. As ℓ increases, the positive part of the effective potential is increasing, with respect to a fixed negative potential. You can take a derivative of V_{eff} with respect to r to find that it has a minimum at

$$V_{\text{min}} = -\frac{\hbar}{2ma_0^2} \frac{1}{\ell(\ell+1)}.$$

What this means is that with a fixed E (remember were not raising the energy by operating with a_ℓ), as ℓ increases, theres a point where $V_{\text{min}} > E$ and no more bound states exist.

Thus, there must be a maximum value of ℓ , say L , such that $a_L R_L = 0$. The energy of this solution is

$$H_L R_L = \frac{\hbar^2}{ma_0^2} a_L^\dagger \underbrace{(a_L R_L)}_{=0} - \frac{\hbar^2}{2ma_0^2(L+1)^2} R_L = -\frac{\hbar^2}{2ma_0^2(L+1)^2} R_L.$$

So now define $n \equiv L+1$ we see the highest energy is given by

$$E_n = -\frac{\hbar^2}{2ma_0^2 n^2} = -\frac{e^2}{8\pi\epsilon_0 a_0 n^2} = -\frac{E_1}{n^2} \approx -\frac{13.6}{n^2} \text{eV}$$

where E_1 is the Bohr energy.

Now we can generate all the other states with the same energy by acting with the lowering operator, $a_\ell^\dagger$, down to $\ell = 0$. Since we know that $L = n-1$, then for a given n ,

$$0 \leq \ell \leq n-1 \quad -\ell \leq m \leq \ell.$$

Therefore, the degeneracy of a given E_n , denoted by $d(n)$, is given by

$$\sum_{\ell=0}^{n-1} \sum_{m=-\ell}^{\ell} 1 = \sum_{\ell=0}^{n-1} (2\ell+1) = n^2.$$

By using three quantum numbers: n, ℓ, m , we can completely specify the system:

$$\psi_{n,\ell,m}(r, \theta, \phi) = R_{n,\ell}(r) Y_\ell^m(\theta, \phi).$$

The important thing to note here is that by choosing n , we automatically determine a maximal energy E_n , and hence L . Why did I change the notation from $R_\ell(r)$ to $R_{n,\ell}(r)$ here? It is because that previously we assumed that R_ℓ satisfy the equation $\hat{H}_\ell R_\ell = E R_\ell$ for a given E . But this E is completely determined by n . So what I really should have said is that $R_{n,\ell}$ is the solution to the equation

$$\hat{H}_\ell R_{n,\ell} = E_n R_{n,\ell}.$$

So how should I determine the ground state? Well, the ground state corresponds to the energy E_1 , so $n = 1$. By choosing $n = 1$ and by the restriction $0 \leq \ell \leq n - 1$, we force $\ell = 0$. Then by $-\ell \leq m \leq \ell$ we force $m = 0$. So the ground state has no degeneracy. And it is given by

$$\psi_{1,0,0}(r, \theta, \phi) = R_{1,0}(r) Y_0^0(\theta, \phi).$$

But what is this ground state anyway? We know that for a given n , there exists $L = n - 1$ such that $a_L R_{n,L} = 0$. That is,

$$a_{n-1} R_{n,n-1} = 0.$$

We use the definition of the raising operator a and get

$$\frac{a_0}{\hbar\sqrt{2}} \left(\hbar \frac{d}{dr} + \hbar \frac{1}{r} - \frac{n\hbar}{r} + \frac{\hbar}{na_0} \right) R_{n,n-1} = 0.$$

The solution is given by

$$R_{n,n-1} = C r^{n-1} e^{-r/na_0}.$$

Substitute in $n = 1$ we get

$$R_{1,0} = C e^{-r/a_0}.$$

Now we normalize $R_{1,0}$ to find C .

$$1 = \int_0^\infty |R_{1,0}|^2 r^2 dr = \int_0^\infty |C|^2 r^2 e^{-2r/a_0} dr.$$

This suggests that $C = 2/a_0^{3/2}$. By the table of spherical harmonics we know $Y_0^0 = 1/\sqrt{4\pi}$. Thus the ground state of the Hydrogen atom is given by

$$\psi_{1,0,0}(r, \theta, \phi) = \frac{1}{\sqrt{\pi a_0^3}} e^{-r/a_0}.$$

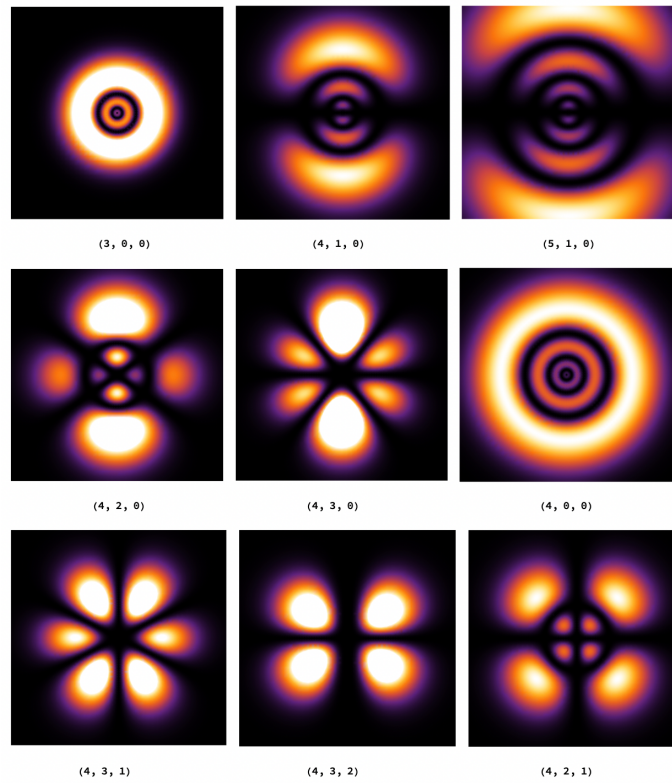


Figure 6.1: Density plot: We can do the density plot of the Hydrogen atom. The brightness of the cloud is proportional to $|\psi|^2$. We use the following representation of the wave function: $(n, \ell, m) = R_{n,\ell}(r)Y_\ell^m(\theta, \phi)$. For example, $(4, 2, 0) = R_{4,2}(r)Y_2^0(\theta, \phi)$. The brighter the color in a particular region, the more probable it is to find the electron in that particular region.

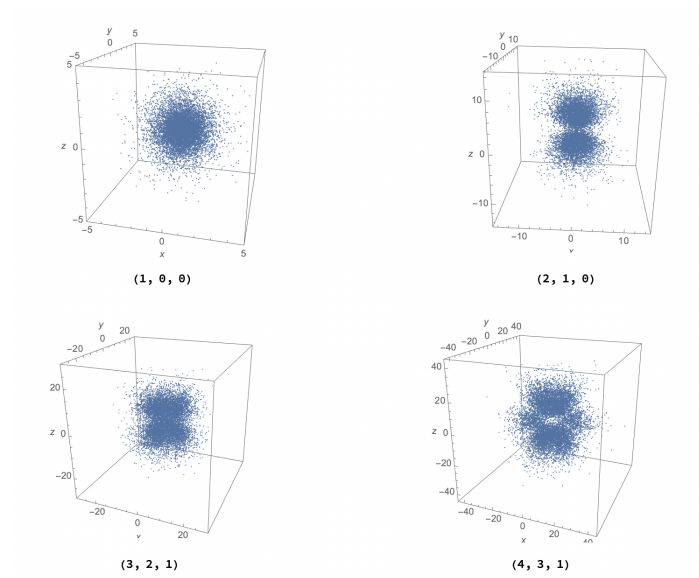


Figure 6.2: 3D plot: We can also do a 3D visualization. each dot of the cloud represents a possible result of a measurement of the position of the electron in an individual atom. By imagining that the measurement is repeated many times (here we do 10,000 times) in different atoms at the same quantum state, you can get a plot representing the probability density function associated with that state.

Chapter 7

Operator Theory in Quantum Mechanics

Now we come to the analysis side of quantum mechanics. Since observables are self-adjoint operators in quantum mechanics, it is no surprise that the theory of operators is extensively used in quantum mechanics if one wants to do rigorous math related to quantum theory. However, most operators in quantum mechanics are unbounded, like the position and momentum operator. To fully study the theory of unbounded operators will take a whole book. Therefore, in this chapter we only include a rapid introduction to theory of bounded operators to give interested reader a taste of the subject, and refer curious readers who wants to specialize in this direction any functional analysis textbook like [7]. Although this chapter contains only less than 30 pages of material, it is extremely technical. For readers who are not that interested in formal mathematics, it is safe to skip this chapter.

7.1 Basic Theory of Bounded Operators

We assume basic knowledge of Banach space and bounded operators. In what follows, we will use the result that if $T : X \rightarrow Y$ is an operator between Banach spaces, then it is continuous iff it is bounded (Note that in finite dimensional case, a linear operator is always continuous, however this is [not](#) always true for infinite dimensional case). A proof of this is given in Prop 7.7. Another result we will assume is the Baire category theorem. We refer the readers to any functional analysis textbooks for details, such as [5].

We use $\text{Cl}[U]$ to denote the closure of a set U , and we use $B(x, r)$ to denote an open ball centered at x of radius r . Finally, we use $B(X)$ to denote the set of bounded operators $T : X \rightarrow X$.

The first result we want to discuss is the open mapping theorem, the proof of which, however, requires the following lemma:

Lemma 7.1. *Let $T : X \rightarrow Y$ be a continuous bounded operator between Banach spaces, and is onto. Then there exists $\delta > 0$ such that $B(0, \delta) \subset T(B(0, 1))$. That is, image of the unit ball under T contains an open ball around the origin.*

Proof. Since T is surjective, we can write $Y = \bigcup_n \text{Cl}[T(B(0, n))]$. Then by the Baire category theorem, there exists some n_0 such that $\text{Cl}[T(B(0, n_0))]$ contains an open ball. But since $\text{Cl}[B(0, n_0)] = n_0 \text{Cl}[T(B(0, 1))]$, this says

that we have an open ball inside $\text{Cl}[T(B(0,1))]$. We choose $r > 0, v_0 \in Y$, such that $B(v_0, 4r) \subset \text{Cl}[T(B(0,1))]$.

But we still want to remove the closure condition and prove the claim in the lemma. We will divide this into several steps:

- (1) Choose $v_1 = Tu_1 \in T(B(0,1))$ such that $\|v_0 - v_1\| < 2r$. This can be done since points in $\text{Cl}[T(B(0,1))]$ is arbitrarily close to points in $T(B(0,1))$. Thus now we have

$$B(v_1, 2r) \subset B(v_0, 4r) \subset \text{Cl}[T(B(0,1))].$$

- (2) We now show that $\text{Cl}[T(B(0,1))]$ contains an open ball centered at 0, getting one step closer to what we want.

Note for any $\|v\| < r$, we have

$$\frac{1}{2}(2v + v_1) \in \frac{1}{2}B(v_1, 2r) \subset \frac{1}{2}\text{Cl}[T(B(0,1))] = \text{Cl}[T(B(0,1/2))].$$

Therefore (recall our choice of u_1 in step 1),

$$v = -T(u_1/2) + \frac{1}{2}(2v + v_1)$$

is an element in $-T(u_1/2) + \text{Cl}[T(B(0,1/2))]$. By linearity, this means that

$$v \in \text{Cl}[T(-u_1/2 + B(0,1/2))].$$

Since by our choice $\|u_1\| < 1$, we must have $-u_1/2 + B(0,1/2) \subset B(0,1)$. Hence by above $v \in \text{Cl}[T(B(0,1))]$. Therefore, we have shown that $B(0,r) \subset \text{Cl}[T(B(0,1))]$.

- (3) Finally, we prove $B(0,r/2) \subset T(B(0,1))$, which we conclude the proof of the lemma.

Note that $B(0, 2^{-n}r) = 2^{-n}B(0, r) \subset 2^{-n}\text{Cl}[T(B(0,1))] = \text{Cl}[T(B(0, 2^{-n}))]$ by step 2.

Now take $\|v\| < r/2$. Put $n = 1$ in above we get $v \in \text{Cl}[T(B(0,1/2))]$. Hence there exists $b_1 \in B(0,1/2)$ such that $\|v - Tb_1\| < r/4$, again, because points in the closure are arbitrarily close to points in the actual set. Now take $n = 2$ in the last paragraph, since $v - Tb_1 \in \text{Cl}[T(B(0,1/4))]$, there exists $b_2 \in B(0,1/4)$ such that $\|v - Tb_1 - Tb_2\| < r/8$. Now iterate this procedure with $n = 3, 4, \dots$ and we get a sequence $(b_k)_{k=1}^\infty \subset X$ such that $\|b_k\| < 2^{-k}$ and

$$\|v - \sum_{k=1}^n Tb_k\| < 2^{-n-1}r.$$

Since the series $\sum_{k=1}^\infty b_k$ is absolutely convergent, and X is a Banach space, so the series is convergent. Put $b = \sum_{k=1}^\infty b_k$. Now since

$$\|b\| = \lim_{n \rightarrow \infty} \left\| \sum_{k=1}^n b_k \right\| \leq \lim_{n \rightarrow \infty} \sum_{k=1}^n \|b_k\| = \sum_{k=1}^\infty \|b_k\| < \sum_{k=1}^\infty 2^{-k} = 1.$$

Hence we have $b \in B(0, 1)$. Furthermore, since T is a bounded continuous operator, we have

$$Tb = \lim_{n \rightarrow \infty} T\left(\sum_{k=1}^n b_k\right) = \lim_{n \rightarrow \infty} \sum_{k=1}^n Tb_k = v$$

because by our construction $\|v - Tb_1 - Tb_2 - \cdots - Tb_k\| \rightarrow 0$ as $k \rightarrow \infty$. Therefore, we conclude that $v \in T(B(0, 1))$. This shows that $B(0, r/2) \subset T(B(0, 1))$, as desired.

Choosing $\delta = r/2$ concludes the proof of the lemma. $\square$

Now we use the lemma above to prove the important result we promised:

Theorem 7.2 (Open Mapping Theorem). *If $T : X \rightarrow Y$ is a continuous bounded operator between Banach spaces, and is onto, then T is an open mapping.*

Proof. Now we use the lemma above to conclude the proof of the open mapping theorem. Let $U \subset X$ be open, and $y = Tx$ is an arbitrary point in $T(U)$. Since U is open, there exists $\epsilon > 0$ such that $B(x, \epsilon) \subset U$. By the lemma above, there exists some $\delta > 0$ such that $B(0, \delta) \subset T(B(0, 1))$. So this gives

$$B(y, \epsilon\delta) = y + \epsilon B(0, \delta) \subset y + \epsilon T(B(0, 1)) = Tx + \epsilon T(B(0, 1)) = T(x + \epsilon B(0, 1))$$

which is equal to $T(B(x, \epsilon)) \subset T(U)$. This concludes the proof. $\square$

Using the open mapping theorem, we can deduce several interesting results:

Theorem 7.3 (Inverse Mapping Theorem). *Let $T : X \rightarrow Y$ a continuous bijective operator between Banach spaces, then T^{-1} is continuous.*

Proof. This follows trivially from the result that T is open, therefore for any $M \subset Y$ open, we have $(T^{-1})^{-1}(M) = T(M)$ is open. $\square$

Theorem 7.4 (Closed Graph Theorem). *If $T : X \rightarrow Y$ is a linear operator between Banach space, and the graph*

$$\Gamma(T) = \{(x, Tx) : x \in X\} \subset X \times Y$$

is closed, then T is bounded.

Proof. Note $\Gamma(T) \subset X \times Y$ is a Banach space with the norm $\|(x, Tx)\| = \|x\| + \|Tx\|$. This is because if (x_n, Tx_n) is a Cauchy sequence then both x_n and Tx_n are Cauchy. So $x_n \rightarrow x$ and $Tx_n \rightarrow y$. But since the graph is closed, $y = Tx$ and $(x_n, Tx_n) \rightarrow (x, y)$.

Define $P_1 : \Gamma(T) \rightarrow X$ the projection to X , which is clearly a bijection. Therefore by the inverse mapping theorem, P_1^{-1} is continuous. So if $x_n \rightarrow x$ then $P_1^{-1}x_n \rightarrow P_1^{-1}x$, in other words $(x_n, Tx_n) \rightarrow (x, Tx)$. Hence $Tx_n \rightarrow Tx$, showing that T is continuous. Hence T is bounded by Prop 7.7. $\square$

Corollary 7.5 (Hellinger-Toeplitz). *Let $A : H \rightarrow H$ be an operator and H a Hilbert space. Suppose $(x, Ay) = (Ax, y)$ for all $x, y \in H$, then A is bounded.*

Proof. We show that $\Gamma(A)$ is closed. Suppose $(x_n, Ax_n) \rightarrow (x, y)$. Then for all $z \in H$, we have $(z, y) = \lim_n (z, Ax_n) = \lim_n (Az, x_n) = (Az, x) = (z, Ax)$. Therefore $y = Ax$. Now apply the closed graph theorem, and the fact that continuity is equivalent to boundedness in a Banach space. $\square$

In general, we denote $D(T)$ the domain of an operator T from a subspace of a Banach space to another Banach space, i.e., $T : D(T) \subset X \rightarrow Y$.

Definition 7.6. We say that T is bounded if there exists constant $M > 0$ such that $\|Tx\| \leq M\|x\|$ for all $x \in D(T)$.

A simple observation is

Proposition 7.7. The following statements are equivalent:

1. T is continuous;
2. T is continuous at 0;
3. T is bounded.

Proof. 1 and 2 are equivalent by linearity of T . Suppose 2, then for all $\epsilon > 0$, there exists $\delta > 0$ such that $\|x\| < \delta$ implies $\|Tx\| < \epsilon$. For all $x \in D(T)$, we write $y = r \cdot \frac{x}{\|x\|}$ where $r < \delta$ is fixed, then $\|Ty\| < \epsilon$, so $\|Tx\| \cdot \frac{r}{\|x\|} < \epsilon$, which implies $\|Tx\| < \frac{\epsilon}{r}\|x\|$. Then $M = \epsilon/r$. Hence T is bounded. So 2 implies 3. Finally, every bounded linear operator is continuous, so 3 implies 1. $\square$

Clearly, for a bounded operator, T can be extended continuously to $\text{Cl}[D(T)]$. From now on, we will always assume that

$$\text{Cl}[D(T)] = X.$$

We define the operator norm

$$\|T\| := \sup_{\|\psi\|=1} \|T\psi\|.$$

It is easy to see that

- $\|T\psi\| \leq \|T\| \cdot \|\psi\|$ for any $\psi \in X$.
- The triangle inequality: $\|T + S\| \leq \|T\| + \|S\|$.
- The space of linear bounded operators in a Banach space is itself a Banach space, with the operator norm.

Definition 7.8. Let $T : X \rightarrow Y$ be a linear operator, where X is a normed space. We can define the dual space X^* of X to be the set of all continuous linear maps $x^* : X \rightarrow \mathbf{C}$.

Exercise: X^* is always a Banach space even if X is just a normed space (not necessarily Banach).

Definition 7.9. We can define the adjoint $T^* : Y^* \rightarrow X^*$ of T as follows: $T^*(y^*) : x \mapsto y^*(Tx)$, for any $y^* \in Y^*$.

Theorem 7.10 (Hahn-Banach). *If V is a subspace of a linear space X , and we have a sublinear function $p : X \rightarrow \mathbf{R}$, i.e., $p(x + y) \leq p(x) + p(y)$. Then if $T : V \rightarrow \mathbf{C}$ is linear such that $|Tx| \leq p(x)$, there exists a linear extension $T' : X \rightarrow \mathbf{C}$ of T such that $|T'(x)| \leq p(x)$.*

Proof. Omitted. Check any functional analysis text for details. The proof is an application of Zorn's lemma. $\square$

Corollary 7.11. *Given $x \in X$, there exists $y^* \in X^*$ with norm 1, such that $y^*(x) = \|x\|$.*

Proposition 7.12. $T^* : Y^* \rightarrow X^*$ is a bounded operator.

Proof. By definition, we have $\|T^*y^*(x)\| = \|y^*(Tx)\| \leq \|y^*\| \cdot |Tx| \leq \|y^*\| \cdot \|T\| \cdot \|x\|$. Taking the sup over all x with norm 1, we see $\|T^*y^*\| \leq \|y^*\| \cdot \|T\|$. In other words, $\|T^*\| \leq \|T\|$. $\square$

Definition 7.13. Given $T : X \rightarrow X$ a linear bounded operator. We define the **resolvent** set $\rho(T)$ of T , as

$$\rho(T) = \{\lambda \in \mathbf{C} \mid \lambda I - T : X \rightarrow X \text{ has a bounded inverse}\}.$$

Note that this definition makes sense for a linear operator defined on $D(T) \subset X$. We define for $\lambda \in \rho(T)$, the resolvent of T at λ to be

$$R_\lambda(T) = (\lambda I - T)^{-1}.$$

We define the **spectrum** $\sigma(T)$ of T to be $\sigma(T) = \mathbf{C} \setminus \rho(T)$. In other words,

$$\sigma(T) = \{\lambda \in \mathbf{C} \mid \lambda I - T : X \rightarrow X \text{ has NO bounded inverse}\}.$$

There are 3 possibilities, in all situations $\lambda I - T$ fails to have a bounded inverse:

1. If $\ker(\lambda I - T) \neq 0$, we define $\sigma_p(T) := \{\lambda \in \mathbf{C} : \ker(\lambda I - T) \neq 0\}$. Thus this is the set of eigenvalues of T .
2. If $\ker(\lambda I - T) = 0$, but $\text{im}(\lambda I - T)$ is dense in X . Then the inverse $(\lambda I - T)^{-1} : \text{im}(\lambda I - T) \rightarrow X$ is NOT bounded. Because if it were bounded, we can extend $R_\lambda(T)$ to a continuous map, which will be the inverse of $\lambda I - T$. Then λ will be in the **continuous spectrum** of T , denoted by $\sigma_c(T)$.
3. If $\ker(\lambda I - T) = 0$, but $\text{im}(\lambda I - T)$ is NOT dense in X . Then λ will be in the **residue spectrum** of T , denoted by $\sigma_r(T)$.

Therefore, we can write $\sigma(T)$ as a disjoint union of these three things:

$$\sigma(T) = \sigma_p(T) \cup \sigma_c(T) \cup \sigma_r(T).$$

In the finite dimensional case, by the rank-nullity theorem, case 2 and 3 cannot happen. Hence the spectrum $\sigma(T)$ for a finite-dimensional operator is the same as its eigenvalues.

Remark 7.14 (Discrete and Essential Spectrum). There is another decomposition of $\sigma(T)$: we have the discrete spectrum $\sigma_d(T)$, which is the set of eigenvalues λ of T such that there exists $\delta > 0$ so that $B(\lambda, \delta) \cap \sigma(T) = \{\lambda\}$; and there is the essential spectrum $\sigma_e(T)$, which is defined as $\sigma_e(T) := \sigma(T) \setminus \sigma_d(T)$.

Remark 7.15 (Bounded and Scattering States). We denote $B_R := B(0, R)$ and B_R^c its complement, for $R > 0$.

Let H be a Hilbert space, and $T : D(T) \subset H \rightarrow H$ a closed operator. We can decompose H into two parts: first we have bound states H_b , where H_b is the set of eigenvectors of T , and continuous space $H_c := (H_b)^\perp$. Therefore

$$H = H_b \oplus H_c.$$

Bound states in physics will correspond to elements in H_b . These are the states $\psi_0 \in H_b$, such that for all $\epsilon > 0$, there exists $R > 0$, so that we have

$$\frac{1}{L} \int_0^L \|e^{iTt}\psi_0\|_{L^2(B_R^c)} < \epsilon$$

for all L . The scattering states in physics are characterized by states $\psi_0 \in H_c$, where given $\epsilon > 0$, for all $R > 0$ we have

$$\lim_{L \rightarrow \infty} \frac{1}{L} \int_0^L \|e^{iTt}\psi_0\|_{L^2(B_R)} < \epsilon.$$

Example 7.16. Define $\ell_1 := \{(x_n)_n : \sum_j |x_j| < \infty\}$. Define the left shift and right shift via

$$L(x_1, x_2, \dots) = (x_2, x_3, \dots)$$

and

$$R(x_1, x_2, \dots) = (0, x_1, x_2, \dots).$$

Then one can show $R = L^*$, and $\ell_1^* = \ell_\infty$. This example is worked out in detail in [5] p.192-194.

Now we study some properties of the resolvent.

Proposition 7.17 (1st Resolvent Identity). We have

$$R_\lambda(T) - R_\mu(T) = -(\lambda - \mu)R_\lambda(T)R_\mu(T)$$

and

$$R_\lambda(T)R_\mu(T) = R_\mu(T)R_\lambda(T).$$

Proof. The second identity comes from the fact that $(\lambda - T)(\mu - T) = (\mu - T)(\lambda - T)$ by expansion, hence take the inverse we get the equation. The first identity also comes from straight computation:

$$\begin{aligned} R_\lambda(T) - R_\mu(T) &= (\lambda - T)^{-1} - (\mu - T)^{-1} \\ &= (\lambda - T)^{-1}[(\mu - T) - (\lambda - T)](\mu - T)^{-1} \\ &= R_\lambda(T)(\mu - \lambda)R_\mu(T). \end{aligned}$$

□

Proposition 7.18 (Von-Neumann Series). Given $T \in B(X)$ (recall that this is the set of bounded linear operator from X to itself) where X is Banach. Then if $\|T\| < 1$, then $I - T$ is invertible, and

$$(I - T)^{-1} = \sum_{n=0}^{\infty} T^n =: S.$$

Proof. Since $\|T\| < 1$, we claim that $I - T$ is injective: If $x \in X$ such that $x = Tx$, then

$$\|x\| = \|Tx\| \leq \|T\| \cdot \|x\| < \|x\|$$

which would imply $x = 0$. Hence $\ker(I - T) = 0$, as desired. Now denote

$$S_n := \sum_{j=0}^n T^j.$$

Then by the triangle inequality, $(S_n)_{n=1}^{\infty}$ is a Cauchy sequence, and $S_n \rightarrow S$. Then

$$(I - T)S_n = \sum_{j=0}^n T^j - \sum_{j=0}^n T^{j+1} = I - T^{n+1}.$$

which converges to I as $n \rightarrow \infty$, and note that $(I - T)S_n = S_n(I - T)$, hence we have

$$(I - T)S = S(I - T) = I.$$

This concludes the proof. $\square$

The first application of the proposition above is

Theorem 7.19. Let $T \in B(X)$ where X is a Banach space. Then $\rho(T)$ is open and $\rho(T) \neq \emptyset$. Also we have $\sigma(T) \neq \emptyset$.

Proof. To show that $\rho(T) \neq \emptyset$, take λ such that $|\lambda| > \|T\|$. Then $\lambda I - T = \lambda(I - T/\lambda)$. Since $\|T/\lambda\| < 1$ by assumption, by the previous proposition, $I - T/\lambda$ is invertible, hence $\lambda \in \rho(T)$. Hence $\rho(T) \neq \emptyset$, and $\sigma(T) \subset \text{Cl}[B(0, \|T\|)]$.

Now assume that $\lambda_0 \in \rho(T)$. Then for any $\lambda \in \mathbf{C}$, we have

$$\lambda I - T = (\lambda - \lambda_0)I + \underbrace{(\lambda_0 I - T)}_{\text{invertible}} = (\lambda_0 I - T)[-(\lambda_0 - \lambda)R_{\lambda_0}(T) + I].$$

If $|\lambda - \lambda_0| \cdot \|R_{\lambda_0}(T)\| < 1$, then $I - (\lambda_0 - \lambda)R_{\lambda_0}(T)$ is invertible. Hence $\lambda I - T$ is invertible, i.e., $\lambda \in \rho(T)$. And we have

$$R_{\lambda}(T) = [I + (\lambda - \lambda_0)R_{\lambda_0}(T)]^{-1}R_{\lambda_0}(T) = R_{\lambda_0}(T)[I + (\lambda - \lambda_0)R_{\lambda_0}(T)]^{-1}.$$

This implies $B(\lambda_0, \epsilon) \subset \rho(T)$ for $\epsilon < 1/\|R_{\lambda_0}(T)\|$. Hence $\rho(T)$ is open.

Now it remains to show $\sigma(T) \neq \emptyset$. From the 1st resolvent identity, we have that the map $R_{\lambda}(T) : \rho(T) \rightarrow B(X)$ is analytic, in the sense that the derivative exists:

$$\lim_{\lambda \rightarrow \mu} \frac{R_{\lambda}(T) - R_{\mu}(T)}{\lambda - \mu} = R_{\lambda}(T)R_{\mu}(T) \rightarrow R_{\mu}(T)^2.$$

If $\sigma(T) = \emptyset$, then $\rho(T) = \mathbf{C}$. Define, for fixed $y^* \in X^*$, $x \in X$, the map

$$\begin{aligned} f : \mathbf{C} &\rightarrow \mathbf{C} \\ \lambda &\mapsto y^*(R_\lambda(T)x). \end{aligned}$$

If $|\lambda| > \|T\|$, we have

$$\|R_\lambda(T)\| = \frac{1}{|\lambda|} \|(I - T/\lambda)^{-1}\| \leq \frac{1}{|\lambda|} \cdot \frac{1}{1 - \|T\|/|\lambda|} \xrightarrow{\lambda \rightarrow \infty} 0.$$

Hence f is analytic in $\mathbf{C}$ and goes to zero for large λ . Thus f is bounded, so f is constant by Liouville's theorem, which is a contradiction. $\square$

Proposition 7.20 (2nd Resolvent Identity). We have

$$R_\lambda(T) - R_\lambda(S) = R_\lambda(T)(T - S)R_\lambda(S).$$

Proof. Straightforward calculation:

$$R_\lambda(T) - R_\lambda(S) = (\lambda I - T)^{-1}[(\lambda I - S) - (\lambda I - T)](\lambda I - S)^{-1} = R_\lambda(T)(T - S)R_\lambda(S).$$

This concludes the proof. $\square$

Definition 7.21. Given a subspace $M \subset X$, we define $M^\perp = \{x^* \in X^* : M \subset \ker x^*\}$.

This notation is begging us to relate it with the orthogonal complement of a subspace in linear algebra. Indeed: If V is a finite-dimensional inner product space, and $W \subset V$ is a subspace, then $W^\perp = \{x^* \in V : W \subset \ker x^*\}$ by definition. By the Riesz representation theorem, x^* can be identified with the linear functional $\langle y, \cdot \rangle : V \rightarrow \mathbf{C}$ for some $y \in V$, and by definition y is in the orthogonal complement of W . Therefore there is an isomorphism between $W^\perp$ and the orthogonal complement of W in V .

Proposition 7.22. Let T^* be defined as in Definition 7.9. Then we have

$$\ker(T^*) = \operatorname{im}(T)^\perp.$$

Proof. We have $T^*y^* = 0$ iff $T^*y^*(x) = 0$ for all x , iff $y^*(Tx) = 0$ for all x , iff $y^* \in \operatorname{im}(T)^\perp$. $\square$

Proposition 7.23. Show that if $V \subset X$ is a subspace of a Banach space, then $\operatorname{Cl}[V] = X$ iff $V^\perp = 0$.

Proof. ($\Leftarrow$) If $\operatorname{Cl}[V] \neq X$, then there exists $y \in X$, $y \notin \operatorname{Cl}[V]$. So exists $y^* \in X^*$ such that $y^*(y) = 1$ and $y^*(V) = 0$. Hence $y^* \in V^\perp$.

($\Rightarrow$) If there exists $y^* \in V^\perp$ nonzero, then $y^*(V) = 0$, so $y^*(x) = 0$ for all $x \in \operatorname{Cl}[V]$. But since y^* is not identically zero by assumption, there must exist $x \in X$ such that $y^*(x) \neq 0$. Thus $\operatorname{Cl}[V] \neq X$. $\square$

Remark 7.24 (Banach Space adjoint and Hilbert space adjoint). If $T \in B(X, Y)$ and X, Y Banach, then

$$(\lambda - T)^*y^*(x) = y^*((\lambda - T)x) = \lambda y^*(x) - y^*(Tx) = (\lambda - T^*)y^*(x).$$

Hence

$$(\lambda - T)^* = \lambda - T^*.$$

That is, the adjoint of the operator $\lambda - T$ is actually $\lambda - T^*$ in a Banach space. But this is somehow inconsistent with our intuition in quantum mechanics, where if $T : H \rightarrow H$ is a bounded operator between Hilbert spaces, then we should have $(\lambda - T)^\dagger = \lambda^* - T^\dagger$ where $T^\dagger$ denote the Hermitian adjoint of T . The reason is the following: Given a Hilbert space H , by the Riesz representation theorem on Hilbert space, there is an embedding of H onto its dual H^* , given by $\Phi : H \rightarrow H^*$, where $\Phi(x) = (x, \cdot) : H \rightarrow \mathbf{C}$. Note that this is an antilinear map. Then the Hermitian adjoint $T^\dagger : H \rightarrow H$ is actually defined such that the following diagram commute:

$$\begin{array}{ccc} H^* & \xrightarrow{T^*} & H^* \\ \Phi \uparrow & & \uparrow \Phi \\ H & \xrightarrow{T^\dagger} & H \end{array}$$

In other words, $T^\dagger := \Phi^{-1}T^*\Phi$. Now let $x, y \in H$. Then we have

$$(T^\dagger y, x) = (\Phi^{-1}T^*\Phi y, x) = \Phi(\Phi^{-1}T^*\Phi y)(x) = (T^*\Phi y)x = \Phi(y)(Tx) = (y, Tx).$$

Hence, we see that indeed $T^\dagger : H \rightarrow H$ as the usual property in quantum mechanics. Now using this property, one can show that

$$(\lambda - T)^\dagger = \lambda^* - T^\dagger$$

because the inner product $(\cdot, \cdot) : H \times H \rightarrow \mathbf{C}$ is antilinear in the 1st slot. Sometimes, we abuse the notation a little bit and do not distinguish between T^* and $T^\dagger$.

Note that if T is unbounded, then self-adjoint is different from the operator having the property $(Tu, v) = (u, Tv)$ for all $u, v \in H$ (for bounded operators on a Hilbert space two notions are equivalent). The latter property is called symmetric.

Theorem 7.25. *Let $T \in B(H)$ where H is a Hilbert space, and $T^* = T$, then $\sigma_r(T) = \emptyset$, and $\sigma(T) \subset \mathbf{R}$.*

Proof. Given $\lambda = \alpha + i\beta \in \mathbf{C}$, $\beta \neq 0$ and $x \in H$. Then

$$\begin{aligned} \|(T - \lambda)x\|^2 &= \|(T - \alpha)x - i\beta x\|^2 \\ &= \|(T - \alpha)x\|^2 + \beta^2\|x\|^2. \end{aligned} \quad (*)$$

The second equality is just straightforward computation using inner product, using the (anti)linearity of the inner product and the assumption that $T^* = T$.

Now, if $\beta \neq 0$, we claim that $T - \lambda$ is injective: If $x \in \ker(T - \lambda)$, then since $\|(T - \lambda)x\|^2 = 0$, we must have $\beta^2\|x\|^2 = 0$ hence $x = 0$, proving injectivity.

Also, note that if $\beta \neq 0$, then $\lambda^* = \alpha - i\beta$ also has nonzero imaginary part. Hence by the same argument, we know $\ker(T - \lambda^*) = 0$. By Proposition 7.22, we conclude then $\text{im}(T - \lambda)^\perp = 0$ for $(T - \lambda)^* = T - \lambda^*$ by the remark above and self-adjointness. Then by Proposition 7.23, we conclude that $\text{im}(T - \lambda I)$ is dense in H . Hence $T - \lambda I : H \rightarrow H$ is injective and has a dense image, thus we can define its inverse $(T - \lambda)^{-1}$.

By (*) we have $\|(T - \lambda)x\|^2 \geq \beta^2\|x\|^2$. Now replacing x by $(T - \lambda)^{-1}x$, we have

$$\|(T - \lambda)^{-1}x\| \leq \frac{1}{\beta}\|x\|.$$

This shows that $(T - \lambda)^{-1}$ is bounded. Therefore, we have proven that if $\lambda \in \mathbf{C}$ is not real, then $T - \lambda$ has a bounded inverse, thus $\lambda \notin \sigma(T)$. This shows that $\sigma(T) \subset \mathbf{R}$.

Finally, we show that $\sigma_r(T) = \emptyset$. We argue by contradiction: If not, then pick $\lambda \in \sigma_r(T) \subset \mathbf{R}$, as the spectrum must be a real subset. Then $(\lambda - T)^{-1}$ exists, but its domain, i.e., $\text{im}(\lambda - T)$ is not dense in H . Then by Proposition 7.23, we must have $\text{im}(\lambda - T)^\perp \neq 0$. But again by Proposition 7.22, we must have $\ker(\lambda - T) \neq 0$. Hence $\lambda \in \sigma_p(T)$ by definition, which is a contradiction, since $\sigma_p(T) \cap \sigma_r(T) = \emptyset$. $\square$

Proposition 7.26. If $T = T^*$ then for $\lambda, \mu \in \sigma_p(T)$ distinct reals, the corresponding eigenvectors of T are orthogonal.

Proof. Let $Tx = \lambda x$ and $Ty = \mu y$. Since $(y, Tx) = \lambda(y, x)$ and $(Ty, x) = \mu(y, x)$. So we have $(\lambda - \mu)(y, x) = 0$. $\square$

Definition 7.27. For $T : D(T) \subset H \rightarrow H$ bounded (or unbounded), and if $D(T)$ is dense in H , then we define $D(T^*) \subset H$ as follows:

$$D(T^*) = \{y \in H : x \mapsto (y, Tx), x \in D(T) \text{ is bounded}\}.$$

Proposition 7.28. Let $T \in B(H)$ such that $D(T)$ is dense in H . Then $T = T^*$ (i.e., T is self-adjoint) iff $D(T) = D(T^*)$.

Proof. Note that $D(T) \subset D(T^*)$ follows directly from the Cauchy-Schwarz inequality even without the condition $T = T^*$: $|(y, Tx)| \leq \|y\| \cdot \|Tx\| \leq \|y\| \cdot \|T\| \cdot \|x\|$. Hence the map $x \mapsto (y, Tx)$ has norm at most $\|y\| \cdot \|T\| < \infty$, for T is bounded.

Now suppose $D(T) = D(T^*)$. Then by denseness of $D(T)$, the map $x \mapsto (y, Tx)$ extends to the whole H . Again this map is a bounded linear functional on H by the previous paragraph, hence by the Riesz representation theorem, there exists an element $T^*y \in H$ such that $(T^*y, x) = (y, Tx)$, and T^*y is uniquely determined by this condition since $D(T)$ is dense. It is not hard to show that T^* defined this way agrees with Definition 7.9. $\square$

The Proposition above motivates us to make the following definition of adjoint for unbounded operator:

Definition 7.29. When $T : H \rightarrow H$ is unbounded, and $D(T^*)$ is dense, the operator T^* with domain $D(T^*)$ defined by

$$(T^*y, x) = (y, Tx)$$

is called the adjoint of T .

7.2 Spectral Theorem for Bounded Operators

Now we start to set up the tools to prove the spectral theorem for bounded operators. The first result we need is

Theorem 7.30 (Spectral theorem for Polynomials). *Let $P(z) = a_0 + a_1z + \cdots + a_nz^n$ be a polynomial, and $T : X \rightarrow X$ an arbitrary bounded operator. Define $P(T) = a_0I + a_1T + \cdots + a_nT^n$. Then we have*

$$P(\sigma(T)) = \sigma(P(T)).$$

Proof. Let $\lambda \in \sigma(T)$. Since $z = \lambda$ is a root of the polynomial $P(z) - P(\lambda)$, we have $P(z) - P(\lambda) = (z - \lambda)Q(z)$ for some polynomial Q . So $P(T) - P(\lambda) = (T - \lambda)Q(T)$. Since $T - \lambda$ has no inverse, neither does $P(T) - P(\lambda)$. That is, $P(\lambda) \in \sigma(P(T))$. This proves one inclusion.

Conversely, let $\mu \in \sigma(P(T))$ and let $\lambda_1, \dots, \lambda_n$ be roots of $P(z) - \mu$, that is, we have the factorization

$$P(z) - \mu = a \prod_{k=1}^n (z - \lambda_k)$$

for some nonzero complex number a . If $\lambda_1, \dots, \lambda_n \notin \sigma(T)$, then we have

$$(P(T) - \mu)^{-1} = a^{-1} \prod_{k=1}^n (T - \lambda_k)^{-1}.$$

So we conclude that there exists at least one k such that $\lambda_k \in \sigma(T)$. That is, $\mu = P(\lambda_k) \in P(\sigma(T))$. This concludes the proof. $\square$

Our first goal is to extend the above result to continuous functions, not just polynomials. Along the way, we need a few results:

Theorem 7.31. *If $T \in B(X)$, then we define the **spectral radius** of T to be*

$$r(T) = \sup_{\lambda \in \sigma(T)} |\lambda|.$$

Then we have

$$r(T) = \lim_{n \rightarrow \infty} \|T^n\|^{1/n}.$$

Proof. By spectral theorem for polynomials, $\sigma(T^n) = \sigma(T)^n$. Note that for any $S \in B(X)$, we have $r(S) \leq \|S\|$ (otherwise suppose $r(S) > \|S\|$. Then there exists $\mu \in \sigma(T)$ such that $\mu > \|S\|$. But then $\mu - S$ would be invertible with the inverse given by the von Neumann series, a contradiction). Hence $r(T^n) \leq \|T^n\|$. But since $\sigma(T^n) = \sigma(T)^n$, by definition of spectral radius,

$$r(T)^n = r(T^n) \leq \|T^n\|.$$

Hence $r(T) \leq \|T^n\|^{1/n}$ for all n . Therefore,

$$r(T) \leq \liminf_n \|T^n\|^{1/n} \leq \limsup_n \|T^n\|^{1/n}.$$

We now want to prove that $\limsup_n \|T^n\|^{1/n} \leq r(T)$, which would show that $r(T) = \lim_n \|T^n\|^{1/n}$. Pick $A \geq \|T\| \geq r(T)$. So $A \in \rho(T)$. And we have

$$R_\lambda(T) = (\lambda - T)^{-1} = \sum_{n=0}^{\infty} \frac{T^n}{\lambda^{n+1}}$$

for all $|\lambda| > A$, $\lambda \in \mathbf{C}$. Now for all $y^* \in X^*$, and for all $x \in X$, we define

$$f(\lambda) = y^*(R_\lambda(T)x) = \sum_{n=0}^{\infty} \frac{y^*(T^n x)}{\lambda^{n+1}}.$$

Clearly f is analytic in $\mathbf{C} \setminus D(0, r(T))$, since it follows from the 1st resolvent formula that f is complex differentiable in $\mathbf{C} \setminus D(0, r(T))$. This shows that

$$f(\lambda) = \sum_{n=0}^{\infty} \frac{y^*(T^n x)}{\lambda^{n+1}}$$

for all $|\lambda| > r(T)$ by the uniqueness of Laurent series. This means that the sequence $y^*(T^n x)/\lambda^{n+1} \rightarrow 0$ as $n \rightarrow \infty$ for all $|\lambda| > r(T)$. In particular,

$$|y^*(T^n x)| \leq C_\lambda |\lambda|^{n+1}$$

for some constant $C_\lambda > 0$ and for all n . But since $|y^*(T^n x)| = (T^*)^n(y^*)(x)$, this shows that $(T^*)^n(y^*)(x)$ is bounded for all x . By the uniform boundedness theorem, this means there exists M such that $(T^*)^n(y^*) \leq M|\lambda|^{n+1}$ for all $y^* \in X^*$ and for all n . Thus $\|(T^*)^n\| \leq C|\lambda|^{n+1}$. And since $T = T^*$, we have $\lim_n \|T^n\|^{1/n} \leq \lim_n C^{1/n} |\lambda|^{1+1/n} = |\lambda|$ for all $|\lambda| > r(T)$. Thus $\limsup_n \|T^n\|^{1/n} \leq |\lambda|$ for all $|\lambda| > r(T)$. Thus $\limsup_n \|T^n\|^{1/n} \leq r(T)$. This proves the theorem. $\square$

We now show that if it happens that $T \in B(H)$ is a bounded self-adjoint operator on a Hilbert space, then the formula for spectral radius above reduces to $r(T) = \|T\|$. To prove this result we first need a

Lemma 7.32. *If $A \in B(H)$, then $\|A^*\| = \|A\|$ and $\|A^*A\| = \|A\|^2$.*

Proof. Recall that the operator norm is computed as $\|A\| := \sup_{\|\psi\|=1} \|A\psi\|$. Furthermore, for any $\phi \in H$, one has $\|\phi\| = \sup_{\|\chi\|=1} |(\chi, \phi)|$. Thus we have

$$\|A\| = \sup_{\|\phi\|=\|\psi\|=1} |(\phi, A\psi)|.$$

From this, we get that

$$\|A^*\| = \sup_{\|\phi\|=\|\psi\|=1} |(\phi, A^*\psi)| = \sup_{\|\phi\|=\|\psi\|=1} |(A\phi, \psi)| = \sup_{\|\phi\|=\|\psi\|=1} |(\psi, A\phi)| = \|A\|.$$

Meanwhile, we have $\|A^*A\| \leq \|A^*\|\|A\| = \|A\|^2$. On the other hand, we have

$$\|A^*A\| = \sup_{\|\phi\|=\|\psi\|=1} |(\phi, A^*A\psi)| = \sup_{\|\phi\|=\|\psi\|=1} |(A\phi, A\psi)| \geq \sup_{\|\psi\|=1} |(A\psi, A\psi)| = \|A\|^2.$$

This proves the lemma. $\square$

Theorem 7.33. *If $T \in B(H)$ and if $T = T^*$, then we have $r(T) = \|T\|$.*

Proof. Applying the previous lemma repeatedly and using $T^* = T$, we see that $\|T^{2^n}\| = \|T\|^{2^n}$ for all n . Then the claim follows immediately from the spectral radius formula (Theorem 7.31), since

$$r(T) = \lim_n \|T^n\|^{1/n} = \lim_n \|T^{2^n}\|^{1/2^n} = \lim_n \|T\|^{2^n/2^n} = \|T\|.$$

This concludes the proof. $\square$

Theorem 7.34. Suppose $T \in B(H)$ is self adjoint: $T^* = T$. Then there exists a unique bounded linear map $\Psi : \mathcal{C}(\sigma(T), \mathbf{R}) \rightarrow B(H)$ from the space of continuous real functions on $\sigma(T)$ to $B(H)$, defined by $\Psi : f \mapsto f(T)$, such that when $f(\lambda) = \lambda^m$ is a polynomial, we have $f(T) = T^m$. The map Ψ is called the **(real-valued) functional calculus for T** . The continuous functional calculus also satisfies the following properties:

- (1) *Multiplicativity:* For all f, g we have $(fg)(T) = f(T)g(T)$.
- (2) *The map $\Psi : f \mapsto f(T)$ is an isometry.*
- (3) *For all $f \in \mathcal{C}(\sigma(T), \mathbf{R})$, the operator $f(T)$ is self-adjoint.*
- (4) *If $f \geq 0$ on $\sigma(T)$, then $f(T) \geq 0$, i.e., $(\psi, f(T)\psi) \geq 0$ for all $\psi \in H$.*
- (5) *$\sigma(f(T)) = f(\sigma(T))$.*

Proof. Note that if T is self-adjoint, and p a real-valued polynomial, then $p(T)$ is also self-adjoint, where $p(T)$ is defined as in Theorem 7.30. Therefore, we may define Ψ as follows: First, for a polynomial $p \in \mathcal{C}(\sigma(T), \mathbf{R})$, we define $\Psi : p \mapsto p(T)$ as in Theorem 7.30. Now, using Theorem 7.30 and Theorem 7.33, we have that

$$\|p(T)\| = r(p(T)) = \sup_{\lambda \in \sigma(p(T))} |\lambda| = \sup_{\lambda \in \sigma(T)} |p(\lambda)| = \|p\|_\infty.$$

That is, the map $\Psi : p \mapsto p(T)$ is an isometry from the space of polynomials on $\sigma(T)$ (with supremum norm) to the space of bounded operators on H . Since T is bounded, $\sigma(T)$ is compact (for it is closed and bounded by Theorem 7.19). Now, by the Stone Weierstrass theorem, polynomials are dense in $\mathcal{C}(\sigma(T), \mathbf{R})$. Hence the map $\Psi : p \mapsto p(T)$ extends uniquely to a bounded linear map $\Psi : \mathcal{C}(\sigma(T), \mathbf{R}) \rightarrow B(H)$.

Now we prove the listed properties:

- (1) Since the statement holds for polynomials, by taking limits, then the statement also holds for all $f, g \in \mathcal{C}(\sigma(T), \mathbf{R})$. The detailed proof is an simple application of triangle equality and is left as an exercise.
- (2) Clearly Ψ is linear. If $p_n \rightarrow f$ uniformly, on one hand we have $\|p_n(T)\| = \|p_n\|_\infty$ by the previous discussions, but since $\|f_n(T) - p_n(T)\| \rightarrow 0$, we have $\|p_n(T)\| \rightarrow \|f(T)\|$ by the reversed triangle inequality. Since we also have $\|p_n\|_\infty \rightarrow \|f\|_\infty$, we have that $\|f(T)\| = \|f\|_\infty$.
- (3) Since for any real polynomial p , the operator $p(T)$ is self-adjoint because $T^* = T$, the statement also follows by taking limits: Take a sequence p_n of polynomials such that $p_n \rightarrow f$ in supremum norm. Then by (2) we have $f(T) = \lim_n p_n(T)$. Therefore

$$(f(T)\phi, \psi) = \lim_n (p_n(T)\phi, \psi) = \lim_n (\phi, p_n(T)\psi) = (\phi, f(T)\psi).$$

Hence we conclude that $f(T)$ is self-adjoint.

- (4) If $f \geq 0$ on $\sigma(T)$, that means $f : \sigma(T) \rightarrow \mathbf{R}$ is continuous and nonnegative. Then there exists $g \in \mathcal{C}(\sigma(T), \mathbf{R})$ such that $f = g^2$ (just take the square root of f). Then by multiplicativity, $f(T) = g(T)^2$, with $g(T)$ self-adjoint. Then

$$(\psi, f(A)\psi) = \|g(T)\psi\|^2 \geq 0$$

for all $\psi \in H$, as desired.

- (5) Postponed. This is proved in Theorem 7.36.

□

Theorem 7.35. *If $T : D(T) \subset H \rightarrow H$ self-adjoint, closed, and bounded, then $\lambda \in \sigma(T)$ iff there exists a sequence $\psi_n \in D(T)$, $\|\psi_n\| = 1$, such that $\|(T - \lambda)\psi_n\| \rightarrow 0$.*

Proof. If $\lambda \in \sigma(T) = \sigma_p(T) \cup \sigma_c(T)$, then either there exists $\psi \neq 0$ such that $(T - \lambda)\psi = 0$ or there exists $(T - \lambda)^{-1} : \text{im}(T - \lambda) \rightarrow H$ that is not bounded. In the first case, pick $\psi_n = \psi/\|\psi\|$ to be the constant sequence, done. In the latter case, since there exists $\psi_n \in \text{im}(T - \lambda)$ with $\|\psi_n\| = 1$ such that $\|(T - \lambda)^{-1}\psi_n\| \geq n$, or $\|(T - \lambda)^{-1}\psi_n\| \geq n\|\psi_n\|$. Now define $\phi_n := (T - \lambda)^{-1}\psi_n \in D(T)$. Then

$$\|(T - \lambda)\phi_n\| = \|\psi_n\| \leq \frac{1}{n}\|\phi_n\|.$$

Now let $u_n := \phi_n/\|\phi_n\|$. We see that $\|u_n\| = 1$ and $\|(T - \lambda)u_n\| \rightarrow 0$.

Conversely, if $\lambda \notin \sigma(T)$, then $(T - \lambda)^{-1}$ is bounded. If there exists a sequence (ψ_n) such that $\|\psi_n\| = 1$, and $(T - \lambda)\psi_n \rightarrow 0$ by assumption, then on one hand by continuity of the resolvent $\|(T - \lambda)^{-1}(T - \lambda)\psi_n\| \rightarrow 0$, but clearly this is just the sequence $\|\psi_n\|$, which is constantly 1, contradiction. □

Theorem 7.36. *Let $T^* = T$ and $T \in B(H)$. Then for any $f \in \mathcal{C}(\sigma(T), \mathbf{R})$, we have $\sigma(f(T)) = f(\sigma(T))$.*

Proof. We first show $\sigma(f(T)) \subset f(\sigma(T))$. If $\lambda \notin f(\sigma(T))$, then consider $g = (f - \lambda)^{-1} \in \mathcal{C}(\sigma(T), \mathbf{R})$. We can define $g(T)$. Since $g \cdot (f - \lambda) = (f - \lambda) \cdot g = \text{id}$ on $\mathcal{C}(\sigma(T), \mathbf{R})$, we have by multiplicativity that $g(T)(f(T) - \lambda) = (f(T) - \lambda)g(T) = I$. Thus $g(T) = (f(T) - \lambda)^{-1}$. Thus $\lambda \in \rho(f(T))$.

Conversely, let $\lambda \in f(\sigma(T))$. Then there exists $\mu \in \sigma(T)$ such that $f(\mu) = \lambda$. Since f is continuous, given any n there exists $\delta_n > 0$ such that $|f(x) - \lambda| < 1/n$ for all $x \in (\mu - \delta_n, \mu + \delta_n) \cap \sigma(T)$. Hence we may construct (see Fig. 7.1) a continuous function $g_n \in \mathcal{C}(\sigma(T), \mathbf{R})$, with $0 \leq g_n \leq 1$, where

$$g_n = \begin{cases} 1 & x \in (\mu - \delta_n/2, \mu + \delta_n/2) \\ 0 & |x - \mu| \geq \delta_n. \end{cases}$$

Then $\|(f - \lambda)g_n\|_\infty \leq 1/n$. Hence

$$\|(f(T) - \lambda)g_n(T)\| \leq \frac{1}{n}.$$

Since $\|g_n(T)\| = \|g_n\|_\infty = 1$, there exists $u_n \in H$, $\|u_n\| = 1$, such that $\|g_n(T)u_n\| \geq 1/2$. Thus define $\psi_n = g_n(T)u_n/\|g_n(T)u_n\|$. Then we have

$$\|(f(T) - \lambda)\psi_n\| = \frac{\|(f(T) - \lambda)g_n(T)u_n\|}{\|g_n(T)u_n\|} \leq \frac{\|(f(T) - \lambda)g_n(T)\|}{\|g_n(T)u_n\|} \leq \frac{2}{n} \rightarrow 0.$$

Hence by the previous theorem, we have $\lambda \in \sigma(f(T))$. Thus $\sigma(f(T)) = f(\sigma(T))$. $\square$

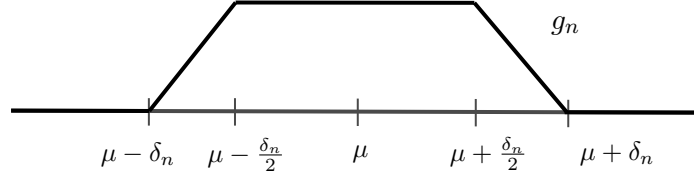


Figure 7.1: The construction of the function g_n .

One can also develop the complex-valued continuous functional calculus for T , i.e., we can consider functions $f \in \mathcal{C}(\sigma(T), \mathbf{C})$. But this requires little additional efforts, since for any complex valued continuous function $f : \sigma(T) \rightarrow \mathbf{C}$, we can always write $f = u + iv$ with functions $u, v \in \mathcal{C}(\sigma(T), \mathbf{R})$. Then we simply define

$$f(T) := u(T) + iv(T)$$

where $u(T), v(T)$ are defined as above. Note that $f(T)^* = u(T) - iv(T)$ since $u(T), v(T)$ are both self-adjoint. From here we can develop the complex-valued functional calculus theory for T similarly to what we did above, and it will not make much difference.

Remark 7.37 (Power Series Expansion and Residue Theorem). If f is analytic on Ω and $\sigma(A) \subset \Omega$. Then if we write $f(z) = \sum_{n=0}^{\infty} a_n(z - z_0)^n$ where R is the radius of convergence of the power series, then we may define

$$f(A) := \sum_{n=0}^{\infty} a_n(A - z_0 I)^n.$$

This definition of $f(A)$ coincides with our definition of $f(A)$ in Theorem 7.34. In fact, one can even show that if γ is a closed cycle in Ω with $\sigma(A)$ contained in the interior part of the complex plane separated by γ (by Jordan curve theorem), then we have a version of the residue theorem:

$$f(A) = \frac{1}{2\pi i} \int_{\gamma} f(z)(zI - A)^{-1} dz.$$

Note that since γ does not touch $\sigma(A)$, for every point $z \in \gamma$, the residue $R_z(A) = (zI - A)^{-1}$ is well-defined. Below, we shall first prove the case for a polynomial function, the general statement will be a consequence of the spectral theorem ([proof here](#)).

Proposition 7.38 (Residue Theorem for Polynomials). Let γ be a closed cycle in Ω with $\sigma(A)$ contained in the interior part of the complex plane separated by γ . Then for any polynomial p , we have that

$$p(A) = \frac{1}{2\pi i} \int_{\gamma} p(z)(z - A)^{-1} dz.$$

Proof. By linearity of integral, it suffices to assume that $p(z) = z^n$. Note that for $n \geq 1$, we have

$$\begin{aligned} \frac{1}{2\pi i} \int_{\gamma} z^n (z - A)^{-1} dz &= \frac{1}{2\pi i} \int_{\gamma} z^{n-1} (z - A + A)(z - A)^{-1} dz \\ &= \frac{1}{2\pi i} \int_{\gamma} z^{n-1} dz + A \frac{1}{2\pi i} \int_{\gamma} z^{n-1} (z - A)^{-1} dz \\ &= \dots \\ &= A^n \frac{1}{2\pi i} \int_{\gamma} (z - A)^{-1} dz. \end{aligned}$$

Hence it suffices to show that $\frac{1}{2\pi i} \int_{\gamma} (z - A)^{-1} dz = I$. Now take $R > \|A\|$. Then since integral of a closed form is homotopy invariant (for a more specific proof, see Fig.7.2), and using von Neumann series, we have

$$\frac{1}{2\pi i} \int_{\gamma} (z - A)^{-1} dz = \frac{1}{2\pi i} \int_{\partial B(0,R)} (z - A)^{-1} dz = \frac{1}{2\pi i} \sum_{n=0}^{\infty} \int_{\partial B(0,R)} \frac{dz}{z^{n+1}} A^n.$$

Note that when $n \geq 1$, the integral $\int_{\partial B(0,R)} dz/z^n = 0$ since the integrand $1/z^m$ for $m \geq 2$ is holomorphic. When $n = 0$, we have $\frac{1}{2\pi i} \int_{\partial B(0,R)} dz/z = 1$ by standard computation. Thus the only term survived in the summand above is just the $n = 0$ term, which is exactly the identity operator I , as desired. $\square$

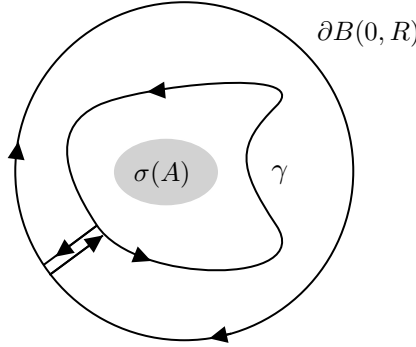


Figure 7.2: A standard argument in complex analysis as suggested in the figure, shows that integral along γ of a holomorphic operator valued function is the same as integral along $\partial B(0, R)$. The argument can be made precise by arguing that the integration along the pass suggested in the figure is zero by Cauchy integral theorem, and we can make the width of the bridge arbitrarily small, then the result will be the integral along γ minus integral along $\partial B(0, R)$.

Definition 7.39. Let $\mathcal{B}(\mathbf{R})$ denote the Borel sets of $\mathbf{R}$. A map $P : \mathcal{B}(\mathbf{R}) \rightarrow B(H)$ satisfying (i)-(iv) is called a projection-valued measure:

- (i) $P(\Omega)$ is an orthogonal projection for any $\Omega \in \mathcal{B}(\mathbf{R})$.
- (ii) $P(\Omega_1 \cap \Omega_2) = P(\Omega_1)P(\Omega_2)$
- (iii) $P(\cup_{n=1}^{\infty} \Omega_n) = \sum_{n=1}^{\infty} P(\Omega_n)$ if $\Omega_m \cap \Omega_n = \emptyset$ for all $m \neq n$. Here, the limit is to be understood as the strong limit. ¹

¹Recall that $T_n \rightarrow T$ strongly if $T_n \psi \rightarrow T \psi$ for all $\psi \in X$ for $T_n, T \in B(X)$.

(iv) $P(\emptyset) = 0$ and $P(\mathbf{R}) = I$.

Example 7.40 (Finite-Dimensional Case). Given a projective valued measure P , one can write an operator A as

$$A = \int_{\mathbf{R}} \lambda dP_{\lambda}.$$

We will in fact show that P is supported in $\sigma(A)$, so the integral is really an integral on $\sigma(A)$. In finite dimensional case, for a self-adjoint operator $A = A^*$, the spectrum $\sigma(A)$ is discrete, so the integral above becomes a sum, and for each $\lambda \in \sigma(A)$, we know P_{λ} is an orthogonal projection by (i). In fact, P_{λ} is the projection onto the eigenspace $P_{E_{\lambda}}$. Thus for a self-adjoint operator we have

$$A = \sum_{\lambda \in \sigma(A)} \lambda P_{E_{\lambda}}.$$

In physicist's dirac notation, this is simply

$$A = \sum_{\lambda \in \sigma(A)} \lambda |\lambda\rangle\langle\lambda|.$$

Construction 7.41 (The measure $\mu_{\psi, \phi}$). Recall that we defined a map (real functional calculus for A)

$$\Psi : \mathcal{C}(\sigma(A), \mathbf{R}) \rightarrow B(H) \quad f \mapsto f(A).$$

For any $\psi \in H$, we define a continuous linear form $\ell_{\psi} : \mathcal{C}(\sigma(A), \mathbf{R}) \rightarrow \mathbf{C}$ given by

$$\ell_{\psi}(f) = (\psi, f(A)\psi).$$

By definition, it is easy to see that $\|\ell_{\psi}\| = \|\psi\|^2$, and $\ell_{\psi}(\text{id}_{\mathcal{C}(\sigma(A), \mathbf{R})}) = \|\psi\|^2$. By Riesz representation theorem, for a compact set K , we have $\mathcal{C}(K)^* \cong m(K)$, where $m(K)$ is the set of Borel, regular, complex measures.² Therefore, for any $\psi \in H$, there exists a unique such measure μ_{ψ} , such that

$$(\psi, f(A)\psi) = \int_{\sigma(A)} f(\lambda) d\mu_{\psi}(\lambda) = \int_{\mathbf{R}} f d\mu_{\psi}.$$

Note that in the last inequality, we abused the notation and extended μ_{ψ} to the whole line by defining a new measure μ'_{ψ} with $\mu'_{\psi}(\Omega) = \mu_{\psi}(\Omega \cap \sigma(A))$ and still call μ'_{ψ} the old measure μ_{ψ} . Note that for any $f \geq 0$, by Theorem 7.34 part (4), we have $(\psi, f(A)\psi) \geq 0$. Hence by definition μ_{ψ} is a positive measure.

Now, we define, for $\Omega \in \mathcal{B}(\mathbf{R})$,

$$(\psi, P(\Omega)\psi) := \int \chi_{\Omega} d\mu_{\psi} = \mu_{\psi}(\Omega).$$

In the above expression, the domain of the integration is on the real line. From now on, any integral with no integration domain specified is assumed to be an integral over the entire real line.

²More specifically to our situation, the statement is that the isomorphism is given by ℓ_{ψ} to μ_{ψ} , where μ_{ψ} is the unique measure satisfying $\ell_{\psi}(f) = \int_K f d\mu_{\psi}$. The full statement is in Theorem 8.5, p.158 of [7].

Now we wish to define $(\phi, P(\Omega)\psi)$ for any $\phi, \psi \in H$. This can be done using the polarization identity:

$$\begin{aligned}\operatorname{Re}(\psi, \phi) &= \frac{1}{4}(\|\psi + \phi\|^2 - \|\psi - \phi\|^2) \\ \operatorname{Im}(\psi, \phi) &= \frac{1}{4}(\|\psi + i\phi\|^2 - \|\psi - i\phi\|^2).\end{aligned}$$

In other words, we define a complex measure $\mu_{\psi, \phi}$ to be

$$\operatorname{Re} \mu_{\psi, \phi} = \frac{1}{4}(\mu_{\psi+\phi} - \mu_{\psi-\phi}) \quad \operatorname{Im} \mu_{\psi, \phi} = \frac{1}{4}(\mu_{\psi+i\phi} - \mu_{\psi-i\phi}).$$

Then we may define

$$(\phi, P(\Omega)\psi) := \int \chi_\Omega d\mu_{\phi, \psi} = \mu_{\phi, \psi}(\Omega).$$

Viewing the equation above as a map $\phi \mapsto \mu_{\phi, \psi}(\Omega)$ for fixed $\psi \in H$ and $\Omega \in \mathcal{B}(\mathbf{R})$, by the Riesz representation theorem, this uniquely determines an element $P(\Omega)\psi \in H$. Thus the above definition determines $P(\Omega)$ completely. Now it remains to check that $P(\Omega)$ has the desired properties.

Remark 7.42. Given any projection-valued measure P , there exists a unique self-adjoint operator $A \in B(H)$ associated to it. Again we define a complex valued measure $\mu_{\psi, \phi}$ to be $\mu_{\psi, \phi}(\Omega) = (\psi, P(\Omega)\phi)$. Then we define

$$(\psi, A\phi) := \int_{\mathbf{R}} \lambda d\mu_{\psi, \phi}.$$

This completely determines A by the Riesz representation theorem, and one can show that $A = A^*$.

Theorem 7.43. *For all $A \in B(H)$ self-adjoint, there exists a unique projection-valued measure P^A on $\mathcal{B}(\mathbf{R})$ that is supported on $\sigma(A)$.*

Proof. If f is bounded, Borel function on $\sigma(A)$, we define $f(A)$ to be the unique operator in $B(H)$ such that

$$(\psi, f(A)\phi) := \int_{\sigma(A)} f(\lambda) d\mu_{\psi, \phi}(\lambda).$$

Now we define a new measure $\mu'_{\psi, \phi}$ such that $\mu'_{\psi, \phi}(\Omega) = \mu_{\psi, \phi}(\Omega \cap \sigma(A))$. Thus this is an extension of $\mu_{\psi, \phi}$ that is supported on $\sigma(A)$. And we identify this new measure with $\mu_{\psi, \phi}$. Therefore, we may write

$$(\psi, f(A)\phi) := \int_{\mathbf{R}} f d\mu_{\psi, \phi}$$

for any continuous complex valued function defined on $\mathbf{R}$. Define, for all $\Omega \in \mathcal{B}(\mathbf{R})$, the operator $P(\Omega)$ such that

$$(\psi, P(\Omega)\phi) = \int_{\mathbf{R}} \chi_\Omega d\mu_{\psi, \phi}.$$

Now we verify that $P : \mathcal{B}(\mathbf{R}) \rightarrow B(H)$ constructed this way satisfies (i)-(iv).

(i) Note that for any $f, g \in \mathcal{C}(\sigma(A), \mathbf{C})$, we have

$$\begin{aligned}(\psi, f(A)g(A)\phi) &= \int fg d\mu_{\psi, \phi} = \int f d\mu_{\psi, g(A)\phi} \\(f(A)^*\psi, g(A)\phi) &= \int g d\mu_{f(A)^*\psi, \phi} = \int d\mu_{f(A)^*\psi, g(A)\phi}.\end{aligned}$$

Therefore, by uniqueness of Riesz representation theorem, we have that

$$d\mu_{f(A)\psi, g(A)\psi} = f^*(\lambda)g(\lambda)d\mu_{\psi, \phi}.$$

That is, $d\mu$ is antilinear with respect to the first slot, and linear with respect to the second slot. Now, recall that

$$(\psi, P(\Omega)\phi) = \int \chi_\Omega d\mu_{\psi, \phi}.$$

Hence

$$(\psi, P(\Omega)^2\phi) = \int \chi_\Omega^* \chi_\Omega d\mu_{\psi, \phi} = \int \chi_\Omega^2 d\mu_{\psi, \phi} = \int \chi_\Omega d\mu_{\psi, \phi} = (\psi, P(\Omega)\phi).$$

Since ψ, ϕ are arbitrary, this shows that $P(\Omega)^2 = P(\Omega)$. Hence $P(\Omega)$ is a projection. It is also clear that

$$(\psi, P(\Omega)^*\phi) = (P(\Omega)\psi, \phi) = \int d\mu_{P(\Omega)\psi, \phi} = \int \chi_\Omega^* d\mu_{\psi, \phi} = \int \chi_\Omega d\mu_{\psi, \phi} = (\psi, P(\Omega)\phi).$$

Hence $P(\Omega) = P(\Omega)^*$. This in fact implies that $P(\Omega)$ is orthogonal. For details, see Proposition A.57 on p.541 of [7].

(iii) Let $\Omega = \cup_n \Omega_n$ with $\Omega_m \cap \Omega_n = \emptyset$ for $m \neq n$. Note that for all bounded complex-valued continuous function f , we have

$$\|f(A)\psi\|^2 = (f(A)\psi, f(A)\psi) = (\psi, f(A)^*f(A)\psi) = \int |f|^2 d\mu_{\psi, \psi}.$$

Thus for any $\psi \in H$ we have

$$\begin{aligned}\|P(\Omega)\psi - \sum_{j=1}^n P(\Omega_j)\psi\|^2 &= \int |\chi_\Omega - \sum_{j=1}^n \chi_{\Omega_j}|^2 d\mu_\psi \\&= \int \chi_{\Omega \setminus \cup_{j=1}^n \Omega_j} d\mu_\psi \\&= \int_{\Omega \setminus \cup_{j=1}^n \Omega_j} d\mu_\psi \\&= \mu_\psi(\Omega \setminus \cup_{j=1}^n \Omega_j) \\&= \mu_\psi(\Omega) - \sum_{j=1}^n \mu_\psi(\Omega_j)\end{aligned}$$

which converges to zero. Here we identify $d\mu_{\psi, \psi}$ with $d\mu_\psi$. This proves (iii).

(iv) Since $(\psi, P(\emptyset)\psi) = \int \chi_\emptyset d\mu_\psi = 0$, and $(\psi, P(\mathbf{R})\psi) = \int_{\mathbf{R}} d\mu_\psi = (\psi, \psi)$ for all $\psi \in H$ by the first equation in the proof of (iii). This shows that $P(\emptyset) = 0$ and $P(\mathbf{R}) = I$.

(ii) $P(\Omega_1 \cap \Omega_2) = P(\Omega_1)P(\Omega_2)$ is clear since $\chi_{\Omega_1}\chi_{\Omega_2} = \chi_{\Omega_1 \cap \Omega_2}$. To be more specific, recall that

$$(\psi, P(\Omega)\phi) = \int \chi_\Omega d\mu_{\psi, \phi}.$$

But since by (iv) we already know $P(\mathbf{R}) = I$, hence

$$(\psi, P(\Omega)\phi) = (\psi, P(\mathbf{R})P(\Omega)\phi) = \int \chi_{\mathbf{R}} d\mu_{\psi, P(\Omega)\phi} = \int d\mu_{\psi, P(\Omega)\phi}.$$

Therefore, by the uniqueness part of Riesz representation theorem as we did in (i), one sees that

$$\chi_\Omega d\mu_{\psi, \phi} = d\mu_{\psi, P(\Omega)\phi}.$$

Then it follows that for any $\psi, \phi \in H$, we have

$$\begin{aligned} (\psi, P(\Omega_1 \cap \Omega_2)\phi) &= \int \chi_{\Omega_1 \cap \Omega_2} d\mu_{\psi, \phi} \\ &= \int \chi_{\Omega_1} \chi_{\Omega_2} d\mu_{\psi, \phi} \\ &= \int \chi_{\Omega_1} d\mu_{\psi, P(\Omega_2)\phi} \\ &= (\psi, P(\Omega_1)P(\Omega_2)\phi). \end{aligned}$$

Hence it follows that $P(\Omega_1 \cap \Omega_2) = P(\Omega_1)P(\Omega_2)$.

This concludes the proof. □

We have showned that given any self-adjoint operator $A \in B(H)$ self-adjoint, there exists a unique projection-valued measure P such that (using formal notation):

$$(\psi, f(A)\phi) = \int_{\sigma(A)} f(\lambda) d(\psi, P\phi).$$

Now we prove the converse:

Theorem 7.44. *Given P a projection-valued measure, there exists a unique $A \in B(H)$ self-adjoint, such that*

$$A = \int_{\mathbf{R}} \lambda dP.$$

Proof. Define, for any simple function $f = \sum_{j=1}^n \alpha_j \chi_{f^{-1}(\alpha_j)}$ the integral

$$\int_{\mathbf{R}} f dP = \sum_{j=1}^n \alpha_j P(f^{-1}(\alpha_j)).$$

Note: We are defining $\int_{\mathbf{R}} \chi_{\Omega} dP = P(\Omega)$ and extending linearly. Before defining $\int f dP$ for general f , we first make a quick observation: For any $\psi \in H$, we have

$$\begin{aligned} \left\| \int_{\mathbf{R}} f dP\psi \right\|^2 &= \left(\int_{\mathbf{R}} f dP\psi, \int_{\mathbf{R}} f dP\psi \right) = \sum_{j,k} \alpha_j^* \alpha_k (P(f^{-1}(\alpha_j))\psi, P(f^{-1}(\alpha_k))\psi) \\ &= \sum_{j,k} \alpha_j^* \alpha_k (\psi, P(f^{-1}(\alpha_j))P(f^{-1}(\alpha_k))\psi) \\ &= \sum_{j,k} \alpha_j^* \alpha_k (\psi, P(f^{-1}(\alpha_j) \cap f^{-1}(\alpha_k))\psi) \\ &= \sum_j |\alpha_j|^2 (\psi, P(f^{-1}(\alpha_j))\psi) \\ &= \int |f|^2 d\mu_{\psi} \end{aligned}$$

where the second to last step is because $P(f^{-1}(\alpha_j) \cap f^{-1}(\alpha_k))$ is zero if $j \neq k$, for $P(\emptyset) = 0$. Thus we know that

$$\left\| \int f dP\psi \right\|^2 = \int |f|^2 d\mu_{\psi} \leq \|f\|_{\infty}^2 \mu_{\psi}(\mathbf{R}) = \|f\|_{\infty}^2 \|\psi\|^2.$$

Also, if $f = g$ a.e. μ_{ψ} for some ψ , then $\int f dP\psi = \int g dP\psi$. Now since, every bounded measurable function f is the uniform limit of simple functions $(\varphi_n)_n$, then by our previous calculation, we have

$$\left\| \int \varphi_n dP\psi - \int \varphi_m dP\psi \right\|^2 = \int |\varphi_n - \varphi_m|^2 d\mu_{\psi} \rightarrow 0$$

we can define

$$\int f dP\psi = \lim_{n \rightarrow \infty} \int \varphi_n dP\psi$$

and it is well-defined. Since $P((-a, a)) = I$ for some $a > 0$, it follows that

$$A := \int \lambda dP$$

is a bounded operator. Since $(-a, a) \subset \mathbf{R}$, it is easy to see that $A^* = A$, and by construction,

$$\int f(\lambda) d\mu_{\psi}^P = (\psi, f(A)\psi) = \int f(\lambda) d\mu_{\psi}^A$$

for every f continuous. Thus the projective-valued measure $\mu_{\psi}^P = \mu_{\psi}^A$, in other words, the spectral valued measure P is exactly the spectral valued measure μ_{ψ}^A associated to A given by our previous construction. $\square$

The two theorems above shows that there is a bijection between bounded self-adjoint operators in H and projection-valued measures. Now we study some applications of this consequence:

Lemma 7.45. *Assume $A \in B(H)$, $A = A^*$. Then $\lambda \in \sigma(A)$ iff $P^A(B(\lambda, \epsilon)) \neq 0$ for all $\epsilon > 0$.*

Proof. Assume $\lambda \in \sigma(A)$, and there exist $\epsilon > 0$ such that $PB(\lambda, \epsilon) = 0$. Then consider the function

$$f(x) = \frac{1}{x - \lambda} \chi_{B(\lambda, \epsilon)^c}(x).$$

Then f is a bounded Borel function on $\mathbf{R}$. Thus

$$(x - \lambda)f(x) = \chi_{B(\lambda, \epsilon)^c}$$

implies

$$(A - \lambda)f(A) = \chi_{B(\lambda, \epsilon)^c}(A) = f(A)(A - \lambda).$$

But

$$\chi_{B(\lambda, \epsilon)^c}(A) = \int \chi_{B(\lambda, \epsilon)^c} dP = P(B(\lambda, \epsilon)^c) = P(\mathbf{R}) = I$$

since by assumption $PB(\lambda, \epsilon) = 0$. Therefore, we have

$$(A - \lambda)f(A) = I = f(A)(A - \lambda).$$

Moreover,

$$\|f(A)\| \leq \|f\|_\infty = \sup_{|x - \lambda| \geq \epsilon} \frac{1}{|x - \lambda|} \leq \frac{1}{\epsilon}.$$

Hence $f(A) = (A - \lambda)^{-1}$ is a bounded operator, so $\lambda \in \rho(A)$, which contradicts our assumption that $\lambda \in \sigma(A)$.

Conversely, suppose $PB(\lambda, \epsilon) \neq 0$ for any $\epsilon > 0$. If $\lambda \in \rho(A)$, then since $\rho(A)$ is open, there exists $\delta > 0$ such that $(\lambda - \delta, \lambda + \delta) \cap \sigma(A) = \emptyset$. But then $P(\lambda - \delta, \lambda + \delta) = 0$ since P is supported on $\sigma(A)$, which is again a contradiction. $\square$

Definition 7.46 (Essential and Discrete Spectrum). We define

$$\sigma_e(A) = \{\lambda \in \sigma(A) : \dim \operatorname{im} P(B(\lambda, \epsilon)) = \infty \text{ for all } \epsilon > 0\}.$$

and

$$\sigma_d(A) = \sigma(A) \setminus \sigma_e(A) = \{\lambda \in \sigma(A) : \text{exists } \epsilon > 0 \text{ s.t. } \dim \operatorname{im} P(B(\lambda, \epsilon)) < \infty\}$$

Lemma 7.47. $\sigma_e(A)$ is closed.

Proof. Let $\lambda \notin \sigma_e(A)$, then there exists $\epsilon > 0$ such that $PB(\lambda, \epsilon)$ is finite dimensional. If $\mu \in B(\lambda, \epsilon)$, then there exist $\delta > 0$ such that $B(\mu, \delta) \subset B(\lambda, \epsilon)$, so $P(B(\mu, \delta))H \leq P(B(\lambda, \epsilon))H$. Thus $P(B(\mu, \delta))$ is finite-dimensional. Thus $\mu \notin \sigma_e(A)$. Hence $B(\lambda, \epsilon) \subset \sigma_e(A)^c$. Hence $\sigma_e(A)$ is closed. $\square$

Remark 7.48. $\sigma_d(A)$ is not necessarily closed, an example is when $A = -\Delta - 1/|\mathbf{x}|$ (Hydrogen atom), where $\sigma_d(A) = (\lambda_n)_{n=1}^\infty$, with λ_n increasingly converging to zero, and $\sigma_e(A) = [0, +\infty)$. But clearly $0 \notin \sigma_d(A)$.

Theorem 7.49. $\lambda \in \sigma_d(A)$ iff both the following conditions are satisfied:

(i) $\lambda \in \sigma(A)$ is isolated: i.e., there exist $\epsilon > 0$ such that $B(\lambda, \epsilon) \cap \sigma(A) = \{\lambda\}$.

(ii) λ is an eigenvalue of finite multiplicity, i.e., $\dim \ker(A - \lambda) < \infty$.

Proof. Assume first $\lambda \in \sigma_d(A)$.

Then exists $\epsilon > 0$ such that $PB(\lambda, \epsilon) \neq 0$ is finite dimensional. If λ is not isolated, then there exists a sequence $(\lambda_n) \subset \sigma(A)$ converging to λ , and $\lambda_n \neq \lambda_m$ for $n \neq m$, and for all n we have $\lambda_n \neq \lambda$. So there exist $\epsilon_n \rightarrow 0$ such that $B(\lambda_n, \epsilon_n) \subset B(\lambda, \epsilon)$ for all n , and $B(\lambda_n, \epsilon_n) \cap B(\lambda_m, \epsilon_m) = \emptyset$. Therefore $P(B(\lambda_n, \epsilon_n)) \neq 0$ for any n . Thus there exists $\psi_n \neq 0$, such that $\psi_n \in \text{im}(P(B(\lambda_n, \epsilon_n)))$. But since these balls are disjoint, we know (ψ_n) are orthogonal. WLOG we may assume $\|\psi_n\| = 1$ for all n . Hence $P(B(\lambda, \epsilon))$ contains all ψ_n , which means it cannot have a finite-dimensional image. This is a contradiction. Hence λ has to be isolated. This proves (i).

To show λ is an eigenvalue, there exist $\epsilon > 0$ such that $B(\lambda, \epsilon) \cap \sigma(A) = \{\lambda\}$. Let $\psi \in P(B(\lambda, \epsilon))H$ nonzero. For every $0 < \delta < \epsilon$, we have $\psi = P(B(\lambda, \delta))\psi$. Then

$$\|(A - \lambda)\psi\|^2 = \|(A - \lambda)P(B(\lambda, \delta))\psi\|^2 = \int_{B(\lambda, \delta)} |x - \lambda|^2 d\mu_\psi \rightarrow |\lambda - \lambda|^2 \mu_\psi(\{\lambda\}) = 0$$

as $\delta \rightarrow 0$ by the dominated convergence theorem. This shows that λ is an eigenvalue of A with ψ a corresponding eigenvector.

The prove that $\dim(A - \lambda) = P(B(\lambda, \epsilon))H$ is finite dimensional is left as an exercise.

Conversely, if $\lambda \in \sigma(A)$ is an isolated eigenvalue with finite multiplicity, then there exists $\epsilon > 0$ such that $B(\lambda, \epsilon) \cap \sigma(A) = \{\lambda\}$. And if $\psi \in \ker(A - \lambda) = P(B(\lambda, \epsilon))H$ (check this equality!) is finite-dimensional, then $\lambda \in \sigma_d(A)$. $\square$

Corollary 7.50. $\lambda \in \sigma_e(A)$ iff λ is an eigenvalue of infinite multiplicity or λ is an accumulation point of $\sigma(A)$.

Theorem 7.51 (Weyl's Criterion). Let $A \in B(H)$, $A = A^*$. Then $\lambda \in \sigma(A)$ iff there exists $(\psi_n) \subset D(A)$ such that $\|\psi_n\| = 1$ for all n , and $\|(A - \lambda)\psi_n\| \rightarrow 0$.

Furthermore, $\lambda \in \sigma_e(A)$ iff the condition above is satisfied, and furthermore we may choose (ψ_n) to be orthonormal.

Proof. The first part of the statement is proven in Theorem 7.35.

Suppose $\lambda \in \sigma_e(A)$. Then for all $\epsilon > 0$, we have $P(B(\lambda, \epsilon))$ infinite dimensional. Pick $\epsilon_n \rightarrow 0$, such that $\psi_n \in P(B(\lambda, \epsilon_n))$ and $\psi_n \perp \psi_m$ for $n \neq m$. Then

$$\|(A - \lambda)\psi_n\|^2 = \int_{|x - \lambda| < \epsilon_n} |x - \lambda|^2 d\mu_\psi \leq \epsilon_n^2 \|\psi_n\|^2 = \epsilon_n$$

which tends to zero. Conversely, if $P(B(\lambda, \epsilon))$ were finite dimensional for some ϵ , then there exists n_0 such that $\psi_n \perp \text{im } P(B(\lambda, \epsilon))$ for all $n \geq n_0$. Then

$$\|(A - \lambda)\psi_n\|^2 = \|(A - \lambda)(I - P(B(\lambda, \epsilon)))\psi_n\|^2 \geq \epsilon^2 \|\psi_n\|^2 = \epsilon^2.$$

So $\|(A - \lambda)\psi_n\|^2$ cannot tend to zero. This concludes the proof. $\square$

Theorem 7.52 (Stone's Formula). Let $A \in B(H)$, $A = A^*$. Then

$$\lim_{\epsilon \rightarrow 0} \frac{1}{2\pi i} \int_a^b (A - \lambda - i\epsilon)^{-1} - (A - \lambda + i\epsilon)^{-1} d\lambda = \frac{1}{2}P([a, b]) + \frac{1}{2}P(a, b)$$

where the limit is to be understood as strong limit.

Proof. Take $\psi \in H$. Then

$$\begin{aligned} & \frac{1}{2\pi i} \int_a^b [(A - \lambda - i\epsilon)^{-1} - (A - \lambda + i\epsilon)^{-1} d\lambda] \psi \\ &= \frac{1}{2\pi i} \int_a^b \int_{\mathbf{R}} \left(\frac{1}{x - \lambda - i\epsilon} - \frac{1}{x - \lambda + i\epsilon} \right) dP\psi d\lambda \\ &= \frac{1}{2\pi i} \int_{\mathbf{R}} \int_a^b \left(\frac{1}{x - \lambda - i\epsilon} - \frac{1}{x - \lambda + i\epsilon} \right) d\lambda dP\psi \end{aligned}$$

where the last step is by Fubini's theorem since the integrand is square integrable. Our next task is to calculate the integral inside the last line above.

Note that

$$\begin{aligned} \frac{1}{2\pi i} \int_a^b \left(\frac{1}{x - \lambda - i\epsilon} - \frac{1}{x - \lambda + i\epsilon} \right) d\lambda &= \frac{1}{\pi} \int_a^b \frac{\epsilon}{(x - \lambda)^2 + \epsilon^2} d\lambda \\ &= \frac{1}{\pi} \arctan \left(\frac{\lambda - x}{\epsilon} \right) \Big|_a^b \\ &= \frac{1}{\pi} \left[\arctan \left(\frac{b - x}{\epsilon} \right) - \arctan \left(\frac{a - x}{\epsilon} \right) \right] \end{aligned}$$

which goes to

$$\frac{1}{2} (\chi_{(a,b)} + \chi_{[a,b]}) = \begin{cases} 0 & x \notin [a, b] \\ 1 & x \in (a, b) \\ \frac{1}{2} & x = a, b \end{cases}$$

as $\epsilon \rightarrow 0$. Now note that

(*) If $f_\epsilon \rightarrow f$ pointwise for $\epsilon > 0$, and $|f_\epsilon| \leq M$ for some constant M , then we have

$$\| \int f_\epsilon dP\psi - \int f dP\psi \|^2 = \int |f_\epsilon - f|^2 d\mu_\psi \rightarrow 0$$

by the dominated convergence theorem.

Using (*), and put $f_\epsilon(x) = \frac{1}{2\pi i} \int_a^b \left(\frac{1}{x - \lambda - i\epsilon} - \frac{1}{x - \lambda + i\epsilon} \right) d\lambda$ and $f(x) = \frac{1}{2} (\chi_{(a,b)} + \chi_{[a,b]})$. Then we have that

$$\lim_{\epsilon \rightarrow 0} \frac{1}{2\pi i} \int_a^b [(A - \lambda - i\epsilon)^{-1} - (A - \lambda + i\epsilon)^{-1} d\lambda] \psi = \int \frac{1}{2} (\chi_{(a,b)} + \chi_{[a,b]}) dP\psi$$

which is equal to

$$\frac{1}{2} (P[a, b] + P(a, b)) \psi.$$

This concludes the proof. $\square$

Definition 7.53. We say $\psi \in H$ is cyclic for A if $\text{span} \{A^n \psi : n = 0, 1, 2, \dots\} = H$.

Theorem 7.54 (Multiplicative Form of the Spectral Theorem). *Let $A \in B(H)$, $A = A^*$, then there exists a sequence of measures $(\mu_n)_{n=1}^\infty$, and a unitary operator*

$$U : H \rightarrow \bigoplus_{n=1}^{\infty} L^2(\mathbf{R}, d\mu_n)$$

such that

$$(UAU^{-1}\psi)_n = \lambda\psi_n$$

and

$$(Uf(A)U^{-1}\psi)_n = f(\lambda)\psi_n.$$

Proof. Define $U : \text{span}\{A^n\psi\} \rightarrow L^2(\mathbf{R}, d\mu_n)$ as follows: For a given polynomial p_k , define $U(p_k(A)\psi) = p_k$. Note that U is an isometry since

$$\|Up_k(A)\psi\|^2 = \int |p_k|^2 d\mu_\psi = \|p_k(A)\psi\|^2.$$

Therefore U is an isometry. This is well-defined, since if $p(A)\psi = q(A)\psi$, then

$$0 = \|p(A)\psi - q(A)\psi\|^2 = \int |p - q|^2 d\mu_\psi.$$

This shows that $p = q$ a.e. $d\mu_\psi$. Now extend U to $\text{Cl}[\text{span}\{A^n\psi\}]$. Then by construction

$$(UAU^{-1}p)(\lambda) = UA(p(A)\psi) = U(Ap(A)\psi) = \lambda p(\lambda)$$

and

$$UAU^{-1} = \lambda I.$$

Now we claim that given H separable, we have $H = \bigoplus_{n=1}^{\infty} H_n$, where H_n are cyclic subspaces, i.e., there exists ψ_n such that $H_n = \text{Cl}[\text{span}\{A^k\psi_n\}]$: First we take any $\psi \in H$ nonzero, and set $H_1 = \text{span}\{A^k\psi\} \leq H$. Then clearly $AH_1 \leq H_1$. Since A is self-adjoint, we have $A : H_1^\perp \rightarrow H_1^\perp$. Then since $H = H_1 \oplus H_1^\perp$, and repeat this construction on $H_1^\perp$. Then the decomposition of H into H_n 's follows.

For every n , we define $U_n : H \rightarrow L^2(\mathbf{R}, d\mu_n)$ as above, with $\mu_n = \mu_{\psi_n}$ where each ψ_n is chosen as in the last paragraph.. Then take $U = \bigoplus_{n=1}^{\infty} U_n : H \rightarrow \bigoplus_{n=1}^{\infty} L^2(\mathbf{R}, d\mu_n)$. $\square$

Index

- 4-velocity, 42
- action, 11
- antisymmetric, 39
- antisymmetrize, 40
- atlas, 55
- Bianchi identity, 75
- bosonic, 112
- character, 113
- chart, 55
- Christoffel symbol, 72
- Clebsch-Gordan Decomposition Theorem, 113
- commutator, 59
- configuration space, 8
- connection coefficients, 65
- continuous, 54
- contract, 39
- cotangent space, 59
- covariant derivative, 65
- decomposable, 112
- diffeomorphic, 55
- diffeomorphism, 55
- differential forms, 62
- directional covariant derivative, 69
- dust, 43
- Einstein Equivalence Principle, 52
- Einstein tensor, 76
- Einstein's Equivalence Principle, 49
- Einstein's field equation, 84
- electric field, 13
- energy-momentum tensor, 46
- Euler-Lagrange equation, 11
- event, 25
- exterior derivative, 63
- fermionic, 112
- flux, 46
- force, 12
- Galilean Transformation, 21
- Gallilean group, 12
- gauge transformation, 86
- geodesic, 70
- geodesic deviation equation, 78
- geodesic equation, 70
- gravitational mass, 49
- Hamiltonian, 9, 10
- Hamiltonian flow, 9
- Heisenberg's uncertainty principle, 140
- homeomorphic, 54
- homeomorphism, 54
- indecomposable, 112
- inertial frame of reference, 20
- inertial mass, 49
- inverse metric, 39
- irreducible, 91
- isotropic, 10
- kinetic energy, 12
- Lagrangian, 11
- Lagrangian submanifold, 10
- length contraction, 34
- Levi-Civita connection, 68
- Levi-Civita tensor, 62
- Lie algebra, 94
- Lie bracket, 59
- Lie group, 93
- light cone, 27
- Liouville volume form, 10
- Liouville's Theorem, 10
- local Lorentz frame, 52
- Lorentz group, 32
- Lorentz transformations, 31
- magnetic field, 13
- metric tensor, 61

Minkowski spacetime, 48
 momentum map, 10

 number density, 43

 observables, 9
 open sets, 54
 orthogonal, 39

 parallel transport, 69
 partial derivative operator, 57
 Pauli Exclusion Principle, 112
 perfect fluid, 46
 phase space, 8
 Poisson manifold, 10
 Poisson's equation, 16
 potential, 12
 principle of least action, 11
 proper time, 28

 raise and lower indices, 39
 rapidity, 32
 representation, 91
 rest energy, 43
 Ricci scalar, 76
 Ricci tensor, 75
 Riemann curvature tensor, 72

 simple harmonic oscillator, 136
 smooth, 55
 smooth manifold, 55
 smoothly compatible, 55
 spacetime, 25
 spherical harmonics, 109
 symmetric, 39
 symmetrize, 40
 symplectic, 10
 symplectic manifold, 8
 symplectomorphism, 10

 tangent space, 35
 time dilation, 34
 time independent Schrödinger equation, 133
 topological manifold, 54
 topological space, 54
 topology, 54
 torsion tensor, 67
 torsion-free, 67
 trace, 40
 trace-reversed metric perturbation, 87

 unitary, 91
 universal cover, 103

 variation, 11
 vector field, 59

 Weak Equivalence Principle, 49, 52
 wedge product, 62
 world line, 25

Bibliography

- [1] John Lee, *Introduction to Smooth Manifolds*, Springer, 2012.
- [2] John Lee, *Introduction to Riemannian Manifolds*, Springer, 2019.
- [3] Peter Woit, *Quantum Theory, Groups and Representation*, Springer, 2017.
- [4] Brian Hall, *Lie Groups, Lie Algebras, and Representations: An Elementary Introduction*, Springer, 2015.
- [5] Brian Hall, *Quantum Theory for Mathematicians*, Springer, 2013.
- [6] David Griffiths, *Introduction to Quantum Mechanics*, 3ed, Cambridge University Press, 2018.
- [7] Michael Reed and Barry Simon, *Methods of Modern Mathematical Physics I: Functional Analysis*, Academic Press, 1980.
- [8] Rafael de la Madrid Modino, [*Quantum Mechanics in Rigged Hilbert Space Language*](#), 2001.
- [9] V.I. Arnold, *Mathematical Methods of Classical Mechanics*, Springer, 1978.
- [10] Sean M. Carroll, *Spacetime and Geometry: An Introduction to General Relativity*, Cambridge University Press, 2019.
- [11] James Hartle, *Gravity: An Introduction to Einstein's General Relativity*, Cambridge University Press, 2002.
- [12] Theodore Frankel, *The Geometry of Physics: An Introduction*, Cambridge University Press, 2011.
- [13] Maxim Jeffs, *Classical mechanics and Symplectic Geometry*, 2022.
- [14] Matthew D. Schwartz *Quantum Field Theory and The Standard Model*, Cambridge University Press, 2013.
- [15] Anthony Zee, *Group Theory in a Nutshell for Physicists*, Princeton University Press, 2016.
- [16] Tianyi Wang, *Topics on Smooth Manifolds and Differential Topology*.
- [17] Mikio Nakahara, *Geometry, Topology, and Physics*, CRC Press, 2003.