



European IT Certification Curriculum Self-Learning Preparatory Materials

EITC/QI/QIF
Quantum Information Fundamentals



This document constitutes European IT Certification curriculum self-learning preparatory material for the EITC/QI/QIF Quantum Information Fundamentals programme.

This self-learning preparatory material covers requirements of the corresponding EITC certification programme examination. It is intended to facilitate certification programme's participant learning and preparation towards the EITC/QI/QIF Quantum Information Fundamentals programme examination. The knowledge contained within the material is sufficient to pass the corresponding EITC certification examination in regard to relevant curriculum parts. The document specifies the knowledge and skills that participants of the EITC/QI/QIF Quantum Information Fundamentals certification programme should have in order to attain the corresponding EITC certificate.

Disclaimer

This document has been automatically generated and published based on the most recent updates of the EITC/QI/QIF Quantum Information Fundamentals certification programme curriculum as published on its relevant webpage, accessible at:

<https://eitca.org/certification/eitc-qi-qif-quantum-information-fundamentals/>

As such, despite every effort to make it complete and corresponding with the current EITC curriculum it may contain inaccuracies and incomplete sections, subject to ongoing updates and corrections directly on the EITC webpage. No warranty is given by EITCI as a publisher in regard to completeness of the information contained within the document and neither shall EITCI be responsible or liable for any errors, omissions, inaccuracies, losses or damages whatsoever arising by virtue of such information or any instructions or advice contained within this publication. Changes in the document may be made by EITCI at its own discretion and at any time without notice, to maintain relevance of the self-learning material with the most current EITC curriculum. The self-learning preparatory material is provided by EITCI free of charge and does not constitute the paid certification service, the costs of which cover examination, certification and verification procedures, as well as related infrastructures.

TABLE OF CONTENTS

Getting started	5
Overview	5
Introduction to Quantum Mechanics	7
Introduction to double slit experiment	7
Double slit experiment with waves and bullets	9
Conclusions from the double slit experiment	11
Introduction to Quantum Information	13
Qubits	13
Geometric representation	15
Photon polarization	17
Uncertainty principle	19
Quantum Entanglement	21
K-level system and bra-ket notation	21
Systems of two qubits	23
Entanglement	24
EPR Paradox	26
Bell and EPR	28
Rotational invariance of Bell state	30
CHSH inequality	32
Bell and local realism	34
Quantum Information processing	36
Time evolution of a quantum system	36
Unitary transforms	37
Single qubit gates	39
Two qubit gates	41
Quantum Information properties	44
No-cloning theorem	44
Bell state circuit	45
Quantum Teleportation	47
Quantum Teleportation using CNOT	48
Quantum Measurement	50
Introduction to Quantum Computation	51
N-qubit systems	51
Universal family of gates	53
Reversible computation	55
Conclusions from reversible computation	57
Quantum Algorithms	59
Fourier sampling	59
Applying Fourier sampling	61
Simon's Algorithm	63
Conclusions from Simon's Algorithm	65
Simon's algorithm in terms of the double slit experiment	67
Extended Church-Turing Thesis	69
Quantum Fourier Transform	71
QFT overview	71
N-th roots of unity	73
Discrete Fourier Transform	74
N-th Dimensional Quantum Fourier Transform	76
Properties of Quantum Fourier Transform	78
Shor's Quantum Factoring Algorithm	80
Period finding	80
Shor's Factoring Algorithm	82
QFT circuit	85
Grover's Quantum Search Algorithm	87
Needle in a haystack	87
Grover's Algorithm	89
Implementing Grover's Algorithm	91

Observables and Schrodinger's equation	93
Introduction to observables	93
Observables properties	95
Schrodinger's equation	96
Introduction to implementing qubits	98
Continuous quantum states	98
Schrodinger's equation for a 1D free particle	100
Particle in a box	102
Implementing qubits	104
Introduction to Quantum Complexity Theory	105
Limits of quantum computers	105
Adiabatic quantum computation	107
BQP	109
Introduction to spin	111
Spin as a qubit	111
Bloch Sphere	112
Stern-Gerlach experiment	114
Pauli spin matrices	116
Manipulating spin	118
Larmor precession	118
Spin resonance	120
Classical control	122
Summary	124
Summary	124

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: GETTING STARTED****TOPIC: OVERVIEW**

Welcome to this course on quantum mechanics and quantum computation. In this overview, we will discuss what you can expect to learn from this course and how it is organized.

Quantum computation is based on the remarkable discovery that quantum systems are exponentially powerful. The goal of quantum computation is to harness this exponential power to solve computational problems. To understand the power of quantum systems, let's consider a small quantum system of a few hundred particles. If we could harness all the computational power inherent in this system, each cycle of a resulting quantum computer could carry out 2 to the power of 500 steps.

2 to the power of 500 is an impossibly large number, much larger than estimates for the total number of particles in the universe or the age of the universe in femtoseconds. This means that even if we could use the entire resources of the classical universe, we could not match the computational power of a quantum computer.

However, harnessing this power is challenging. In this course, we will discuss the challenges of quantum computation. First, we need to pick the right computational problems that can be sped up by quantum computation. One famous example is the factoring problem, where we want to factorize a given number into its prime power factors.

Designing a quantum algorithm is also a tricky task. Quantum algorithms look very different from classical algorithms and have different design principles. We will explore concepts like the content Fourier transform and a new style of designing algorithms.

We will also discuss the limits of quantum computers, the problems that cannot be solved quickly even with quantum computers. Additionally, building a quantum computer is a difficult challenge that many scientists around the world are working on. We will briefly touch upon the types of systems and principles involved in designing quantum computers.

To study quantum computation, it is necessary to learn the basics of quantum mechanics. In this course, we will introduce quantum mechanics in terms of qubits, which are the simplest quantum systems. Describing quantum mechanics in terms of qubits simplifies the presentation and allows us to quickly delve into the most counterintuitive aspects of quantum mechanics.

Within three to four weeks, we will cover the basic notions of quantum computation, including quantum algorithms. We will also explore entanglement, one of the most mysterious aspects of quantum systems. Topics such as Bell inequalities and quantum teleportation will be discussed, allowing you to grapple with the counterintuitive aspects of quantum mechanics.

By the end of this course, you will have a solid understanding of quantum mechanics and quantum computation, as well as the challenges and possibilities they present.

Quantum Information - Quantum Information Fundamentals - Getting started - Overview

Quantum mechanics is a theory that exploits the most counterintuitive aspects of the mechanics. It is important to deeply understand these counterintuitive aspects in order to grasp quantum mechanics. Niels Bohr, the physicist who discovered the structure of the atom, emphasized the counterintuitive nature of quantum mechanics, stating that anyone who is not shocked by it has not truly understood it.

In this course, we will approach quantum mechanics by focusing on simple systems that illustrate its most counterintuitive aspects. This way of learning quantum mechanics can be beneficial, regardless of whether you are interested in quantum computation or not. For those who haven't studied quantum mechanics before, this approach might be the right way to start studying the subject. Later on, if you're interested, you can take a standard course in quantum mechanics to learn more advanced topics.

Even for those who have already studied quantum mechanics, this treatment can deepen your appreciation of

the subject. The required background for this course has been designed to be as broadly accessible as possible. The main prerequisite is a solid background in basic linear algebra. Additionally, you should be able to analyze the running time of simple algorithms, such as sorting or multiplying integers.

The course will adopt a Kanban approach to mathematical concepts and notation. Similar to the just-in-time manufacturing approach, the course will introduce new concepts as naively as possible, allowing you to build an intuition for them before delving into precise mathematical notation. This approach aims to prevent overwhelming you with mathematical notation while grappling with the challenging concepts.

It is important to bring your imagination and an ability to think critically to this course. Some of the concepts covered will be mind-bending, and your willingness to grapple with them will enhance your learning experience.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM MECHANICS****TOPIC: INTRODUCTION TO DOUBLE SLIT EXPERIMENT**

The double slit experiment is a fundamental experiment in quantum mechanics that illustrates the strange behavior of nature at the atomic level. It was initially used to demonstrate the wave nature of light, but later became important in understanding the behavior of electrons as well.

In the early 20th century, there was a lot of confusion about the nature of light. Newton believed that light was made up of particles, while evidence suggested that light actually traveled as waves. This confusion was resolved when Einstein discovered the photoelectric effect, which showed that light is transmitted in discrete packets called photons.

Similar confusion surrounded electrons, which were initially believed to be particles. However, evidence showed that they also behaved like waves in phenomena such as electron diffraction. This led to a growing confusion about whether atomic particles were wave-like or particle-like.

The laws of quantum mechanics, discovered in the mid-1920s, resolved this confusion. They showed that atomic particles are neither waves nor particles in the classical sense. Instead, they exhibit their own strange quantum mechanical behavior.

The double slit experiment is a way to describe and understand this quantum mechanical behavior. In this experiment, a source emits either light or electrons as discrete particles at a very low intensity. These particles then pass through a screen with two slits and are detected on a backstop screen at various points.

When only one slit is open, the probability of detecting the particle at a particular point on the backstop screen follows a certain curve. This curve represents the behavior we would expect if the particle went through that specific slit. When the other slit is open and both slits are available for the particle to pass through, the probability of detection at each point on the backstop screen is the sum of the probabilities from each individual slit.

However, what is observed in the double slit experiment is an interference pattern. This means that the probability of detection at certain points on the backstop screen is much higher than what would be expected from the sum of the probabilities of each individual slit.

This interference pattern is the mystery of the double slit experiment. It raises the question of how it is possible for the particle to be influenced by the presence or absence of the other slit. If the particle went through one slit, how could it matter whether the other slit was open or closed? This mystery highlights the counterintuitive behavior of atomic particles described by quantum mechanics.

Understanding the double slit experiment is crucial in grasping the fundamental concepts of quantum mechanics. It provides a simple context to develop intuition about the strange behavior of atomic particles. However, some individuals may find it challenging to relate to this experiment. For those with a computer science background, a very simple description of quantum bits (qubits) will be explored in the next lecture, which will be self-contained and easier to follow.

The double slit experiment is a fundamental experiment in quantum mechanics that demonstrates the wave-particle duality of atomic particles. The interference pattern observed in this experiment raises questions about the nature of particles and their behavior at the atomic level.

Quantum mechanics is a fundamental theory that describes the behavior of atomic particles. It poses a mystery: how can nature make these particles behave in such a peculiar way? This question has puzzled scientists for decades. To illustrate this mystery, we can turn to a quote from the renowned physicist Richard Feynman, who once said, "I can safely say that no one understands quantum mechanics."

Quantum mechanics introduces a new way of understanding the physical world at the atomic scale. It challenges our classical intuition and requires us to think in terms of probabilities and wave-like behavior. One of the most famous experiments that highlights the peculiarities of quantum mechanics is the double-slit

experiment.

In the double-slit experiment, a beam of particles, such as electrons or photons, is directed towards a barrier with two slits. Behind the barrier, a screen captures the pattern produced by the particles after passing through the slits. Surprisingly, when the particles are sent one by one, they create an interference pattern on the screen, as if they were behaving like waves.

This phenomenon contradicts our classical understanding of particles as localized objects. In classical physics, we would expect each particle to pass through one slit and create two separate patterns on the screen. However, in the quantum realm, particles can behave as both particles and waves simultaneously. This duality is a fundamental aspect of quantum mechanics.

The double-slit experiment challenges us to question the nature of reality. It suggests that particles do not have definite properties until they are observed. The act of measurement collapses the wave-like behavior into a specific outcome. This concept is known as wavefunction collapse.

The mystery of quantum mechanics lies in the fact that we cannot fully comprehend the underlying mechanisms that govern the behavior of atomic particles. Nature seems to operate in a way that defies our everyday intuition. Scientists continue to explore and study quantum mechanics, hoping to unlock its secrets and gain a deeper understanding of the fundamental nature of our universe.

Quantum mechanics is a perplexing theory that describes the behavior of atomic particles. The double-slit experiment exemplifies the peculiarities of quantum mechanics, showcasing the wave-particle duality and the mystery of wavefunction collapse. While the theory remains enigmatic, scientists strive to unravel its intricacies and shed light on the fundamental workings of our world.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM MECHANICS****TOPIC: DOUBLE SLIT EXPERIMENT WITH WAVES AND BULLETS**

The double slit experiment is a classic experiment in quantum mechanics that helps us understand the behavior of subatomic particles like electrons and photons. In this experiment, we compare and contrast the behavior of these particles with classical particles, which we model as bullets and waves.

Let's start by considering the behavior of bullets. Imagine a machine gun as our source, randomly firing bullets in different directions. We assume that bullets are indestructible and that our detector, a box of sand for example, always detects a whole number of bullets. If we place two detectors at different locations, we never see a bullet arriving simultaneously in both detectors due to the firing rate of the machine gun.

When we study the probability of arrival of the bullets as a function of position (X), we observe a curve ($P1(X)$) when only the first hole is open, another curve ($P2(X)$) when only the second hole is open, and the sum of these two curves when both holes are open. This behavior is what we would expect with bullets.

Now let's repeat the experiment with waves. Imagine a pond where waves are generated by an object vibrating at a constant rate. When these waves encounter a barrier with two slits, they start spreading out and the crests and troughs of the waves completely match up due to the equidistant source. At the backstop, we detect the energy of the wave by observing the oscillation of a cork in the water.

When we plot the intensity or energy (I) as a function of position (X), we observe a function ($I1(X)$) when only one slit is open, another function ($I2(X)$) when only the second slit is open, and an interference pattern ($I12(X)$) when both slits are open. This interference pattern is the same as what we observed with electrons and photons. However, in the case of waves, we have a clear explanation for why the intensity when both slits are open ($I12(X)$) is not equal to the sum of the intensities when each slit is open ($I1(X) + I2(X)$).

The reason for this is that the energy of the wave is proportional to the square of the height of the wave at a given position. When both slits are open, the height of the water at a position (X) is the sum of the heights due to the waves from the first slit and the waves from the second slit. However, the energies do not add up in the same way. This is why we observe the interference pattern.

To intuitively understand this interference pattern, let's consider the midpoint between the two slits. Here, when a crest arrives from the first slit, a corresponding crest arrives from the second slit, causing the water to move up by the sum of these two crests. Similarly, when a trough arrives from the first slit, a corresponding trough arrives from the second slit, causing the water to move down by the sum of these two troughs. This results in a particularly big wave and a particularly big trough at this midpoint.

The double slit experiment with waves and bullets demonstrates the wave-particle duality of subatomic particles. While bullets behave like classical particles, waves exhibit interference patterns that can only be explained through the wave nature of particles.

In the double slit experiment, we observe a phenomenon known as interference, which occurs when waves interact with each other. This experiment can be conducted with water waves, as well as with particles such as electrons and photons.

When water waves pass through two slits, they create an interference pattern on a screen. This pattern consists of alternating regions of constructive and destructive interference. Constructive interference occurs when the crests of two waves align, resulting in a higher amplitude and more energy. Destructive interference, on the other hand, happens when the crest of one wave aligns with the trough of another wave, leading to cancellation and a lower amplitude.

The interference pattern observed in the double slit experiment with water waves can be explained by considering the distance traveled by the waves from each slit to a particular point on the screen. If the two waves have traveled the same distance, they will be in phase and produce constructive interference. However, if one wave has traveled a longer distance than the other, they will be out of phase and produce destructive interference.

In the case of electrons and photons, which are transmitted as discrete packets of energy, the interference pattern is also observed. However, since these particles behave more like bullets, it is intriguing that they still exhibit interference. When one of the slits is closed, the probability of the particle arriving at a specific point on the detector is given by the probability amplitude. This probability amplitude can be positive or negative, similar to the height of water waves.

The total probability amplitude of the particle being detected at a certain point is the sum of the probability amplitudes for each slit. However, the probability of detection is not equal to the sum of the probabilities for each slit individually. This is analogous to the water wave case, where the probability of detecting the particle is the square of the total amplitude, which is not equal to the sum of the squares of the individual amplitudes.

Mathematically, this scenario is similar to the water wave case. However, the challenge lies in interpreting the positive and negative probability amplitudes. The question of what nature is doing behind the scenes to make this happen remains unanswered. Physicists have come to accept that certain questions about nature cannot be answered satisfactorily.

The double slit experiment with waves and particles demonstrates the phenomenon of interference. Whether it is water waves or particles like electrons and photons, interference patterns emerge when waves interact with each other. The mathematics behind these patterns is similar, but the interpretation of probability amplitudes for particles poses challenges that have yet to be fully understood.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM MECHANICS****TOPIC: CONCLUSIONS FROM THE DOUBLE SLIT EXPERIMENT**

In the previous material, we discussed the double slit experiment and observed the peculiar behavior of elementary particles. We saw that these particles exhibit both particle-like and wave-like characteristics. While they behave like discrete particles, we also observe interference patterns when both slits are open. This raises the question of how to make sense of the probability amplitudes associated with the particles, which can be complex numbers with positive, negative, or imaginary values.

To address this, we need to develop a new intuition about how particles behave in the quantum realm. We can start by examining a proposition that assumes particles have a definite trajectory. According to this proposition, if an electron arrives at point X, it must have either gone through slit 1 or slit 2. However, this proposition is proven false by the interference pattern observed when both slits are open.

To further investigate this proposition, we can design an experiment to determine which slit the electron passes through. By placing a source of light near each slit, we can detect whether the electron goes through slit 1 or slit 2. If we close one of the slits and count the number of electrons detected at point X, we can create separate curves, p_1 prime of X and p_2 prime of X, representing the number of electrons detected for each slit. Interestingly, these curves resemble the individual curves obtained when each slit is closed.

However, when both slits are open, the probability of an electron ending up at point X is not simply the sum of p_1 prime and p_2 prime. Instead, we observe a new curve, p_{1+2} prime of X, which represents the total number of electrons detected at point X. This curve does not match the interference pattern we expect. It seems that if we can determine which slit the electron passes through, the interference pattern disappears.

To further investigate this phenomenon, we can consider the effect of light on the electrons. When we place a source of light near the slits, the electrons are disturbed, altering their trajectory slightly. This disturbance smooths out the interference pattern, resulting in a smooth curve. To minimize the disturbance caused by light, we can reduce the intensity of the light source. However, we encounter a new complication.

Light itself is quantized, meaning it comes in distinct packets called photons. As we decrease the intensity of the light, there are times when no photons are emitted while the electron passes through the slits. Consequently, we may miss detecting some electrons. This leads to a situation where we sometimes detect electrons passing through one slit, sometimes through the other, and sometimes we don't detect them at all.

The double slit experiment reveals the wave-particle duality of elementary particles. When both slits are open, particles exhibit interference patterns, suggesting wave-like behavior. However, when we attempt to determine which slit the particle passes through, the interference pattern disappears. Additionally, the disturbance caused by light, which is quantized as photons, further complicates the experiment.

In the double slit experiment, we observe a phenomenon called interference, which is a key characteristic of quantum systems. When electrons pass through two slits, they create an interference pattern on a screen behind the slits. However, if we try to detect which slit each electron goes through, the interference pattern disappears.

This tells us something important about quantum systems. They are delicate and easily disturbed by measurement. When we try to observe or measure an electron, it disrupts the system, and the experiment is no longer the same. This is known as Heisenberg's uncertainty principle, which states that it is impossible to detect which slit an electron goes through without disturbing the interference pattern.

The uncertainty principle implies that there is no way to design an apparatus that can detect the path of an electron without affecting its behavior. The more accurately we try to determine the path, the more we disturb the electron and destroy the interference pattern. In other words, there is a trade-off between knowledge of the electron's path and the ability to observe interference.

Another important insight from the double slit experiment is that the behavior of quantum particles is inherently probabilistic. Even if we had complete knowledge of the initial conditions of the electron, we still cannot predict

which slit it will go through. If we could predict the path, we would not observe interference.

This randomness in measurement outcomes is a fundamental aspect of quantum mechanics. It is not due to a lack of knowledge but is inherent to the nature of quantum systems. In future lectures, we will explore these concepts more deeply and discuss them in the context of qubits, which are fundamental units of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM INFORMATION****TOPIC: QUBITS**

The basic unit of quantum information is called a qubit. To understand how a qubit works, let's consider a thought experiment using an electron. In an atom, the energy of an electron is quantized, meaning it can only have certain discrete energy levels. For example, it can be in the ground state or one of the excited states, depending on its energy level.

To represent a bit of information using an electron, we can ensure that its energy level is high enough to be in either the ground state or the first excited state, but not high enough for any higher energy state. We can encode the bit by assigning 0 to the ground state and 1 to the excited state.

However, in quantum mechanics, the electron doesn't definitively choose one state. Instead, it exists in a superposition of both states, with complex amplitudes. This means that the electron is partly in the ground state and partly in the excited state. We can express this superposition using complex amplitudes α and β , where α represents the amplitude of being in the ground state and β represents the amplitude of being in the excited state.

The complex amplitudes can take on different values, such as $1/\sqrt{2}$ and $-1/\sqrt{2}$. It's important to note that the state of the electron must be normalized, meaning the square of the magnitude of α plus the magnitude of β must equal 1.

When we measure the state of the electron, it quickly "collapses" into either the ground state or the excited state. This is known as a measurement. The act of measurement causes the electron to choose a definite state.

A qubit is the basic unit of quantum information. It can be represented by the state of an electron, which exists in a superposition of the ground and excited states with complex amplitudes. When measured, the qubit collapses into one of the two states.

When dealing with quantum information, it is important to understand the concept of qubits. A qubit is the basic unit of quantum information and can be represented by a two-level quantum system. These two levels are often referred to as the ground state (0) and the excited state (1). However, unlike classical bits which can only be in one of these two states at a time, qubits can exist in a superposition of both states simultaneously.

In a superposition, a qubit can be in a combination of the ground and excited states, with certain probabilities assigned to each state. These probabilities are represented by complex numbers called amplitudes. The probability of the qubit being in the ground state is determined by the squared magnitude of the amplitude α , while the probability of it being in the excited state is determined by the squared magnitude of the amplitude β .

When a measurement is performed on a qubit, the superposition collapses and the qubit "chooses" to be in either the ground or excited state with certain probabilities. This collapse occurs because the act of measurement disturbs the state of the qubit. It is important to note that the probabilities of the qubit being in either state must add up to 1, which is why the state is normalized.

The idea of a qubit existing in a superposition and the collapse of the superposition upon measurement can be difficult to interpret. Many interpretations have been proposed, but none have gained widespread consensus. However, the mathematical framework of quantum mechanics allows us to work with and manipulate qubits effectively, even if the underlying interpretation remains elusive.

Another interesting aspect of qubits is that their state is complex and requires a large amount of information to fully describe when not being observed. In the case of a qubit, this means specifying two complex numbers, which represents an infinite number of bits of information. However, when the qubit is observed or measured, it simplifies to either the ground or excited state, which can be represented by a single classical bit.

Qubits are the fundamental units of quantum information. They can exist in a superposition of the ground and excited states, with probabilities determined by complex amplitudes. Upon measurement, the qubit collapses

into either the ground or excited state with certain probabilities. The interpretation of qubits in a superposition remains a topic of debate, but the mathematical framework allows us to work with them effectively. Additionally, the state of a qubit is complex when not being observed, but simplifies to a classical bit when measured.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM INFORMATION****TOPIC: GEOMETRIC REPRESENTATION**

A quantum bit, or qubit, is the state of a system such as an electron in a hydrogen atom when it is confined to its ground or excited state. The general state of a qubit is a superposition of the ground and excited states, which can be written as $\alpha|0\rangle + \beta|1\rangle$, where α and β are complex numbers that are normalized, meaning that the magnitude squared of α plus the magnitude squared of β equals 1.

To specify the state of a qubit, we need two complex numbers. One way to represent this state is by stacking the two numbers on top of each other, like $\begin{pmatrix} \alpha \\ \beta \end{pmatrix}$. This representation suggests that the state is a vector in a two-dimensional vector space.

The vector space representing the qubit state is two-dimensional and complex because the entries are allowed to be complex. The vector representing the state is normalized, meaning that the magnitude squared of α plus the magnitude squared of β is equal to the square of the length of the vector, which is 1.

The state 0 corresponds to the vector $(1, 0)$, while the state 1 corresponds to the vector $(0, 1)$. If we plot these vectors, we can see that they correspond to the ground state and excited state, respectively. Other vectors representing different qubit states can also be plotted on the unit circle.

We have learned that a qubit is a unit vector in a two-dimensional complex vector space. The notation used to represent qubit states, called ket notation or Dirac's ket notation, is another way of writing vectors. This notation allows us to name the states as 0 and 1, representing the encoding of information, while also acknowledging that the qubit is a superposition of these states, represented as a vector.

Now that we have a geometric interpretation of qubit states, let's understand what it means to measure a qubit. We can define other states, such as the state Ψ , which makes an angle θ with the 0 state. In a two-dimensional real space, the state can be written as $\cos\theta|0\rangle + \sin\theta|1\rangle$.

A qubit is a unit vector in a two-dimensional complex vector space. The geometric interpretation allows us to visualize qubit states as vectors on the unit circle. The ket notation represents both the encoding of information as 0 and 1 and the superposition of these states as a vector. Measuring a qubit involves determining the angle it makes with the 0 state.

In the study of quantum information, it is important to understand the concept of measurement and its interpretation. When a measurement is performed on a qubit, it collapses into one of its possible states. This collapse occurs with a certain probability, which can be calculated using the cosine squared and sine squared of the angle of the measurement.

Let's consider a qubit in a superposition state, represented as $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$, where α and β are complex numbers. When a measurement is made on this qubit, it will be projected onto either the ground state $|0\rangle$ or the excited state $|1\rangle$. The probability of the qubit being projected onto the ground state is given by the square of the cosine of the angle it makes with the ground state, which is $\cos^2\theta$. Similarly, the probability of the qubit being projected onto the excited state is given by the square of the sine of the angle it makes with the ground state, which is $\sin^2\theta$.

It is important to note that the act of measurement disturbs the system, causing the qubit to actually become the measured state. Therefore, after the measurement, the qubit will be in either the ground state or the excited state.

Another way to interpret measurement in quantum information is through a geometric representation. In this representation, the state of the qubit is projected onto either the ground state or the excited state with certain probabilities. The probability of projection onto the ground state is $\cos^2\theta$, where θ is the angle between the state vector and the ground state. Similarly, the probability of projection onto the excited state is $\sin^2\theta$, where θ is the angle between the state vector and the excited state.

This measurement process can be thought of as a projection onto a standard basis, which consists of the ground state and the excited state. By measuring the state of the qubit, it is projected onto one of these two states with a probability determined by the angle it makes with each state.

Furthermore, it is possible to perform measurements in other bases besides the standard 0-1 basis. By choosing a different orthogonal basis, such as the U - $U^\perp$ basis, the measurement process remains the same. The state of the qubit will be projected onto the U state with probability $\cos^2 \theta$ and onto the $U^\perp$ state with probability $\sin^2 \theta$.

In this context, measuring the state in a different basis means determining whether the qubit is in a specific superposition of the ground and excited states. For example, if the state is $\frac{1}{\sqrt{2}} |0\rangle + \frac{1}{\sqrt{2}} |1\rangle$ and the $U^\perp$ state is $-\frac{1}{\sqrt{2}} |0\rangle + \frac{1}{\sqrt{2}} |1\rangle$, the measurement is aimed at determining which of these two states the qubit is in.

Quantum mechanics allows for this kind of measurement, where the state is measured in a basis other than the standard basis. This flexibility in measurement is a fundamental aspect of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM INFORMATION****TOPIC: PHOTON POLARIZATION**

In the context of quantum information, an important concept to understand is photon polarization. Photons, which are particles of light, possess a property called polarization, which can carry information in the form of a quantum bit or qubit.

To visualize this concept, we can think of light as an electromagnetic wave traveling in a certain direction. The electrical field associated with the light wave oscillates in an orthogonal direction to its movement. The orientation of these electrical field oscillations determines the polarization state of the photon.

If the electrical field oscillations are horizontally oriented, we can assign a polarization state of 0 to the photon. Conversely, if the oscillations are vertically oriented, the polarization state is assigned as 1. When the polarization is at a diagonal angle, the state of the qubit is a superposition of 0 and 1, represented as $\frac{1}{\sqrt{2}} * 0 + \frac{1}{\sqrt{2}} * 1$.

To measure the polarization of a qubit, a polarizing filter or lens is used. This filter has a specific orientation, such as vertical or horizontal. If a vertically polarized photon passes through a vertically oriented filter, it is transmitted. However, if the photon is horizontally polarized, it is blocked.

When a qubit is in a superposition state, such as a combination of vertical and horizontal polarization, the probability of transmission is determined by the cosine squared of the angle (theta) between the polarization and the orientation of the filter. If the qubit is transmitted, its new state becomes the original polarization state. If the qubit is blocked, its new state is orthogonal to the original polarization state.

By changing the orientation of the lens, the basis for measurement can be altered. For example, if the lens is oriented at a 45-degree angle, photons with diagonal polarization will be transmitted, while those with a different orientation will be blocked.

To illustrate this concept, let's consider an experiment with two lenses. The lens in the back is vertically oriented, while the one in front is horizontally oriented. When a beam of light passes through these lenses, the photons within the beam will either be transmitted or blocked based on their polarization. In the case of a single photon, its original polarization state determines its interaction with the lenses. If the photon is transmitted through the back lens, its new state becomes vertically polarized. However, when it encounters the front lens, which is horizontally oriented, the photon is blocked.

This experiment demonstrates how the interaction between polarizing lenses and photons can result in the transmission or blocking of light, depending on the polarization state. In the quantum context, where only a single photon is considered, the blocking of the photon leads to a dark spot in the observed area.

Photon polarization is a fundamental concept in quantum information. Photons can carry information in the form of qubits, where the polarization state represents the quantum state. Polarizing filters or lenses can be used to measure the polarization of a qubit, transmitting or blocking photons based on their polarization state. By changing the orientation of the lens, the basis for measurement can be altered. Understanding photon polarization is crucial for further exploration of quantum information and its applications.

When studying quantum information, one important concept to understand is photon polarization. In this context, polarization refers to the orientation of the electric field of a photon. The polarization of a photon can be vertical, horizontal, or any combination in between.

To illustrate the effect of interposing a lens at a 45-degree angle between two other lenses, let's consider a scenario. Initially, a photon is vertically polarized. When it encounters the back lens, it has a probability of being transmitted through, which is given by the cosine squared of the angle between the polarization and the orientation of the lens. In this case, the photon becomes vertically polarized.

Next, the photon encounters the middle lens, which is oriented at a 45-degree angle. The photon has a 50% chance of being transmitted through and a 50% chance of being blocked. If it is transmitted, its polarization

changes to a combination of vertical and horizontal, specifically $\frac{1}{\sqrt{2}}$ over the square root of 2 times vertical plus $\frac{1}{\sqrt{2}}$ over the square root of 2 times horizontal.

Finally, the photon reaches the front lens, which is horizontally oriented. If the photon is transmitted through the lens, it has a 50% chance of this happening, its polarization becomes purely horizontal.

The overall effect of interposing the lens in the middle is that the photon has a quarter chance of being transmitted after considering the transmission probabilities of the first and second lenses. This means that the amount of light coming through is faint, or there is some chance of transmission.

Understanding the behavior of photons and their polarization is crucial in the field of quantum information, as it forms the foundation for various applications such as quantum communication and quantum computing.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM INFORMATION****TOPIC: UNCERTAINTY PRINCIPLE**

The uncertainty principle is a fundamental concept in quantum mechanics that states that we cannot know both the position and velocity of a particle with perfect accuracy. This principle also applies to qubits, which are the basic units of quantum information. In this lecture, we will explore how the uncertainty principle applies to qubits and what it means in the context of quantum information.

To understand the uncertainty principle, let's first revisit the double-slit experiment discussed in the previous lecture. In this experiment, an electron passes through one of two slits and creates an interference pattern on a screen. However, if we try to determine which slit the electron went through, we disturb the system and destroy the interference pattern. This is because the act of measuring the position of the electron changes its velocity or momentum.

Heisenberg formulated the uncertainty principle by stating that we can never know both the position and velocity of a particle with perfect accuracy. This applies to the double-slit experiment as we were trying to determine the position of the electron. In our attempt to do so, we inadvertently changed its velocity or momentum, leading to the destruction of the interference pattern.

Now, let's apply this uncertainty principle to qubits. A qubit can exist in two different bases, the zero-one basis and the plus-minus basis. In the zero-one basis, a qubit can be in either the state 0 or 1. In the plus-minus basis, a qubit can be in a superposition of the states 0 and 1, represented by the vectors $|+\rangle$ and $|-\rangle$ respectively.

If we were to measure a qubit in the zero-one basis, we would determine its bit value, whether it is 0 or 1. Similarly, if we were to measure a qubit in the plus-minus basis, we would determine its sign value, whether it is + or -. These measurements can be thought of as determining the position (bit value) and velocity (sign value) of the qubit.

The question now arises: can we ever know both the bit value and sign value of a qubit with perfect accuracy? The answer is no. If we know the bit value of a qubit perfectly, it must be in either the state 0 or 1. Similarly, if we know the sign value perfectly, it must be either + or -. However, a qubit can also exist in superposition states, where both the bit value and sign value are uncertain.

The uncertainty principle applies to qubits as well, stating that we cannot know both the bit value and sign value of a qubit with perfect accuracy. This principle highlights the inherent uncertainty and probabilistic nature of quantum information.

In the field of Quantum Information, one of the fundamental concepts is the Uncertainty Principle. This principle states that it is impossible to perfectly know both the bit value and the sine value of a quantum state. To understand this principle, let's consider the example of a state called 'side'. As the state 'side' tries to get closer to either 0 or 1, it gets farther from the states plus and minus. This is because the states plus and minus make a 45-degree angle with each other. Therefore, if 'side' is close to 0, it must make at least a 22.5-degree angle with both plus and minus, and the same applies if it is close to 1. This leads us to the uncertainty principle, which states that the bit value and the sine value cannot be perfectly known at the same time.

A similar situation arises when considering the position and velocity of momentum. However, in this case, we are working in a more complex vector space. Working with qubits, which are quantum bits, allows us to understand the uncertainty principle in a simpler setting.

To quantify the uncertainty in knowing the bit value and the sine value, we can define a measure called 'spread'. In the standard basis (0 and 1 basis), the spread is defined as the absolute value of α_0 plus the absolute value of α_1 , where α_0 and α_1 are the amplitudes of the state 'side' in the 0 and 1 basis. In the sign basis (plus/minus basis), the spread is defined as the absolute value of β_0 plus the absolute value of β_1 , where β_0 and β_1 are the amplitudes of the state 'side' in the plus and minus basis.

When the bit value is known perfectly, the spread is 1, as we can determine whether α_0 or α_1 is equal to 1 and the other is equal to 0. On the other hand, when we don't know the bit value at all (e.g., in the state

plus), the spread is square root 2. The claim is that the spread can only be small (i.e., 1) if the bit value is known perfectly. The farther the spread is from 1, the less certain we are about the bit value. The same principle applies to the spread in the plus/minus basis.

The uncertainty principle for the spread in the standard basis and the spread in the sign basis of any qubit states that their product is at least square root 2. This means that both values cannot be 1 simultaneously. At least one of them must be at least the fourth root of 2.

The uncertainty principle in Quantum Information states that it is impossible to perfectly know both the bit value and the sine value of a quantum state. This principle applies to qubits and is quantified by the spread in the standard and sign bases. The spread represents the uncertainty in knowing the bit value and the sine value, and the uncertainty principle states that their product is at least square root 2.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: K-LEVEL SYSTEM AND BRA-KET NOTATION**

In this lecture, we will discuss systems of two qubits. These quantum systems exhibit a property known as entanglement, which plays a critical role in quantum computation. Before diving into entanglement, let's formalize our understanding of qubits and superpositions.

Recall that the energy of an electron in an atom is quantized, meaning it can only have specific energy levels. Let's assume the electron is in the ground state or one of the excited states up to the $K-1$ th excited state. In a classical system, this electron could store K bits of information, denoted as $0, 1, \dots, K-1$.

The superposition principle, a fundamental axiom of quantum mechanics, states that the general state of the system is a linear superposition of these allowable states. In other words, the system can be in a state represented as a linear combination of 0 through $K-1$, each with an amplitude α_J , a complex number. These amplitudes are normalized, meaning the sum of their magnitudes squared equals 1 for J ranging from 0 to $K-1$.

Interpreting this state is challenging because it's not easy to grasp the meaning of, for example, the electron being in the ground state with an amplitude of $-1/2$ or $1/2 + i/2$. However, the measurement axiom provides a way to interpret it. When we measure the system, the probability of observing outcome J is the magnitude squared of α_J . The normalized state guarantees that with probability 1 , we will observe an outcome J between 0 and $K-1$. Additionally, a measurement disturbs the system, and the new state, denoted as ψ' , will be the J th excited state if the measurement outcome is J .

Let's consider a quick example with $K = 3$. We have a three-state system, and our state could be represented as 0 with amplitude $1/2 + i/2$, 1 with amplitude $-1/2$, and 2 with amplitude $i/2$. If we measure the system, the probability of observing 0 is $1/2$, and the new state will be 0 . The probability of observing 1 is $1/4$, and the new state will be 1 . The probability of observing 2 is also $1/4$, and the new state will be 2 .

To better understand the concept of superposition, let's consider a geometric interpretation of the quantum state. The superposition principle states that the state of a K -level quantum system is a unit vector in a K -dimensional complex vector space, also known as a Hilbert space. This vector space has an orthonormal basis consisting of the states $|0\rangle, |1\rangle, \dots, |K-1\rangle$. The state ψ can be represented as a unit vector in this vector space, written as $\alpha_0|0\rangle + \alpha_1|1\rangle + \alpha_2|2\rangle$, where $\alpha_0, \alpha_1, \alpha_2$ are complex numbers. If we consider a three-dimensional complex vector space, we can write it as $\alpha_0, \alpha_1, \alpha_2$ in standard vector notation.

When we measure the system, the state vector ψ gets projected onto one of the basis states. If we are measuring in the standard basis $|0\rangle, |1\rangle, |2\rangle$, the state ψ gets projected onto the state $|0\rangle$ with a probability equal to the cosine squared of the angle θ_0 it makes with the state $|0\rangle$. In general, the probability of the outcome being 0 is cosine squared θ_0 . If that's the outcome, the state ψ gets projected onto the state $|0\rangle$. This holds true for each vector in the orthonormal basis, and the probability of projection is given by cosine squared θ , where θ is the angle it makes with the particular vector.

To define the angle θ or cosine θ between two different vectors, we use the inner product of the vectors. If we have two vectors ψ and ϕ , both complex vectors, the cosine θ is defined as the inner product of ψ and ϕ divided by the product of their magnitudes.

We have discussed the concept of entanglement and the formalization of qubits and superpositions. We explored the superposition principle, which states that the general state of a quantum system is a linear superposition of allowable states. We also examined the measurement axiom, which determines the probabilities of observing outcomes and the resulting disturbance to the system. Finally, we looked at the geometric interpretation of the quantum state, where the state is represented as a unit vector in a complex vector space.

In the study of quantum information, one important concept is quantum entanglement. Quantum entanglement refers to the phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particle. This concept is fundamental to understanding the behavior of quantum systems and has applications in various areas such as quantum

computing and quantum communication.

To describe quantum entanglement, we use the k-level system and the bra-ket notation. In the k-level system, we consider a quantum system that can exist in k different states. In the case of a qubit, which is a two-level system, we have the states $|0\rangle$ and $|1\rangle$, often referred to as the standard basis states. These states form a basis for the qubit system, meaning that any state of the qubit can be expressed as a linear combination of these basis states.

In addition to the standard basis states, we also have the states $|+\rangle$ and $|-\rangle$, which are known as the plus and minus states. The plus state is an equal superposition of the $|0\rangle$ and $|1\rangle$ states, while the minus state is a specific combination of the $|0\rangle$ and $|1\rangle$ states. The minus state can be expressed as $1/\sqrt{2}|0\rangle - 1/\sqrt{2}|1\rangle$.

To understand the concept of entanglement, let's consider an example. Suppose we have a state $|\psi\rangle$ that is given by $1/2|0\rangle + \sqrt{3}/2|1\rangle$. This state forms an angle of 45 degrees with the $|+\rangle$ state and an angle of 60 degrees with the $|-\rangle$ state. The angle between two states can be defined using the cosine of the angle, which in this case is 15 degrees.

If we were to measure the state $|\psi\rangle$ in the plus/minus basis, the probability of observing the plus outcome would be given by the square of the inner product between the $|\psi\rangle$ state and the $|+\rangle$ state. The inner product can be calculated by multiplying the corresponding coordinates of the two vectors and taking the magnitude squared. In this case, the probability of observing the plus outcome is $2 + \sqrt{3}/4$.

Alternatively, we can rewrite the state $|\psi\rangle$ as a linear combination of the plus and minus states. By doing so, we can calculate the probability of observing the plus outcome directly. In this example, the probability of observing the plus outcome is also $2 + \sqrt{3}/4$.

It is worth noting that the probability of observing the minus outcome can be obtained by subtracting the probability of the plus outcome from 1. In this case, the probability of observing the minus outcome is $2 - \sqrt{3}/4$.

Quantum entanglement is a fundamental concept in quantum information. It involves the correlation between particles, where the state of one particle cannot be described independently of the state of another particle. The k-level system and bra-ket notation are used to describe quantum states and calculate probabilities. Understanding quantum entanglement is crucial for various applications in the field of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: SYSTEMS OF TWO QUBITS**

In the study of quantum information, we often encounter systems of two qubits. To illustrate this concept, let's consider the example of a hydrogen atom. We can use the ground or excited state of the electron to represent a bit of information. Since there are two such electrons, we can represent two classical bits of information.

In classical systems, there are four possible states for two bits: 00, 01, 10, and 11. However, in quantum systems, the superposition principle allows for a more complex representation. The quantum state of these two electrons can be described as a superposition of all four possibilities:

$$\alpha 00 + \alpha 01 + \alpha 10 + \alpha 11$$

Here, α represents a complex number for each of the four possibilities, and the sum of the magnitudes squared of these complex numbers is equal to 1, ensuring a normalized state.

When we measure the state of these two qubits, the electrons quickly "make up their minds" and collapse into one of the classical states. The probability of observing a specific state is equal to the magnitude squared of the corresponding complex number. For example, if the state of our system is given by:

$$(1/2 + i/2) |00\rangle + (1/2) |01\rangle - (i/2) |11\rangle$$

The probability of observing 00 would be $1/2$, the probability of observing 01 would be $1/4$, and the probability of observing 11 would also be $1/4$.

Now, let's consider a different scenario. Suppose we have the state:

$$(1/2) |01\rangle + (i/2) |11\rangle$$

If we were to measure only the first qubit, what would we observe? The probability of observing 0 on the first qubit is the same as the probability of observing 0 on both qubits. In this case, it would be:

$$|\alpha_{00}|^2 + |\alpha_{01}|^2$$

To determine the new state, we cross out the possibilities that are not consistent with the observed outcome. In this case, we cross out the possibilities $|01\rangle$ and $|11\rangle$, resulting in the new state:

$$(1/2) |00\rangle$$

To normalize this state, we divide it by the square root of the probability of observing 0, which is:

$$\sqrt{(|\alpha_{00}|^2 + |\alpha_{01}|^2)}$$

Let's apply this process to our example. The probability of observing 0 is:

$$(|1/2|^2 + |i/2|^2) = 1/2 + 1/4 = 3/4$$

The new state, after crossing out the inconsistent possibilities, is:

$$(1/2) |00\rangle$$

To normalize this state, we divide it by the square root of the probability of observing 0, which is $\sqrt{3/4}$.

Systems of two qubits allow for more complex representations of information due to the superposition principle. When measuring the state of these qubits, the probabilities of observing specific outcomes are determined by the magnitudes squared of the corresponding complex numbers. The observed outcome then affects the new state, which can be obtained by crossing out inconsistent possibilities and normalizing the remaining state.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: ENTANGLEMENT**

In the study of quantum information, one of the fundamental concepts is quantum entanglement. Entanglement refers to a phenomenon where two or more quantum systems become so interconnected that their individual states cannot be described independently. In this didactic material, we will explore the concept of entanglement and its implications.

To understand entanglement, let's consider a system of two qubits. We are given the state of each qubit individually. The state of the first qubit is represented as $\alpha_0 |0\rangle + \alpha_1 |1\rangle$, and the state of the second qubit is $\beta_0 |0\rangle + \beta_1 |1\rangle$. Our goal is to determine the state of the composite system.

Informally, we can think of the state of the composite system as the product of the individual qubit states. By multiplying the amplitudes of each possible combination, we obtain the state of the composite system. For example, if the first qubit is in the plus state and the second qubit is in the state $\frac{1}{2} |0\rangle + \frac{\sqrt{3}}{2} |1\rangle$, the composite system is in the state $\frac{1}{2\sqrt{2}} |00\rangle + \frac{\sqrt{3}}{2\sqrt{2}} |01\rangle + \frac{1}{2\sqrt{2}} |10\rangle + \frac{\sqrt{3}}{2\sqrt{2}} |11\rangle$.

Now, let's consider a different scenario. Suppose we are given the state of the two qubits together and we are asked to find the state of each qubit separately. In some cases, like the previous example, we can easily determine the individual qubit states. However, in general, it is not always possible to factorize the composite state into the states of the individual qubits.

To illustrate this, let's examine a simple state: an equal superposition of $|00\rangle$ and $|11\rangle$ with amplitude $1/\sqrt{2}$. If we assume that this state can be factorized as $(\alpha_0 |0\rangle + \alpha_1 |1\rangle)(\beta_0 |0\rangle + \beta_1 |1\rangle)$, we can expand the expression and compare coefficients. By doing so, we find that both $\alpha_0 \beta_0$ and $\alpha_1 \beta_1$ are nonzero. However, the fact that $\alpha_0 \beta_1$ and $\alpha_1 \beta_0$ must both be equal to 0 leads to a contradiction. Therefore, we conclude that this state cannot be factorized into the states of the individual qubits.

This inability to factorize the state of an entangled system highlights a fundamental property of quantum systems. When two quantum systems interact with each other, they can become entangled to the point where their individual states are no longer independent. Even if we separate the entangled systems by a large distance, they remain entangled.

Now, let's consider the measurement of an entangled system. Suppose we measure the first qubit. The probability of observing 0 is $1/2$, and the new state becomes $|00\rangle$. Similarly, the probability of observing 1 is also $1/2$, and the new state becomes $|11\rangle$. The measurement of one qubit affects the state of the other qubit, regardless of the distance between them.

Entanglement is a phenomenon in quantum information where two or more quantum systems become interconnected to the point where their individual states cannot be described independently. This entangled state persists even when the systems are separated by a large distance. The inability to factorize the state of an entangled system highlights the unique nature of quantum systems.

When two qubits are brought together and allowed to interact, they can become entangled. Entanglement is a phenomenon in quantum mechanics where the state of one particle is dependent on the state of another, regardless of the distance between them.

For example, if we measure the first qubit and get outcome 0, the probability that the second qubit will also be measured as 0 is certain. Similarly, if the first qubit gives us outcome 1, the second qubit will also give us 1. This correlation between the two qubits seems mysterious, as it implies that the outcome of one qubit affects the outcome of the other, even when they are far apart.

One way to explain this phenomenon is to imagine that when the qubits were brought together and allowed to interact, they randomly decided to be in the same state. They agreed upon both being in the state 0 or both being in the state 1, each with a probability of $1/2$. So, when we measure one qubit and see 0, the other qubit is also in the state 0.

However, this explanation falls short when we consider the more mysterious properties of entanglement. In the next material, we will explore these properties that cannot be explained by classical intuition.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: EPR PARADOX**

Quantum Entanglement and the EPR Paradox

In this material, we will explore the concept of quantum entanglement and the EPR paradox. Quantum entanglement is a phenomenon in quantum mechanics where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the other particles. The EPR paradox, named after its inventors Einstein, Podolsky, and Rosen, is a thought experiment that highlights the seemingly paradoxical nature of entanglement.

To understand the EPR paradox, let's first revisit the Bell state. The Bell state is an entangled state of two qubits, which can be represented as a superposition of the states $|00\rangle$ and $|11\rangle$. This means that the two qubits are in a simultaneous state of both being in the ground state and being in the excited state.

Interestingly, the Bell state can also be expressed in terms of the plus/minus basis. In this basis, the state of the two qubits can be written as a superposition of $|++\rangle$ and $|--\rangle$ states. The plus state for a qubit corresponds to a superposition of the ground state and the excited state, while the minus state corresponds to a superposition with opposite signs.

To demonstrate this, let's expand the Bell state in the plus/minus basis. We have $1/\sqrt{2}|++\rangle + 1/\sqrt{2}|--\rangle$. By collecting the coefficients, we find that the Bell state can be expressed as $1/\sqrt{2}(|00\rangle + |11\rangle)$.

Now, let's delve into the EPR paradox. The EPR paradox arises when we consider the measurement of the entangled qubits. If we measure the first qubit and find it to be in the state $|0\rangle$, we can be certain that the second qubit will also be in the state $|0\rangle$. Similarly, if the first qubit is measured to be in the state $|1\rangle$, the second qubit will also be in the state $|1\rangle$.

However, the paradox arises when we introduce measurements in the sign basis. In the sign basis, if the first qubit is in the state $|+\rangle$, there is a 50% probability that it will be in the state $|-\rangle$ and a 50% probability that it will be in the state $|+\rangle$. But regardless of the outcome, if the first qubit is in the state $|+\rangle$, the second qubit will also be in the state $|+\rangle$. Similarly, if the first qubit is in the state $|-\rangle$, the second qubit will also be in the state $|-\rangle$.

This paradox led Einstein, Podolsky, and Rosen to question the completeness of quantum mechanics. They argued that since the two entangled qubits can be far apart from each other, any measurement on one qubit should not affect the other qubit. Yet, according to the principles of quantum mechanics, the measurement of one qubit determines the state of the other qubit, regardless of the chosen basis.

The EPR paradox highlights the non-local nature of entanglement, where the state of one particle is intimately connected to the state of another particle, even when they are separated by large distances. This seemingly instantaneous connection between entangled particles challenges our classical intuition about the nature of reality.

Quantum entanglement and the EPR paradox are fundamental concepts in quantum information. Entanglement allows for correlations between particles that defy classical descriptions, and the EPR paradox questions the completeness of quantum mechanics by highlighting the non-local nature of entangled states.

Quantum entanglement is a phenomenon in quantum mechanics where two particles become connected in such a way that the state of one particle is dependent on the state of the other, regardless of the distance between them. This concept was famously discussed by Einstein, Podolsky, and Rosen (EPR) in what is known as the EPR paradox.

According to the uncertainty principle in quantum mechanics, it is not possible to simultaneously know both the bit value and the sign value of a particle. However, in the case of entangled particles, it appears that this principle is violated. When the bit value of the first particle is measured, it disturbs the sign value of the second particle. However, since the two particles can be far apart, measuring the bit value of the first particle does not change the sign value of the second particle.

Einstein, Podolsky, and Rosen concluded that quantum mechanics must be an incomplete theory. They believed that behind the scenes, nature actually defines all the physical quantities and that quantum mechanics only limits the amount of information we can obtain about nature. In an attempt to find a more complete theory, Einstein spent the last 20 years of his life searching, but ultimately failed.

In quantum mechanics, the state of the second particle is not defined by itself because it is entangled with the first particle. Therefore, saying that the bit and sign values of the second particle are well-defined does not make sense.

In the next lecture, we will explore more interesting properties of entanglement and discover that there is much more to learn about this phenomenon. It is intriguing to consider that if Einstein had known about these properties, he might have saved himself 20 years of effort.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: BELL AND EPR**

In the previous lecture, we discussed the concept of entanglement and the EPR paradox, which highlighted Einstein's belief that quantum mechanics is an incomplete theory. In this lecture, we will focus on John Bell's groundbreaking paper from 1965, which demonstrated that entanglement has testable effects that can challenge Einstein's ideas.

To better understand Einstein's beliefs, we need to explore the concept of local realism. Local realism suggests that physics must be local, meaning that physical interactions can only occur through direct proximity or contact. This idea dates back to Isaac Newton, who found the notion of action at a distance in his theory of gravity to be unsettling. Newton reluctantly published his ideas about this theory, stating that it is inconceivable for inanimate matter to operate upon another object without mutual contact.

Realism, on the other hand, asserts that physical entities have a separate reality independent of measurements. Einstein himself supported this idea, stating that matter, such as an electron, possesses properties like spin and location even when not being measured. He even used the example of the moon, expressing his belief that it exists even when he is not observing it.

These concepts are at the core of the quantum mechanics debate. Quantum mechanics suggests that a system can exist in a superposition state, where the properties of interest only manifest when measured. This idea troubled Einstein and posed a challenge to the notion of realism.

Now, let's revisit the EPR paradox briefly. In this thought experiment, two qubits are entangled in a Bell state. When brought close together and interacted with each other, they enter into this entangled state. Subsequently, the qubits are separated by a significant distance. We learned in the previous lecture that the Bell state can be described as an equal superposition of 0 0 & 1 1 or as an equal superposition of plus plus and minus minus.

The concepts of locality and realism come into play here. Locality suggests that since the qubits are far apart, any action performed on one qubit should not affect the state of the other qubit because there hasn't been enough time for light or any other influence to travel between them. Realism, on the other hand, argues that the properties of the qubits, such as the bit value (0 or 1) and the sign value (plus or minus), exist independently of measurement.

Based on these principles, one could reason that by measuring the first qubit, one could determine the values of the second qubit without disturbing it. For example, if the bit value of the first qubit is measured to be 0, then the bit value of the second qubit must also be 0, and the same applies to the sign value. This reasoning suggests that the measurements on one qubit do not disturb the other qubit, and therefore, the properties of the two qubits are unchanged.

In 1965, John Bell published a landmark paper that presented an experiment capable of distinguishing between quantum mechanics and any theory consistent with local realism. This experiment offers a remarkable opportunity to test the predictions of these two theories. According to Bell, the experiment's results can be used to estimate a quantity, denoted as E , which should be less than or equal to $3/4$ if nature behaves in accordance with local realism. Conversely, if nature follows the principles of quantum mechanics, the estimated value of E should ideally be $\cos^2(\pi/8)$, approximately 0.85.

Bell's experiment provided a tangible way to differentiate between quantum mechanics and local realism, something that the EPR paradox alone could not achieve. This experiment marked a significant milestone in the field of quantum information, demonstrating that entanglement has testable effects that can challenge established beliefs about the nature of reality.

Quantum Information - Quantum Information Fundamentals - Quantum Entanglement - Bell and EPR

Quantum mechanics is a fundamental theory that describes the behavior of particles at the subatomic level. It has been extensively tested and has consistently shown results that are inconsistent with classical physics. One of the key phenomena in quantum mechanics is quantum entanglement.

Quantum entanglement refers to the phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particles. This correlation exists even when the particles are separated by large distances. This concept was first introduced by Albert Einstein, Boris Podolsky, and Nathan Rosen in their famous EPR paper in 1935.

To understand the implications of quantum entanglement, we need to discuss Bell's experiment and the Bell inequalities. Bell's experiment was designed to test the predictions of quantum mechanics against the concept of local realism. Local realism suggests that physical properties of particles exist independently of measurement and that there is a limit to how correlated two particles can be.

In Bell's experiment, two entangled particles are measured independently and the correlations between their measurements are analyzed. The results of these measurements are compared with the predictions of local realism. The Bell inequalities, derived by physicist John Bell, provide a mathematical way to quantify the correlations between the measurements.

Numerous experiments have been conducted to test Bell's inequalities, and the results have consistently shown that the predictions of quantum mechanics are in agreement with the experimental data, while local realism fails to explain the observed correlations. This experimental confirmation of quantum entanglement has profound implications for our understanding of the nature of reality.

The concept of quantum entanglement and the violation of Bell's inequalities have paved the way for the development of quantum information and quantum computation. Quantum information science utilizes the unique properties of quantum systems, such as superposition and entanglement, to perform tasks that are not possible with classical information processing.

By understanding the details of Bell's experiment and the violation of Bell's inequalities, we gain insight into the limitations of classical physics and the power of quantum mechanics. This understanding forms the foundation for further exploration and advancements in the field of quantum information.

Quantum entanglement and the violation of Bell's inequalities have revolutionized our understanding of the quantum world. These phenomena highlight the limitations of classical physics and the unique properties of quantum systems. The study of quantum information and quantum computation builds upon these concepts, opening up new possibilities for information processing and technological advancements.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: ROTATIONAL INVARIANCE OF BELL STATE**

Quantum Entanglement and Rotational Invariance of Bell State

In order to understand Bell's experiment and Bell's inequalities, we need to delve deeper into the properties of entanglement. Let's start by revisiting the Bell State from EPR. We know that the Bell State can be written as an equal superposition of 00 & 11 or as an equal superposition of plus plus and minus minus. However, this is just a specific case of a more general property of Bell States - they can be written in any rotated basis.

To understand this, let's consider a basis of 01 for the first qubit (ground and excited states) and a similar basis for the second qubit. We can rotate this basis by an arbitrary angle and obtain a rotated basis, let's call it " u ". The state orthogonal to " u " can be obtained by rotating the zero-one basis by some angle θ . Similarly, we can define the " u perp" basis for the second qubit.

Now, if we express the Bell state in terms of the " u " basis, we find that it can be written as an equal superposition of 00 and 11 or as an equal superposition of " $u u$ " and " u perp u perp". This means that if we measure the first qubit in the " u " basis, the probability of obtaining the outcome " u " is exactly $1/2$. Moreover, if we measure the second qubit in the " u " basis as well, we will always get the same result for both qubits, regardless of the rotated basis we choose to measure in.

Let's now consider a slightly different scenario. We still have the zero-one basis for the first qubit, but now let's measure the second qubit in the " $v v$ perp" basis. We want to know the probability of getting matching outcomes if we measure the first qubit in the " u " basis and the second qubit in the " v " basis. In other words, what's the chance of getting " v " as the outcome of the second measurement if we observe " u " as the outcome of the first measurement?

To answer this question, we can rely on the rotational invariance of the Bell State. When we measure the first qubit and obtain the outcome " u ", the new state of the system becomes " $u u$ ". If we imagine an angle θ between the " u " and " v " bases, then when we measure the second qubit in the " $v v$ perp" basis, the probability of obtaining " v " as the outcome is given by cosine squared θ . The same holds if the outcome of the first measurement was " u perp". In this case, the new state would be " u perp u perp", and the probability of obtaining " v perp" as the outcome of the second measurement would again be cosine squared θ .

We have derived a new principle: if we measure two qubits in two different bases that make an angle θ with each other, the probability of getting matching outcomes on the two measurements is exactly cosine squared θ . Conversely, the probability of getting non-matching outcomes (e.g., " u " and " v perp" or " u perp" and " v ") is sine squared θ .

To establish the rotational invariance, let's consider the zero-one basis and suppose " u " can be written as $a0 + b1$. When we rotate it through 90 degrees, we obtain $-b0 + a1$. Therefore, " u perp" is equal to $-b0 + a1$. By substituting " u " and " v " into the expression, we find that the state $1/\sqrt{2} (u u + u$ perp u perp) is equal to $1/\sqrt{2} (a0 + b1)(a0 + b1) + (-b0 + a1)(-b0 + a1)$. Collecting terms, we find that the amplitude of 00 is $a^2 + b^2$, and the amplitude of 11 is also $a^2 + b^2$. Everything else cancels out, resulting in a normalized state.

We have explored the rotational invariance of the Bell State and its implications for measurements in different bases. The probability of obtaining matching outcomes on two measurements depends on the angle between the bases, while the probability of getting non-matching outcomes is determined by the sine squared of that angle.

In the study of Quantum Information, one of the fundamental concepts is Quantum Entanglement. Quantum Entanglement refers to the phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particle(s).

One specific example of Quantum Entanglement is the Bell state, which is a maximally entangled state. The Bell state can be represented as $(1/\sqrt{2})(|00\rangle + |11\rangle)$, where $|0\rangle$ and $|1\rangle$ represent the two possible states of a

qubit.

It is important to note that the Bell state ($|00\rangle + |11\rangle$) is rotationally invariant with respect to real rotations. This means that if we apply a real rotation to the system, the state remains unchanged. However, in order for this rotational invariance to hold, the coefficients a and b in the Bell state equation must be real numbers.

If we want a state that is invariant under all complex rotations, we need to consider a different Bell state known as the s_y - state. The s_y - state can be represented as $(1/\sqrt{2})(|01\rangle - |10\rangle)$, where $|0\rangle$ and $|1\rangle$ represent the two possible states of a qubit.

The Bell state ($|00\rangle + |11\rangle$) is rotationally invariant with respect to real rotations, but for complete rotational invariance under all complex rotations, we need to consider the s_y - state ($|01\rangle - |10\rangle$).

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: CHSH INEQUALITY**

In the field of Quantum Information, an important concept to understand is Quantum Entanglement. Quantum Entanglement refers to the phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particle, even when they are physically separated. This concept is crucial in various applications of quantum computing and quantum communication.

One way to study Quantum Entanglement is through the CHSH inequality, which is a simplification of Bell's work. The CHSH inequality is demonstrated through a game between two players, referred to as Alice and Bob. In this game, Alice and Bob each receive an input bit, X and Y , respectively, and their task is to output bits A and B , respectively. The challenge is that Alice and Bob are not allowed to communicate with each other during the game.

The inputs, X and Y , are chosen uniformly at random from the set $\{0, 1\}$. The goal of the game is for Alice and Bob to output matching bits, A and B , except when both inputs are 1, in which case they must output non-matching bits. In classical scenarios, the best strategy for Alice and Bob is to target three out of the four possible cases, resulting in a success probability of $3/4$.

However, in the quantum scenario, if Alice and Bob are allowed to share an entangled state, such as a Bell pair, they can potentially achieve a higher success probability. Entanglement is viewed as a resource in quantum computation and quantum information, enabling certain tasks that are not possible classically or can be performed more efficiently. It is important to note that entanglement cannot be used for faster-than-light communication, as stated by the no-signaling theorem.

Instead, entanglement allows Alice and Bob to generate non-local correlations, which can be demonstrated through the CHSH game. In this game, Alice and Bob use their shared entangled state to generate outputs, A and B , that satisfy the condition $x * y = a + b \bmod 2$, or equivalently, $a \oplus b = xy$. By measuring her qubit in one of two bases depending on the value of x , Alice can play the game in such a way that the success probability is cosine squared ($\pi/8$), approximately 0.85.

Quantum Entanglement is a fundamental concept in Quantum Information that allows for the generation of non-local correlations. The CHSH inequality is a demonstration of entanglement's impact on the success probability of a coordination game between two players. Understanding entanglement is crucial for advancements in quantum computing and quantum communication.

Alice and Bob are conducting an experiment to demonstrate quantum entanglement and the violation of the CHSH inequality. They each have a qubit of a Bell state, which is a superposition of two entangled states. Alice chooses to measure her qubit in either the 0-1 basis or the 45-degree rotated basis, depending on the value of x . Similarly, Bob measures his qubit in either the 0-1 basis or the -45-degree rotated basis, depending on the value of y .

To understand why their measurements are relevant, let's visualize it on a circle. If x is 0, Alice measures in the 0-1 basis, and if x is 1, she measures in the 45-degree rotated basis. Similarly, if y is 0, Bob measures in the 0-1 basis, and if y is 1, he measures in the -45-degree rotated basis. The angles of rotation for each basis are $\pi/8$, and these angles are crucial for determining the probabilities of getting the same or different outcomes.

According to the rotational invariance of the Bell state, if Alice measures her qubit in a certain basis and Bob measures his qubit in a basis rotated by θ , the probability of getting the same outcome is cosine squared θ . There are four possible scenarios: $x=0, y=0$; $x=0, y=1$; $x=1, y=0$; and $x=1, y=1$. In each case, the angle between their bases is $\pi/8$, and the probability of getting the same outcome is cosine squared $\pi/8$.

For $x=1, y=1$, the angle between their bases is $3\pi/8$, and the probability of getting different outcomes is 1 minus cosine squared $3\pi/8$, which is equivalent to sine squared $3\pi/8$. Therefore, the chance of meeting the condition in each of the four cases is exactly cosine squared $\pi/8$, indicating that Alice and Bob can succeed with a probability of 0.85.

This result demonstrates that the quantum case allows for a higher success rate than what is possible classically. To perform an experiment based on this concept, Alice and Bob would each have their own apparatus, located far apart from each other. They would create a Bell state and then transport their qubits to their respective apparatus. Random bits x and y would be generated, determining the measurement basis for each qubit. After making the measurements, they would collect the results and repeat the process multiple times to gather statistics.

To test whether x times y is correlated with a plus b modulo 2, they would analyze the correlations and check if they are close to $3/4$, less than $3/4$, or bounded away from $3/4$. Previous experiments have shown that the correlations are bounded away from $3/4$ and consistent with the prediction of quantum mechanics, which is cosine squared $\pi/8$.

It is important to note that these experiments have been conducted multiple times to minimize possible sources of error.

Quantum entanglement is a fundamental concept in the field of quantum information. It refers to the phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particles. This correlation exists even when the particles are separated by large distances.

One important aspect of quantum entanglement is the violation of Bell inequalities, which are mathematical expressions that impose constraints on the correlations that can be observed between entangled particles. The Clauser-Horne-Shimony-Holt (CHSH) inequality is one such Bell inequality.

The CHSH inequality provides a way to test the predictions of quantum mechanics against the principles of local realism. Local realism is the idea that physical properties of objects exist independently of any observation and that these properties can be determined by local measurements. In other words, local realism suggests that there are hidden variables that determine the outcomes of measurements.

However, experiments testing the CHSH inequality have consistently shown violations of the inequality, which implies that local realism is not a valid description of the quantum world. These violations provide strong evidence in favor of the predictions of quantum mechanics.

To ensure the validity of these experimental results, researchers have worked to eliminate various loopholes that could potentially undermine the conclusions. Loopholes are errors or imperfections in the experimental setup that could allow for alternative explanations of the observed correlations.

Some of the loopholes that have been addressed include the detector loophole and the source loophole. The detector loophole refers to imperfections in the detectors used to measure the properties of the entangled particles. The source loophole, on the other hand, relates to imperfections in the entangled particle source.

While individual experiments have successfully eliminated these loopholes individually, no experiment has yet been conducted that eliminates all the loopholes simultaneously. This means that there is still a small chance that these experiments could be consistent with local realism if nature has conspired in a specific way.

However, ongoing efforts are being made to design experiments that can eliminate all the loopholes simultaneously. These experiments, scheduled to be conducted over the next few years, aim to provide even stronger evidence against local realism and further support the predictions of quantum mechanics.

Quantum entanglement and the violation of Bell inequalities, such as the CHSH inequality, provide compelling evidence against the principles of local realism. While loopholes in experimental setups still exist, ongoing research aims to eliminate these loopholes and strengthen the case for the validity of quantum mechanics.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ENTANGLEMENT****TOPIC: BELL AND LOCAL REALISM**

In the study of quantum information, one of the fundamental concepts is quantum entanglement. Quantum entanglement refers to a phenomenon where two or more particles become correlated in such a way that the state of one particle cannot be described independently of the state of the other particles. This concept was famously explored in the context of the Bell inequalities.

The Bell inequalities were derived by physicist John Bell and are used to test the predictions of local realism. Local realism is the idea that physical properties of objects exist independently of measurement and that these properties are determined by hidden variables. In other words, local realism suggests that there is a hidden mechanism behind the scenes that determines the outcomes of measurements.

To understand the Bell inequalities, let's consider the CH SH game played by two players, Alice and Bob. They each have inputs, x and y , and outputs, a and b . The condition they must satisfy is that if $x = y = 1$, then $a \neq b$; otherwise, $a = b$.

To analyze this game, let's consider a specific input scenario where $x = 0$ and $y = 0$. Without loss of generality, we can assume that Alice outputs 0. If Alice and Bob wish to beat the $3/4$ mark, they must be correct in all four input scenarios. This means that Bob is forced to answer 0 in this case.

Similarly, for other input scenarios, we can deduce that Alice and Bob must output matching bits to be correct. This leads to the conclusion that the best they can do is achieve a success probability of $3/4$.

Now, let's explore the concept of local realism more carefully. In the most general situation, we can imagine that Alice and Bob are particles or systems that have been brought together temporarily. In a classical theory that is both local and realistic, these particles would have communicated with each other and stored some information.

This stored information could include instructions on how to react to different experiments and probabilistic scenarios. For example, Alice and Bob could coordinate their choices based on coin tosses or predetermined probabilistic choices. However, once they are far apart, they can no longer communicate with each other.

If Alice and Bob were to beat the $3/4$ bound in this probabilistic scenario, it would contradict the proof we discussed earlier. This is because if their expected outcome beats $3/4$, there must be some setting of the random bits or probabilistic choices under which they achieve a higher success probability.

The concept of quantum entanglement and the Bell inequalities challenge the assumptions of local realism. Quantum entanglement suggests that the state of one particle is intrinsically connected to the state of another particle, regardless of distance. The Bell inequalities provide a way to test the predictions of local realism and demonstrate that certain correlations cannot be explained by hidden variables.

In the study of quantum information, one fundamental concept is quantum entanglement. Quantum entanglement refers to a phenomenon where two or more particles become connected in such a way that the state of one particle cannot be described independently of the state of the other particles. This means that the properties of entangled particles are intrinsically linked, regardless of the distance between them.

One important aspect of quantum entanglement is the violation of Bell's inequality. Bell's inequality is a mathematical expression that sets limits on the correlations between the measurements of entangled particles. According to local realism, a principle that assumes the existence of hidden variables, the correlations between entangled particles should follow certain bounds. However, experiments have shown that these bounds can be violated, indicating that local realism is not a valid description of quantum phenomena.

The CHSH inequality, named after Clauser, Horne, Shimony, and Holt, is a specific form of Bell's inequality that can be used to test the violation of local realism. It involves measuring the correlations between two entangled particles using different measurement settings. If the correlations exceed a certain threshold, it implies that local realism is violated.

To understand this, let's consider an example involving two entangled particles, Alice and Bob. They each have a random bit value, which can be either 0 or 1. Alice and Bob can choose a strategy to measure their respective particles based on their random bit values. The CHSH inequality states that if Alice and Bob have a fixed strategy for their measurements, the probability of obtaining a specific outcome should be greater than $3/4$.

Now, let's assume that Alice and Bob fix their measurement strategies based on their random bit values. When they are far apart and finally meet, they can check if the specific condition mentioned in the CHSH inequality holds with a probability greater than $3/4$. However, experiments have shown that no matter what strategy they choose, they cannot achieve this condition with a probability better than $3/4$. This contradicts the assumption of local realism.

Therefore, the violation of the CHSH inequality implies that nature cannot simultaneously be both local and follow realism. This suggests that while locality may still hold, realism cannot. In other words, the state of quantum systems is undetermined until they are observed. Only when measurements are made, properties such as the value of a bit, the sign, or the position and momentum of a particle emerge.

Quantum entanglement and the violation of Bell's inequality, specifically the CHSH inequality, provide evidence against the validity of local realism in describing quantum phenomena. These concepts highlight the indeterminate nature of quantum systems until measurements are made, challenging our understanding of the fundamental nature of reality.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROCESSING****TOPIC: TIME EVOLUTION OF A QUANTUM SYSTEM**

Good morning. Today, we will discuss the concept of quantum gates, which is an essential aspect of quantum information processing. To understand quantum gates, we need to review the fundamental axioms of quantum mechanics.

The first axiom is the superposition principle, which states that the state of a quantum system is a point on a K -dimensional ball in a K -dimensional complex space. For example, if we have a K -level system like the energy levels of an electron, the state of the system can be represented as a superposition of K basis states. These basis states are labeled from 0 to $K-1$ and correspond to the different energy levels of the electron. The state of the system is a unit vector in this K -dimensional complex vector space, with complex amplitudes representing the coefficients of the superposition.

The second axiom deals with measurements of the quantum system. When we measure the system, we choose an orthonormal basis, which consists of vectors that are mutually perpendicular and have a unit length. The probability of obtaining a specific measurement outcome, say J , is given by the square of the magnitude of the inner product between the measurement basis vector and the state vector of the system. In other words, it is the square of the cosine of the angle between the two vectors. The new state of the system after the measurement is the measurement basis vector corresponding to the obtained outcome.

Now, let's move on to the third axiom, which addresses the time evolution of a quantum system. According to this axiom, the evolution of the system is represented by a rotation of the Hilbert space, which is the complex vector space we discussed earlier. This rotation corresponds to the change in the state of the system over time. In other words, the state vector undergoes a transformation, similar to a spin or rotation in a particular direction within the Hilbert space.

To summarize, quantum gates are a fundamental concept in quantum information processing. They describe how the state of a quantum system evolves over time. This evolution is governed by the principles of superposition, measurement, and time-dependent rotations within the Hilbert space.

When studying quantum information, it is important to understand the time evolution of a quantum system. The state of a quantum system can change over time, and this change can be represented by a rotation in a mathematical space. To better understand this concept, let's consider a two-dimensional space and a quantum state with real coefficients.

In this two-dimensional space, we have a standard basis consisting of the states 0 and 1. To evolve the system, we give a rotation to this space. As a result, the state 0 moves to a new position, denoted as $U(0)$, and the state 1 rotates correspondingly. It is important to note that angles between vectors are preserved during this rotation.

Formally, the rotation of a space is given by a linear transformation, which can be represented by a matrix. Let's consider an example in two dimensions. If we rotate the state vector through an angle θ , the state 0 moves to a new state with coordinates $(\cos \theta, \sin \theta)$, and the state 1 moves to $(\cos \theta, -\sin \theta)$.

This rotation can be described by a linear transformation matrix, $R(\theta)$, whose columns are the vectors $(\cos \theta, \sin \theta)$ and $(-\sin \theta, \cos \theta)$. It is worth mentioning that there is also a transformation for a rotation through $-\theta$, represented by the transpose of $R(\theta)$.

It is interesting to note that $R(\theta)$ and $R(-\theta)$ satisfy the relation $R(\theta) * R(-\theta) = R(-\theta) * R(\theta) = I$, where I is the identity matrix. This means that if you rotate through θ and then rotate through $-\theta$, you will come back to the original state.

These linear transformations that represent rotations in a vector space are called unitary transformations. In the next material, we will study unitary transformations more formally and explore their properties.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROCESSING****TOPIC: UNITARY TRANSFORMS**

Unitary transformations are linear transformations that describe rotations in a complex vector space. In the previous video, we discussed rotations in a two-dimensional real space using a 2×2 matrix called $R(\theta)$, which was defined in terms of sines and cosines. We also looked at rotations through $-\theta$, which was given by the transpose of $R(\theta)$. It was noted that rotating through $-\theta$ undoes the effect of rotating through θ , resulting in $R(\theta)$ times $R(\theta)$ transpose being equal to the identity matrix.

Now, let's describe general unitary transformations in a K -dimensional complex vector space. The linear transformation, or rotation matrix, will be a $K \times K$ matrix with complex entries. A unitary transformation, denoted as U , is a rotation of the vector space if U conjugate transpose times U is equal to U times U conjugate transpose, which is the identity matrix. In other words, U is unitary if and only if this condition is satisfied.

To illustrate this, let's consider a 2×2 matrix U with complex entries a , b , c , and d . The conjugate transpose of U , denoted as $U^\dagger$, is obtained by taking the complex conjugates of the entries and then transposing the matrix. The condition $U^\dagger U$ is the identity matrix can be expressed as $\bar{a} \bar{b} \bar{c} \bar{d}$ times $a \ b \ c \ d$ equals $1 \ 0 \ 0 \ 0$.

Interpreting this condition, we can see that U represents the transformation of the vector space. The 0 state, represented by the vector $[a \ b]$, is mapped to the first column of U , which is $[a \ 0 + b \ 1]$. Similarly, the 1 state, represented by the vector $[c \ d]$, is mapped to the second column of U . The inner product between these two vectors is 0 , indicating that they are orthogonal to each other. Furthermore, the length of these vectors is 1 , as given by the inner product of $a \ b$ with itself and $c \ d$ with itself.

For a general $K \times K$ unitary matrix, the 0 state is mapped to the first column, the 1 state is mapped to the second column, and the $(K-1)$ state is mapped to the K th column. The condition $U^\dagger U$ is equal to the identity matrix ensures that the inner product of each vector with itself is 1 , indicating that they are unit vectors. Additionally, the inner product between different vectors is 0 , indicating orthogonality.

Unitary transformations are linear transformations that describe rotations in a complex vector space. They can be represented by $K \times K$ matrices with complex entries. A unitary transformation is a rotation if it satisfies the condition $U^\dagger U$ is equal to the identity matrix. This condition ensures that the transformation maps the 0 state and the 1 state to orthogonal unit vectors.

A unitary transform is a type of transformation in quantum information processing that preserves the inner products and angles between vectors. When applying a unitary transform, the states 0 through $K-1$ get mapped to orthogonal states, which are also normalized. This means that the length of each column in the transformed matrix is 1 , and the inner product between the transformed vectors remains the same.

To understand the relationship between the original and transformed vectors, we can look at the inner product between them. If the indices of the vectors are not equal, the inner product is 0 , indicating that the columns are orthogonal to each other. This property holds for all entries except when i is equal to j , in which case the inner product is non-zero.

To demonstrate that a unitary transform preserves inner products, let's consider two different states, Φ and Ψ . When we apply the unitary transform U to each state, the claim is that the inner product between Φ and Ψ remains the same as the inner product between $U \Phi$ and $U \Psi$.

To prove this, we can calculate the inner product between the two sets of vectors. The inner product between Φ and Ψ is given by the conjugate transpose of Φ multiplied by Ψ . Similarly, the inner product between $U \Phi$ and $U \Psi$ is given by the conjugate transpose of $U \Phi$ multiplied by $U \Psi$.

Since U is a matrix, $U \Phi$ and $U \Psi$ are vectors. To obtain the conjugate transpose of a vector, we take the conjugate of each entry and transpose the vector. This results in a row vector.

By applying the rules of matrix transpose, we can rewrite the inner product between $U \Phi$ and $U \Psi$ as Φ

multiplied by the conjugate transpose of U multiplied by $U \Psi$. Since U multiplied by its conjugate transpose is the identity matrix, the inner product simplifies to Φ multiplied by Ψ .

Therefore, the inner product between Φ and Ψ is equal to the inner product between $U \Phi$ and $U \Psi$, demonstrating that a unitary transform preserves inner products.

A unitary transform in quantum information processing preserves the inner products and angles between vectors. This property ensures that the transformed vectors maintain their relationship with each other, allowing for accurate calculations and analysis.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROCESSING****TOPIC: SINGLE QUBIT GATES**

Quantum gates are fundamental building blocks in quantum information processing. In classical computing, we have gates that operate on bits of information, such as the NOT gate and the AND gate. In quantum computing, we have quantum gates that operate on qubits, which are the basic units of quantum information.

A quantum gate takes a qubit as input and performs a unitary transformation on it, resulting in a new state for the qubit. The input qubit is represented by a wire, and the output qubit is also represented by a wire. The transformation performed by the gate is a unitary transformation, meaning it preserves the norm of the qubit and is reversible.

Let's look at some examples of single qubit quantum gates. The first example is the bit flip gate, represented by the transformation matrix X:

1.	X =	0 1
2.		1 0

The bit flip gate flips the basis state of the qubit. For example, if the input qubit is in the state $|0\rangle$, the bit flip gate maps it to the state $|1\rangle$, and vice versa. If the input qubit is in a superposition of $|0\rangle$ and $|1\rangle$, the bit flip gate flips the amplitudes of the basis states.

To be considered a quantum gate, the bit flip gate must be a unitary transformation. We can verify this by checking if the product of the gate and its conjugate transpose (also known as the adjoint or dagger) is equal to the identity matrix. In the case of the bit flip gate, X and X dagger are the same matrix, and their product is indeed the identity matrix.

The second example is the phase flip gate, represented by the transformation matrix Z:

1.	Z =	1 0
2.		0 -1

The phase flip gate leaves the basis state $|0\rangle$ unchanged, but it introduces a phase change of -1 to the basis state $|1\rangle$. This means that if the input qubit is in a superposition of $|0\rangle$ and $|1\rangle$, the phase flip gate introduces a phase change to the amplitudes of the basis states.

Similar to the bit flip gate, the phase flip gate must also be a unitary transformation. Again, we can verify this by checking if the product of the gate and its conjugate transpose is equal to the identity matrix. In the case of the phase flip gate, Z and Z dagger are the same matrix, and their product is indeed the identity matrix.

The third example is the Hadamard gate, represented by the transformation matrix H:

1.	H = $1/\sqrt{2}$ *	1 1
2.		1 -1

The Hadamard gate is a particularly important gate in quantum computing. It transforms the basis states $|0\rangle$ and $|1\rangle$ into superpositions of those states. Specifically, the Hadamard gate maps $|0\rangle$ to $(|0\rangle + |1\rangle)/\sqrt{2}$ and $|1\rangle$ to $(|0\rangle - |1\rangle)/\sqrt{2}$.

Similar to the previous gates, the Hadamard gate must also be a unitary transformation. We can verify this by checking if the product of the gate and its conjugate transpose is equal to the identity matrix. In the case of the Hadamard gate, H and H dagger are the same matrix, and their product is indeed the identity matrix.

Quantum gates are essential tools in quantum information processing. They operate on qubits and perform unitary transformations on them. We have seen examples of single qubit gates, such as the bit flip gate, the phase flip gate, and the Hadamard gate. These gates have specific transformation matrices that determine how they affect the qubits. It is important to note that all quantum gates must be unitary transformations to

preserve the norm of the qubit and be reversible.

In the context of quantum information processing, single qubit gates play a crucial role in manipulating the states of individual qubits. One important single qubit gate is the Hadamard gate, denoted as H . The Hadamard gate transforms the basis states 0 and 1 into superposition states known as the plus state and the minus state, respectively.

When the Hadamard gate is applied to the state 0, it maps it to the plus state, which is represented by the vector $[1/\sqrt{2}, 1/\sqrt{2}]$. Similarly, when the Hadamard gate is applied to the state 1, it maps it to the minus state, which is represented by the vector $[1/\sqrt{2}, -1/\sqrt{2}]$. These vectors correspond to the columns of the matrix representing the Hadamard gate.

To ensure that the Hadamard gate is a unitary transformation, we need to verify that the product of the gate and its conjugate transpose, denoted as $H^\dagger H$, is equal to the identity matrix. Since the Hadamard gate is real and symmetric, its conjugate transpose is equal to the gate itself. Thus, $H^\dagger H$ is equal to HH , which is equivalent to H^2 . By checking that H^2 is equal to the identity matrix, we confirm that the Hadamard gate is indeed unitary.

It is noteworthy that the square of the Hadamard gate being the identity is a property shared by other elementary gates as well. This means that if a gate is applied twice, it brings the qubit back to its original state. For example, in the case of the bit flip gate (denoted as X), flipping a bit twice results in the qubit returning to its initial state.

In the case of the Hadamard gate, applying it twice to a qubit results in a change of basis from the standard basis (0 and 1) to the plus and minus basis. This change of basis allows for the encoding of information about the initial state in the phase of the resulting superposition state. By applying another Hadamard gate, the phase information can be recovered and translated back into the original bit information.

From a geometric perspective, the Hadamard gate can be visualized as a rotation about a specific axis. It maps the basis states 0 and 1 to the plus and minus states, respectively, by rotating them around the rotation axis. This rotation axis is located at an angle of $\pi/8$ from the z -axis.

Interestingly, there is a relationship between the Hadamard gate and other single qubit gates, namely the bit flip gate (X) and the phase flip gate (Z). The Hadamard gate swaps the basis states 0 and 1 with the plus and minus states, respectively. On the other hand, the bit flip gate swaps the states 0 and 1, while the phase flip gate swaps the plus and minus states. Therefore, it can be observed that applying a bit flip gate followed by a phase flip gate is equivalent to applying a Hadamard gate.

The Hadamard gate is a fundamental single qubit gate in quantum information processing. It transforms the basis states 0 and 1 into superposition states known as the plus and minus states, respectively. By applying the Hadamard gate twice, the phase information can be encoded and recovered. Additionally, the Hadamard gate has a relationship with other single qubit gates, allowing for the manipulation of qubit states.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROCESSING****TOPIC: TWO QUBIT GATES**

In this didactic material, we will discuss the concept of two qubit gates in quantum information processing. Before we delve into two qubit gates, let's briefly review single qubit gates.

A single qubit gate is a unitary transformation that operates on a single qubit. It takes an input qubit in a certain state and outputs a transformed qubit in a different state. The unitary transformation is represented by a 2×2 complex matrix, denoted as U , with entries A , B , C , and D . The unitary transformation satisfies the condition $U^\dagger U = U U^\dagger = I$, where $U^\dagger$ is the conjugate transpose of U and I is the identity matrix. The transformation of the input qubit is achieved by multiplying the input vector, $[\alpha_0, \alpha_1]$, with the matrix U , resulting in the output vector $[\alpha'_0, \alpha'_1]$. This describes the change in the state of the qubit.

Now, let's move on to two qubit gates. A two qubit gate operates on a pair of qubits and performs a transformation on their joint state. The input consists of two qubits, and the state of the two qubits can be represented as a superposition of all four possibilities, with complex amplitudes α_{00} through α_{11} . The output is also a superposition with different amplitudes α'_{00} through α'_{11} .

Similar to single qubit gates, the two qubit gate is represented by a 4×4 complex matrix, denoted as U , with entries A , B , C , D , and so on. The unitary property of the two qubit gate requires that $U^\dagger U = U U^\dagger = I$. The transformation of the input qubits is achieved by multiplying the input vector, $[\alpha_{00}, \alpha_{01}, \alpha_{10}, \alpha_{11}]$, with the matrix U . Each column of the matrix corresponds to a specific input state, and the resulting output state is a linear combination of these columns.

Let's consider an example of a commonly used two qubit gate called the CNOT gate. The CNOT gate has a control qubit and a target qubit. It leaves the control qubit unchanged and flips the target qubit if and only if the control qubit is 1. In the basis states 0 and 1, the CNOT gate maps 00 to 00, 01 to 01, 10 to 11, and 11 to 10. For a general input state, the CNOT gate performs the corresponding transformation based on the control and target qubits.

It is important to note that the CNOT gate is unitary, satisfying the condition $\text{CNOT}^\dagger \text{CNOT} = \text{CNOT} \text{CNOT}^\dagger = I$. This can be verified by applying the CNOT gate twice, which results in the identity transformation.

Apart from the CNOT gate, there are other ways to create unitary transformations on two qubits. One approach is to apply single qubit gates to each qubit individually. For example, applying a Z gate to the first qubit and a Hadamard gate to the second qubit. The resulting transformation can be considered as a two qubit gate, denoted as U . The specific form of the two qubit gate U can be determined by the combination of single qubit gates applied.

Two qubit gates are unitary transformations that operate on pairs of qubits. They can be represented by complex matrices and perform transformations on the joint state of the input qubits. The CNOT gate is a commonly used example of a two qubit gate, while other two qubit gates can be constructed by combining single qubit gates applied to each qubit individually.

In the study of quantum information, it is important to understand the concept of two-qubit gates and how they can be represented mathematically. One way to represent a two-qubit gate is through a 4×4 linear transformation known as a unitary transformation. This unitary transformation describes the combined effect of applying individual transformations on each qubit.

To better understand this, let's consider two qubits, labeled as qubit 1 and qubit 2. Let's say we apply a transformation, denoted as U_1 , on qubit 1, and another transformation, denoted as U_2 , on qubit 2. We can represent U_1 as a 2×2 matrix with elements a , b , c , and d , and U_2 as a 2×2 matrix with elements e , f , g , and h .

The question now is, what is the 4×4 unitary transformation that describes the combined effect on both qubits? The answer is quite elegant. We can imagine that the four numbers in the resulting transformation matrix will be scaled by a factor, let's call it "a". For example, the element in the top left corner would be a times e , the element in the top right corner would be a times f , the element in the bottom left corner would be a times g ,

and the element in the bottom right corner would be a times h.

To understand this concept further, let's examine how the rows and columns are numbered. The rows and columns can be labeled as 0 or 1, corresponding to the states of the qubits. For example, the top left element corresponds to the state 00, the top right element corresponds to the state 01, the bottom left element corresponds to the state 10, and the bottom right element corresponds to the state 11.

If we consider the first qubit alone, it's as though we have a 2x2 matrix with elements a, b, c, and d. Now, if we fix the first qubit to be 0, we can observe that the matrix on the second qubit shows up as e, f, g, and h. We can reproduce this matrix four times and multiply it by the corresponding entries.

To illustrate this, let's consider the case where the input is 00. If we apply U1 and U2 to this input, we get $a0 + b1$ times $e0 + f1$, which simplifies to ae . Similarly, if we consider the input 01, we get $a0 + b1$ times $g0 + h1$, which simplifies to af . By following this pattern, we can see that the resulting transformation matrix is:

$$\begin{pmatrix} 1/\sqrt{2} & 1/\sqrt{2} & 1/\sqrt{2} & -1/\sqrt{2} \\ 1/\sqrt{2} & -1/\sqrt{2} & -1/\sqrt{2} & 1/\sqrt{2} \end{pmatrix}$$

This matrix represents the unitary transformation for the given values of U1 and U2.

A two-qubit gate can be represented by a 4x4 unitary transformation matrix. The elements of this matrix are obtained by scaling the corresponding elements of the individual transformation matrices applied to each qubit. By understanding the formalism behind this representation, we can better analyze and manipulate quantum information.

In quantum information processing, the fundamental building blocks are qubits, which are analogous to classical bits. However, unlike classical bits, qubits can exist in a superposition of states. For example, a single qubit can be in a superposition of ground and excited states, represented by the state $\beta_{00} + \beta_{11}$, where β_0 and β_1 are complex numbers. This superposition state lives in a two-dimensional complex vector space called Hilbert space H_1 .

When we have multiple qubits, such as two qubits, we need to consider the combined state of the system. The combined state lives in a four-dimensional complex vector space, denoted as H_2 . The amplitudes of the combined state can be written as $\alpha_{0000} + \alpha_{0101} + \alpha_{1010} + \alpha_{1111}$, where α_{00} , α_{01} , α_{10} , and α_{11} are complex numbers. This four-dimensional vector space is obtained by taking the tensor product of the individual qubit Hilbert spaces.

The tensor product operation is used to combine vector spaces. In the case of quantum information processing, we take the tensor product of the Hilbert space H_1 and H_2 to obtain the vector space H . The dimensions of the vector spaces multiply, resulting in a four-dimensional vector space. Formally, if we have a vector u in H_1 and a vector v in H_2 , we can write down the vector $u \otimes v$ in H . The tensor product operation satisfies linearity, meaning that $(u_1 + u_2) \otimes v = u_1 \otimes v + u_2 \otimes v$ and $u \otimes (v_1 + v_2) = u \otimes v_1 + u \otimes v_2$.

In the case of basis vectors 0 and 1, we can represent the tensor product as follows:

- $0 \otimes 0 = 00$ (also denoted as 00 or 00 in cat)
- $0 \otimes 1 = 01$
- $1 \otimes 0 = 10$
- $1 \otimes 1 = 11$

By taking linear combinations of these product vectors, we can obtain a variety of vectors in H . Some of these vectors, known as entangled states, cannot be written as a product of vectors in the individual qubit spaces.

The tensor product operation also extends to unitary transformations. If we have two unitary transformations u_1 and u_2 applied to each of the qubits, the resulting unitary transformation for the two-qubit system is given by $u = u_1 \otimes u_2$. This means that the combined transformation u acts on the combined state of the two qubits.

In terms of inner products, the tensor product inherits the inner product from the individual Hilbert spaces. For elementary tensors u_1v_1 and u_2v_2 , the inner product between these two is equal to the inner product between u_1 and u_2 multiplied by the inner product between v_1 and v_2 . Once the inner product for elementary tensors is

defined, it can be extended to any vector in H by linearity.

To summarize, in quantum information processing, the combination of qubits involves taking the tensor product of the individual qubit Hilbert spaces. This results in a higher-dimensional vector space where the states and transformations of the combined system can be described. The tensor product operation satisfies linearity and inherits the inner product from the individual Hilbert spaces.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROPERTIES****TOPIC: NO-CLONING THEOREM**

Quantum teleportation is a fascinating concept in quantum information. In this lecture, we will explore the topic of quantum teleportation and its implications. Before we dive into that, let's first discuss the no-cloning theorem, which is closely related to quantum teleportation.

The no-cloning theorem addresses the question of whether it is possible to make an exact copy of an unknown quantum state. To illustrate this, let's consider a quantum state represented by a single qubit, denoted as $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$. Here, α and β are complex numbers, and $|0\rangle$ and $|1\rangle$ represent the basis states of the qubit.

Now, suppose we have another qubit in a known state, let's say $|0\rangle$. The question is, can we perform a unitary transformation on these two qubits to achieve a state where both qubits are in the state $|\psi\rangle$?

Formally, we want to find a unitary transformation U such that $U(|\psi\rangle \otimes |0\rangle) = |\psi\rangle \otimes |\psi\rangle$, where $\otimes$ denotes the tensor product. The no-cloning theorem tells us that such a unitary transformation does not exist.

To understand why, let's examine the argument. If the transformation U works for all possible values of α and β , it must work when $\alpha = 1$ and $\beta = 0$, which corresponds to the state $|0\rangle$. In this case, we expect $U(|0\rangle \otimes |0\rangle) = |0\rangle \otimes |0\rangle$.

Similarly, if we consider the state $|1\rangle$ ($\alpha = 0, \beta = 1$), we expect $U(|1\rangle \otimes |0\rangle) = |1\rangle \otimes |1\rangle$.

Now, since U is a linear map, we can apply it to a linear combination of states. So, we can write $U(|\psi\rangle \otimes |0\rangle) = \alpha U(|0\rangle \otimes |0\rangle) + \beta U(|1\rangle \otimes |0\rangle)$.

Expanding this expression, we get $\alpha U(|0\rangle \otimes |0\rangle) + \beta U(|1\rangle \otimes |0\rangle) = \alpha|0\rangle \otimes |0\rangle + \beta|1\rangle \otimes |1\rangle$.

However, this is not the same as $|\psi\rangle \otimes |\psi\rangle$, which is $\alpha^2|0\rangle \otimes |0\rangle + \alpha\beta|0\rangle \otimes |1\rangle + \alpha\beta|1\rangle \otimes |0\rangle + \beta^2|1\rangle \otimes |1\rangle$.

Therefore, we have a contradiction. It is not possible to clone an unknown quantum state perfectly.

The no-cloning theorem has profound implications in quantum information. It implies that quantum information cannot be copied or cloned without altering its state. This property is crucial for various applications in quantum cryptography and quantum computing.

The no-cloning theorem states that it is impossible to make an exact copy of an unknown quantum state. This theorem plays a fundamental role in understanding the unique properties of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROPERTIES****TOPIC: BELL STATE CIRCUIT**

Quantum teleportation is a protocol that allows for the transfer of quantum information from one location to another, even if there is no direct quantum channel between the two locations. The process involves the use of entangled particles, specifically qubits, and measurements.

Let's consider an example scenario where Alice wants to transport a qubit to Bob's lab. Alice has a qubit in a superposition state, represented as $\alpha|0\rangle + \beta|1\rangle$, where α and β are unknown complex numbers. Alice and Bob share a special entangled state called a Bell state, which is represented as $|00\rangle + |11\rangle$.

In the quantum teleportation protocol, Alice performs a measurement on her qubit and the shared Bell state. This measurement yields two classical bits, b_1 and b_2 . Alice then communicates these two bits to Bob, for example, over a phone call. Bob uses the received classical information to perform one of four different quantum gates on his qubit. The choice of the gate depends on the values of b_1 and b_2 . After applying the gate, Bob's qubit is guaranteed to be in the state $\alpha|0\rangle + \beta|1\rangle$, which is the same as the original unknown qubit.

The teleportation process may seem like magic, as Alice's qubit is destroyed in the process, yet Bob is able to reconstruct it. The protocol relies on the entanglement of the shared Bell state and the use of measurements and quantum gates.

To better understand the quantum teleportation protocol, let's break it down into two steps. In the first step, let's assume an unrealistic scenario where Bob's qubit is initially in the state $|0\rangle$, and there is a CNOT gate connecting Alice's and Bob's labs. Alice applies the CNOT gate using her qubit as the control bit. The resulting state is $\alpha|00\rangle + \beta|10\rangle$.

In the second step, we want Bob's qubit to end up in the state $\alpha|0\rangle + \beta|1\rangle$. To achieve this, Alice needs to perform a measurement on her qubit. If she measures in the 0-1 basis, the resulting states for the different measurement outcomes are not desirable for Bob. Instead, Alice can measure her qubit in the plus/minus basis.

By rewriting the joint state of the two qubits in the plus/minus basis, we find that the state becomes $\alpha(1/\sqrt{2})(|0\rangle + |1\rangle) \otimes |0\rangle + \beta(1/\sqrt{2})(|0\rangle - |1\rangle) \otimes |1\rangle$.

In this new representation, Alice can measure her qubit and obtain one of the four possible outcomes: $|0+\rangle$, $|0-\rangle$, $|1+\rangle$, or $|1-\rangle$. She communicates this outcome to Bob, who applies the corresponding quantum gate to his qubit. After this gate operation, Bob's qubit will be in the desired state $\alpha|0\rangle + \beta|1\rangle$.

The quantum teleportation protocol allows for the transfer of quantum information without physically moving the qubit itself. It relies on the principles of entanglement, measurement, and quantum gates to achieve this feat.

In the previous material, we discussed the concept of creating an entangled state known as the Bell state. We explored how Alice can transmit her qubit to Bob using this state, assuming there is a quantum gate connecting their labs. However, in reality, such a gate does not exist. In this didactic material, we will focus on understanding how to overcome this challenge and create the entangled state without direct quantum communication or the presence of a gate between Alice and Bob's labs.

To recap, the Bell state is represented by the equation:

$$|\Psi\rangle = 1/\sqrt{2} (|0\rangle \otimes |0\rangle + |1\rangle \otimes |1\rangle)$$

Where $|0\rangle$ and $|1\rangle$ are the basis states of a qubit. We observed that if Alice measures her qubit in the plus/minus basis, two scenarios arise. If she measures "plus," the new state becomes $|\Psi\rangle = |+\rangle \otimes |0\rangle + |-\rangle \otimes |1\rangle$. On the other hand, if she measures "minus," the new state becomes $|\Psi\rangle = |+\rangle \otimes |1\rangle - |-\rangle \otimes |0\rangle$.

In the first case, we notice that the state of the second qubit (Bob's qubit) is a tensor product state. The first qubit is in the state $|+\rangle$, and the second qubit is in the state $|0\rangle$. This is precisely the desired entangled state,

$|\Psi\rangle$.

However, in the second case, if Alice measures "minus," Bob's qubit is in the state $|+\rangle\otimes|1\rangle - |-\rangle\otimes|0\rangle$, which is not the same as $|\Psi\rangle$. Nevertheless, we can convert it into $|\Psi\rangle$ by applying a certain gate. In this case, we can use the phase flip gate, Z . By applying Z to the state $|+\rangle\otimes|1\rangle - |-\rangle\otimes|0\rangle$, we obtain $|\Psi\rangle$.

To summarize the process, Alice measures her qubit in the plus/minus basis. If the result is "plus," she transmits the value 0 to Bob. If the result is "minus," she transmits the value 1. Bob, upon receiving the value, checks if it is 0 or 1. If it is 0, he leaves his qubit as it is since it is already in the state $|\Psi\rangle$. However, if he receives a 1, he applies the phase flip gate, Z , to his qubit, transforming it into $|\Psi\rangle$.

This demonstrates how Alice can transmit her qubit to Bob without the presence of a quantum gate or direct quantum communication between their labs. It is important to note that the assumption of a gate or direct communication is unrealistic, and in the next material, we will explore alternative methods to create the entangled state $|\Psi\rangle$.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROPERTIES****TOPIC: QUANTUM TELEPORTATION**

In this material, we will discuss a simple quantum circuit that creates a Bell state. The circuit operates on two qubits, each initialized to the state 0. First, a Hadamard gate is applied to the first qubit, followed by a CNOT gate with the first qubit as the control bit and the second qubit as the target bit.

After the Hadamard gate, the state of the first qubit is given by $\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$, while the second qubit remains in the state $|0\rangle$. When the CNOT gate is applied, if the control bit is 0, the second qubit remains unchanged, resulting in the state $\frac{1}{\sqrt{2}}|0\rangle|0\rangle$. However, if the control bit is 1, the target qubit gets flipped, resulting in the state $\frac{1}{\sqrt{2}}|1\rangle|1\rangle$. This final state is known as the Phi plus state, which is a type of Bell state.

If the input qubits are initialized to 0 and 1, the state after the Hadamard gate is still $\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$, but the second qubit is now in the state $|1\rangle$. When the CNOT gate is applied, it only flips the second qubit, resulting in the state $\frac{1}{\sqrt{2}}|0\rangle|1\rangle + \frac{1}{\sqrt{2}}|1\rangle|1\rangle$. This state is called the Psi plus state, which is another type of Bell state.

Similarly, if the input qubits are initialized to 1 and 0, the state after the Hadamard gate is $\frac{1}{\sqrt{2}}|0\rangle - \frac{1}{\sqrt{2}}|1\rangle$, with the second qubit in the state $|0\rangle$. Applying the CNOT gate flips the second qubit, resulting in the state $\frac{1}{\sqrt{2}}|0\rangle|0\rangle - \frac{1}{\sqrt{2}}|1\rangle|1\rangle$. This state is known as the Phi minus state, another type of Bell state.

Lastly, if both input qubits are initialized to 1, the state after the Hadamard gate is $\frac{1}{\sqrt{2}}|0\rangle|1\rangle - \frac{1}{\sqrt{2}}|1\rangle|0\rangle$. Applying the CNOT gate flips the target qubit, resulting in the state $\frac{1}{\sqrt{2}}|0\rangle|1\rangle - \frac{1}{\sqrt{2}}|1\rangle|0\rangle$. This state is called the Psi minus state, which is the singlet state and occurs frequently in nature.

These four states, Phi plus, Psi plus, Phi minus, and Psi minus, are known as the Bell basis states. They form an orthonormal basis for the two-qubit complex vector space. This means that their inner products are zero, or equivalently, they are orthogonal. The orthonormality of the Bell basis states can be demonstrated by computing their inner products or by realizing that the input states used to create these Bell states are the standard orthonormal basis for a four-dimensional vector space. Since the gates used in the circuit are unitary transformations, the output states must also be orthogonal.

This simple quantum circuit with two gates can create the four Bell basis states, which are important in quantum information processing and quantum teleportation.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROPERTIES****TOPIC: QUANTUM TELEPORTATION USING CNOT**

In this didactic material, we will discuss the concept of quantum teleportation using the CNOT gate. Quantum teleportation is a process that allows the transfer of quantum information from one qubit to another without physically moving the qubit itself. We will explore the steps involved in the teleportation protocol and understand how it can be achieved without using a CNOT gate.

To begin, let's recap the initial setup. We have two parties, Alice and Bob, who share an entangled state called a Bell state. Alice also possesses a qubit that she wants to teleport to Bob. The qubit is represented by the state $\alpha|0\rangle + \beta|1\rangle$.

The first step in the teleportation protocol is for Alice to apply a CNOT gate from her qubit to her share of the Bell state. This is done because Alice cannot directly apply a CNOT gate from her lab to Bob's lab. By applying the CNOT gate, the state of the three qubits, including Alice's qubit and the two qubits in the Bell state, is transformed.

Next, let's analyze the state of the three qubits after the CNOT gate is applied. Before the gate, the state was $\alpha|0\rangle \otimes (1/\sqrt{2})(|00\rangle + |11\rangle)$. After the gate, the state becomes $\alpha/\sqrt{2}|000\rangle + \alpha/\sqrt{2}|011\rangle + \beta/\sqrt{2}|100\rangle + \beta/\sqrt{2}|111\rangle$.

Now, Alice measures the second qubit (the middle qubit) and obtains an outcome of either 0 or 1. If she measures 0, the state of the first and third qubits becomes $\alpha|00\rangle + \beta|11\rangle$. If she measures 1, the state becomes $\alpha|01\rangle + \beta|10\rangle$.

Alice then communicates the measurement outcome to Bob. If the outcome is 0, Bob leaves his qubit unchanged, and Alice and Bob now share the state $\alpha|00\rangle + \beta|11\rangle$. This is the desired state, and the teleportation protocol is complete.

If the outcome is 1, Bob applies a bit flip operation to his qubit. This operation transforms the state $\alpha|01\rangle + \beta|10\rangle$ to $\alpha|00\rangle + \beta|11\rangle$, which is the desired state. Again, the teleportation protocol is complete.

To summarize, the quantum teleportation protocol involves the following steps:

1. Alice applies a CNOT gate from her qubit to her share of the Bell state.
2. The state of the three qubits is transformed.
3. Alice measures the second qubit and obtains an outcome of 0 or 1.
4. Alice communicates the measurement outcome to Bob.
5. If the outcome is 0, Bob leaves his qubit unchanged, and the teleportation is complete.
6. If the outcome is 1, Bob applies a bit flip operation to his qubit, and the teleportation is complete.

This protocol allows the teleportation of quantum information from one qubit to another without physically moving the qubit itself. It relies on the principles of entanglement and measurement to achieve the desired result.

In quantum information, one of the most fascinating phenomena is quantum teleportation. Quantum teleportation allows the transfer of quantum states from one location to another, without physically moving the particles involved. This process relies on the principles of entanglement and measurement.

To understand quantum teleportation, let's consider a scenario involving two parties, Alice and Bob. Alice has a qubit in an unknown state, which she wants to teleport to Bob. The process begins with Alice and Bob sharing an entangled pair of qubits.

The first step in the teleportation process is for Alice to apply a CNOT gate to her qubit, using the entangled pair as the control qubit. This gate is a two-qubit gate that flips the second qubit if the first qubit is in the state $|1\rangle$. In this case, Alice's qubit acts as the control qubit, and the entangled pair acts as the target qubit.

Next, Alice performs a measurement on her qubit in the Hadamard basis, also known as the plus-minus basis. This basis is defined by the states $|+\rangle$ and $|-\rangle$, which are superpositions of the classical basis states $|0\rangle$ and $|1\rangle$.

$|1\rangle$. The Hadamard transform is applied to her qubit before the measurement, which is equivalent to measuring in the standard basis.

Alice then communicates the measurement result to Bob. If the measurement outcome is $|1\rangle$, Bob applies a phase flip to his qubit. This phase flip changes the sign of the qubit's state, effectively teleporting the unknown state from Alice's qubit to Bob's qubit.

The reason this process works is due to the entanglement between the qubits shared by Alice and Bob. The CNOT gate followed by the measurement collapses the entangled pair into one of two possible states: either $|00\rangle + |11\rangle$ or $|01\rangle + |10\rangle$. Depending on the measurement outcome, Bob applies a bit flip or a phase flip to his qubit, resulting in the teleportation of the unknown state.

It's important to note that the complex numbers α and β , which specify the unknown state of Alice's qubit, are not explicitly communicated to Bob. This is where the power of entanglement comes into play. The entanglement creates a channel through which these complex numbers implicitly make their way to Bob, allowing for the successful teleportation of the state.

Quantum teleportation is a remarkable process that allows the transfer of quantum states without physically moving particles. It relies on the principles of entanglement and measurement to achieve this feat. By sharing an entangled pair of qubits and performing specific operations, such as CNOT gates and measurements, the unknown state can be teleported from one location to another.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM INFORMATION PROPERTIES****TOPIC: QUANTUM MEASUREMENT**

A measurement in quantum information refers to the process of determining the state of a qubit. The state of a qubit is represented by a unit vector in a complex vector space, with two complex amplitudes for 0 and 1. When a measurement is performed on a qubit, one of the two possibilities, 0 or 1, appears and the state of the qubit changes as a result.

The measurement process is still a topic of mystery and different interpretations have been proposed, such as the Copenhagen interpretation. However, in this material, we will focus on one way to think about measurements.

Imagine a qubit in the state $\alpha|0\rangle + \beta|1\rangle$. To measure this qubit, a measuring apparatus is used. The output of the measurement is 0 with a probability of $|\alpha|^2$ and 1 with a probability of $|\beta|^2$. The question arises: how does the needle in the measuring apparatus point to either 0 or 1 with certain probabilities?

One way to think about this is by entangling the qubit with the needle. We can think of the needle having two states: needle pointing at 0 ($|n0\rangle$) and needle pointing at 1 ($|n1\rangle$). By entangling the qubit with the needle, the combined state becomes $\alpha|0\rangle\otimes|n0\rangle + \beta|1\rangle\otimes|n1\rangle$.

However, macroscopic objects, like the needle, cannot maintain a superposition of states. So, something mysterious happens and the superposition collapses. The needle ends up in either the state $|n0\rangle$ or $|n1\rangle$, with probabilities $|\alpha|^2$ and $|\beta|^2$ respectively.

To understand how the qubit gets entangled with the needle, let's consider another qubit in the state $|0\rangle$. We can entangle these two qubits using a CNOT gate, resulting in the state $\alpha|00\rangle + \beta|11\rangle$. If we have two more qubits in the state $|0\rangle$, we can further entangle them by applying CNOT gates from the first qubit to the third and from the second qubit to the fourth. This leads to the state $\alpha|0000\rangle + \beta|1110\rangle$.

By continuing this entanglement process with more qubits, we can create a cat state, where all qubits are in a superposition of 0 and 1. For example, if we double the number of qubits and perform CNOT gates accordingly, we would end up with a state $\alpha|000\dots0\rangle + \beta|111\dots1\rangle$, where the number of qubits is a macroscopic number.

A measurement in quantum information involves entangling the qubit with a measuring apparatus, such as a needle, and collapsing the superposition to a definite outcome. The entanglement process can be understood by applying CNOT gates between qubits. However, the exact nature of measurements in quantum information is still a topic of ongoing research and interpretation.

Quantum information is a fascinating field that explores the fundamental properties and measurement of quantum systems. One intriguing aspect is the phenomenon of superposition, where a quantum system can exist in multiple states simultaneously. However, when it comes to macroscopic objects, nature seems to dislike superpositions.

The collapse of a superposition is still a mystery. We don't fully understand the exact mechanism behind it. However, what we do know is that somewhere along the way, the superposition collapses, resulting in a measurement outcome. This collapse can be observed, for example, in devices like photomultipliers used to detect the polarization of a photon or in Geiger counters.

At an abstract level, we can think of the collapse as a series of repeated NOT gates, amplifying the quantum bit, or qubit, to a macroscopic scale. It's important to note that this process doesn't involve copying the qubit. Instead, it entangles the qubit with additional qubits, gradually increasing the scale of entanglement. Eventually, the entanglement reaches such a large scale that nature can no longer sustain it, leading to a measurement.

Understanding the nature of measurement in quantum systems is still an ongoing area of research. The perspective of repeated NOT gates and entanglement provides a useful framework to conceptualize this process. It allows us to grasp how a superposition transforms into a measurement outcome.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPUTATION****TOPIC: N-QUBIT SYSTEMS**

Quantum Information - Quantum Information Fundamentals - Introduction to Quantum Computation - N-qubit systems

In this section, we will discuss the fundamentals of quantum computation, specifically focusing on N-qubit systems. Quantum algorithms and circuits will be analyzed in the rest of the course. Before diving into quantum algorithms, it is essential to understand the basic concept that forms the foundation of quantum algorithms and why quantum computers have the potential to be exponentially powerful.

One of the most counterintuitive aspects of quantum mechanics is the concept of superposition, which allows a qubit to exist in a combination of states. A qubit can be represented by a unit vector in a two-dimensional vector space, where the states 0 and 1 correspond to the ground and excited states of an electron in a hydrogen atom.

When we introduce an additional qubit, creating a two-qubit system, the quantum state becomes a superposition of all four possible states. This can be visualized as a unit vector in a four-dimensional complex vector space. Similarly, adding more qubits exponentially increases the number of possibilities. For a three-qubit system, the quantum state is a superposition of all eight possibilities, represented by a unit vector in an eight-dimensional complex vector space.

The exponential growth in dimensionality can be understood by considering the tensor product of the individual Hilbert spaces of each qubit. Taking the tensor product of a two-dimensional Hilbert space with itself three times results in an eight-dimensional complex vector space. This exponential growth continues for larger systems, such as an N-qubit system, where the quantum state is a superposition of all 2^N possibilities. The dimensionality of the complex vector space is 2^N , represented as C^2 tensor C^2 tensor C^2 tensor ... tensor C^2 (N times).

This exponential growth in dimensionality is remarkable, even for relatively small values of N. For example, for $N=500$, the dimensionality of the complex vector space is larger than the number of particles in the universe and the age of the universe in femtoseconds. This exponential growth in dimensionality implies that a quantum computer with N qubits has more computing power than any classical computer that could operate for the age of the universe with an incredibly fast cycle time.

The exponential growth in dimensionality arises from the tensor product of the individual systems and the entanglement between them. When two systems are combined quantumly, their composite system is represented by the tensor product of their Hilbert spaces. The number of parameters required to describe the composite system is the product of the dimensions of the individual systems. This exponential growth is due to the entanglement between the systems, which necessitates describing the state of the composite system as a superposition over all the possibilities.

To summarize, quantum computation with N-qubit systems offers exponential growth in computing power due to the superposition principle and entanglement. The dimensionality of the complex vector space representing the quantum state increases exponentially with the number of qubits, allowing for a vast number of possibilities.

In quantum information, the state of a system in an N-qubit system is described by a superposition over all possible n-bit strings. Each bit string has an associated amplitude, denoted as α_x , and the state is normalized so that the sum of the squares of the magnitudes of α_x is equal to one.

The evolution of the system is achieved through the application of quantum gates. Quantum gates are represented by matrices, typically four by four, that act on a subset of qubits while leaving the rest unchanged. When a gate is applied, the Hilbert space representing the complex vector space of the system is rotated, resulting in a change in the state of the system and the updating of the complex amplitudes.

For example, if a Hadamard gate is applied to a qubit in an n-qubit system, the amplitudes of the paired-up bit strings are affected. The amplitudes of the form $0x$ and $1x$ are updated according to specific

formulas, involving the original amplitudes α_{0X} and α_{1X} . This gate effectively mixes the amplitudes of the paired-up bit strings, updating all the 2^n amplitudes α_X in the system.

This process reveals the remarkable nature of quantum computation. In a system with a large number of qubits, such as a 500-qubit system, nature must keep track of 2^{500} complex numbers, each associated with a specific bit string. Even a simple operation on the qubits, like a Hadamard gate, requires updating all these complex numbers. The sheer magnitude of the numbers involved is staggering, surpassing the number of particles in the universe and the age of the universe in femtoseconds.

One could question how nature is capable of carrying out such an extravagant task. However, an alternative perspective is to consider how we can harness this behavior for our benefit. If nature operates at the quantum level, shouldn't we be utilizing quantum computation instead of classical computation? After all, a computer is essentially a physics experiment, and if nature is already working at the quantum level, we can leverage it to solve problems of interest.

However, accessing the private world of nature's exponential superposition poses a challenge. As soon as we observe or measure the system, we collapse the superposition and only obtain a single outcome. This phenomenon underscores the delicate nature of quantum information and the need for careful handling and measurement techniques.

Quantum information in N-qubit systems involves the superposition of all possible n-bit strings, with associated amplitudes. The evolution of the system is achieved through the application of quantum gates, which rotate the Hilbert space and update the complex amplitudes. Quantum computation takes advantage of this behavior, recognizing that nature already operates at the quantum level. However, accessing and utilizing quantum information poses challenges due to the delicate nature of measurements and observations.

In the realm of quantum information, one of the fundamental concepts is the N-qubit system. An N-qubit system refers to a system composed of N quantum bits or qubits. These qubits can exist in multiple states simultaneously, thanks to the principles of superposition and unitary evolution.

In quantum mechanics, the behavior of nature is often likened to that of individuals who lead rich and private lives. Similarly, in a quantum system, the state of a qubit is described by a complex number known as the probability amplitude. The probability amplitude, denoted by α , represents the likelihood of a particular state. To determine the probability of observing a specific state, we square the magnitude of the probability amplitude.

However, one intriguing aspect of quantum mechanics is the measurement postulate. This postulate suggests that nature conceals its true state and covers its tracks, making it challenging to observe what is happening behind the scenes. This raises the question of whether we can ever truly uncover the secrets of nature.

The field of quantum algorithms and quantum computing aims to address this tension between the measurement postulate and the principles of superposition and unitary evolution. Researchers in this field strive to exploit the exponential power offered by quantum systems despite the limited access granted by the measurement postulate.

By developing quantum algorithms, scientists hope to peel back the veil and gain insight into the inner workings of nature. These algorithms leverage the unique properties of quantum systems to solve complex problems more efficiently than classical computers.

The field of quantum algorithms and quantum computing seeks to explore the exponential power of quantum systems while navigating the limitations imposed by the measurement postulate. By doing so, researchers aim to unravel the mysteries of nature and unlock new possibilities for computation.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPUTATION****TOPIC: UNIVERSAL FAMILY OF GATES**

A quantum circuit is a fundamental model of computing in quantum information. It consists of n qubits, which are initially in the state 0. Each qubit is represented by a wire, and these wires carry qubits of information. Similar to electrons in a hydrogen atom, these qubits can be in a ground or excited state, or even in a superposition of the two. As the circuit progresses, the qubits become entangled.

Quantum circuits are composed of a sequence of gates. These gates can act on one or two qubits. For example, there may be a gate that acts on two qubits and produces two output qubits. There can also be single qubit gates, such as the Hadamard gate or the Pauli Z gate. The circuit consists of these wires, which go from left to right, with gates applied along the way. The entire arrangement is called a quantum circuit.

To obtain classical information from the quantum circuit, we can measure the qubits. We can select specific qubits and use a measuring apparatus to obtain a classical string as the output.

In classical computing, various types of gates can be used, such as AND, OR, and NOT gates. However, it is interesting to note that a NAND gate is sufficient to perform any computation. Similarly, in quantum computing, there exists a universal family of gates that can be used to implement any quantum circuit. These gates include the Hadamard gate, the Pauli X gate, the Pauli Z gate, and the $\pi/8$ rotation gate.

It is worth mentioning that the universality of these gates means that any desired quantum computation can be achieved using only these gates. Even if someone claims to have a more powerful gate, it can be shown that it can be implemented using the universal gates. However, due to the limitations of perfect precision, an approximation of the desired gate is implemented, which is epsilon close to the original gate.

Quantum circuits are composed of qubits represented by wires, with gates applied to perform computations. The universality of certain gates allows for the implementation of any desired quantum circuit. Approximations are used to handle the limitations of perfect precision.

Quantum computation involves the use of quantum systems to perform computational tasks. In order to perform these tasks, we need to be able to manipulate quantum states. One way to do this is by using a universal family of gates.

A universal family of gates is a set of gates that can be combined to create any unitary transformation on a quantum state. These gates can be represented by matrices, where each gate corresponds to a specific matrix. In order to perform a computation, we need to apply a sequence of these gates to our initial quantum state.

The number of gates needed to perform a computation depends on the size of the system and the desired accuracy. If our quantum system has dimension D and we want to be epsilon close to the desired result, the number of gates needed scales roughly like D^2 times some polynomial in $1/\epsilon$. This means that as the size of the system and the desired accuracy increase, the number of gates needed also increases.

This dependence on D^2 is necessary because there are a large number of unitary transformations that can be represented by $D \times D$ matrices. Without a sufficient number of gates, we would not be able to express all of these transformations. Therefore, the dependence on D^2 is essential to ensure that we have enough gates to perform the desired computation.

In practice, we can implement a circuit that uses epsilon accuracy by combining different gates such as CNOT gates, Hadamard gates, X gates, and Z gates. This circuit behaves almost the same as the ideal transformation, with the difference between the two being epsilon close in operator norm. This means that if we apply both the ideal circuit and the epsilon circuit to the same quantum state, the resulting states will be epsilon close to each other in Euclidean norm.

A universal family of gates is essential for performing quantum computations. The number of gates needed depends on the size of the system and the desired accuracy. By combining different gates, we can implement a circuit that behaves almost the same as the ideal transformation. This allows us to perform computations with a

desired level of accuracy.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPUTATION****TOPIC: REVERSIBLE COMPUTATION**

In this material, we will discuss the concept of reversibility in quantum circuits. Quantum computers are always reversible, meaning that any computation performed can be undone. Let's consider a quantum circuit that takes an input state X and produces an output state u times X . We can reverse this process by applying the inverse of the circuit, denoted as $u^\dagger$, to the output state u times X . This will result in the original input state X . The inverse of a unitary quantum gate is simply its conjugate transpose.

To illustrate this, let's examine the gates inside the circuit. Each gate can be reversed by taking its conjugate transpose and applying the gates in the opposite order. When we combine these two circuits, the gates cancel each other out, resulting in the identity map that maps X to itself. This demonstrates the reversibility of quantum computation.

Now, let's consider implementing a classical circuit in a quantum setting. Suppose we have a classical circuit that takes n input bits and produces one output bit based on a boolean function f . To compute this quantumly, we need a quantum circuit that computes f on input X . From the reversibility property, we should also be able to recover the input X from the output. However, this may not always be possible.

For example, let's consider an AND gate. It takes two input bits, a and b , and outputs a bit that is 1 only if both input bits are 1. If we try to recover the input bits from the output, we find that it is not possible. The output bit alone does not provide enough information to determine the values of a and b . This shows that implementing certain functions in a straightforward way is not reversible.

To overcome this, we can modify our approach. Instead of just computing $f(X)$, we can introduce an answer bit B , initially set to 0, to store the output of $f(X)$. The quantum circuit can then output X unchanged and XOR B with $f(X)$. This flips B if and only if $f(X)$ is 1. By computing the inverse of this circuit, we can restore the original state of the output bit.

Now, let's examine how we can implement basic gates in a reversible manner. The NOT gate, which takes a bit as input and outputs its negation, is already reversible. Its quantum analog is the X gate, which performs the same operation on qubits. Similarly, other basic gates can be implemented in a reversible manner.

Reversibility is a fundamental aspect of quantum computation. Quantum circuits can always be reversed, allowing us to undo any computation performed. However, when implementing classical circuits in a quantum setting, we need to consider the reversibility of the functions involved and modify our approach accordingly.

In the realm of quantum information and quantum computation, reversible computation plays a crucial role. Reversible computation refers to a computational process that can be undone, allowing for the retrieval of the original input from the output. This property is highly desirable in quantum computing as it ensures the conservation of information and enables the implementation of quantum algorithms.

To understand reversible computation, let's first consider the XOR (exclusive OR) gate. The XOR gate takes two input bits, A and B , and outputs the result of their logical exclusive OR operation. This gate is reversible since it has an inverse that can recover the original inputs from the output. However, the situation is different for the AND gate, which is not reversible.

To address this issue, we introduce the concept of a controlled swap gate. This gate has three wires: a control wire and two target wires. If the control wire is set to 1, the values of the target wires are swapped; otherwise, they remain unchanged. By using this controlled swap gate, we can compute the AND of two bits in a reversible manner.

Consider the case where the control wire is set to 0. In this scenario, the output is always 0, regardless of the values of the input bits. If the control wire is set to 1 and the input bit A is also set to 1, the output is equal to the value of input bit B . Therefore, this controlled swap gate effectively computes the AND gate, as the output is 1 only when both input bits A and B are 1.

By combining the controlled swap gate with the NOT gate (which is already reversible), we can construct a universal gate for classical computation called the NAND gate. The NAND gate can be used to replace all the AND gates in a classical circuit, making the entire circuit reversible. During this process, it may be necessary to introduce additional fresh bits initialized to zeros, and the resulting circuit may produce some unwanted output bits. However, these extra bits can be disregarded as they are considered "junk bits."

Now, let's explore how this reversible circuit can be used in quantum computing. Suppose we have a reversible circuit and we want to represent it as a unitary transformation. Each gate in the reversible circuit can be seen as a unitary transformation since gates like the NOT gate, CNOT gate, and controlled swap gate are all unitary. Therefore, the reversible circuit can be implemented using a sequence of unitary gates.

By considering the reversible circuit as a unitary transformation, we can also investigate its behavior when provided with input in the form of a superposition over all possible inputs. If we feed in a superposition of inputs, the output of the circuit will be a superposition of outputs. However, it is crucial to note that the unwanted "junk" bits cannot simply be discarded.

In quantum mechanics, the rules state that all possible outcomes must be considered, and discarding the "junk" bits would violate this principle. Instead, we aim to design a reversible circuit that erases the junk bits and leaves the input string intact. This desired circuit takes the input string X and a series of zeros and outputs X along with the output of the original circuit and additional zeros.

This approach is highly preferable as it ensures that no information is lost and allows for further manipulation and exploration of the input state. It is important to emphasize that discarding the junk bits would result in an incomplete representation of the circuit's behavior and could lead to erroneous results.

Reversible computation is a fundamental concept in quantum information and quantum computation. By utilizing reversible gates, such as the controlled swap gate, and constructing universal gates like the NAND gate, we can transform classical circuits into reversible circuits. These reversible circuits can then be represented as unitary transformations, enabling their implementation in quantum computing. It is crucial to preserve all output bits, including the so-called "junk" bits, to maintain the integrity of the circuit's behavior and adhere to the principles of quantum mechanics.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPUTATION****TOPIC: CONCLUSIONS FROM REVERSIBLE COMPUTATION**

In quantum computation, it is crucial to remove junk qubits as they can prevent quantum interference. To illustrate this, let's consider an example. Suppose we have a circuit that takes an input X and outputs X , where X is a bit. In a classical circuit, this would simply be a circuit that does nothing to the input. However, in a quantum circuit, the input would be in a superposition state, denoted by summation αX , and the output would also be in a superposition state, denoted by summation X .

Now, let's introduce a Hadamard gate into the circuit. If we set the input to be in the plus state, represented by $\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$, the output of the first circuit would also be in the plus state. When this is fed into the Hadamard gate, the output would be the zero state. If we were to measure this output, we would observe a 0 with probability 1. This is the desired outcome.

However, let's now consider a classical circuit that takes an input X and outputs X , but also creates some junk qubits that are a function of X . We convert this classical circuit into a reversible circuit, denoted as R_C , by using quantum gates. Let's assume we use the C swap gate as part of a universal family of gates. This reversible circuit takes an input X and some clean bits 0, and outputs the correct answer while creating junk qubits that are a function of X .

Now, if we feed the output of this reversible circuit into a Hadamard gate, something interesting happens. Let's assume the input to the circuit is the plus state, $\frac{1}{\sqrt{2}}|0\rangle + \frac{1}{\sqrt{2}}|1\rangle$. The output of the circuit would be $\frac{1}{\sqrt{2}}|00\rangle + \frac{1}{\sqrt{2}}|01\rangle$. When we apply the Hadamard gate to the first qubit, we get $\frac{1}{2}|00\rangle + \frac{1}{2}|10\rangle$, and this gets transformed to $\frac{1}{2}|01\rangle + \frac{1}{2}|11\rangle$. As a result, we obtain all four possible states.

However, when we measure the first qubit, we observe a 0 and a 1 with equal probability, which is different from the desired outcome. This is because the junk qubits prevent the interference pattern from occurring. In the absence of junk qubits, when we apply the Hadamard gate to the plus state, we get $\frac{1}{2}|0\rangle + \frac{1}{2}|1\rangle$ and $\frac{1}{2}|0\rangle - \frac{1}{2}|1\rangle$. The interference between these two states leads to the desired outcome. But with the presence of junk qubits, these two states cannot interfere with each other, resulting in the wrong results in our quantum computation.

One might think that throwing away the junk qubits would solve the problem. However, this is not possible because the qubits are entangled. Even if we send the qubits far away from each other, they remain entangled. Throwing away a qubit is equivalent to measuring it, which does not help in changing the state of the remaining qubits to a tensor product state.

To overcome this issue, we need to modify the circuit to ensure that junk qubits are not created. Fortunately, there is an elegant solution to this problem. We can change the circuit in such a way that the junk qubits are not generated.

It is essential to remove junk qubits in quantum computation as they can prevent quantum interference. The presence of junk qubits can lead to incorrect results in our computations. Throwing away the junk qubits is not a viable solution as they are entangled with other qubits. Instead, we need to modify the circuit to avoid the creation of junk qubits.

In the field of quantum information, reversible computation plays a crucial role in the development of quantum circuits. In this context, the concept of "junk" arises when a reversible circuit produces unwanted outputs that depend on the input. However, there is a way to eliminate this junk while preserving the desired output.

To achieve this, we can apply the inverse of the circuit to undo the junk and restore all the bits back to 0. However, a problem arises as this inverse operation also reverses the answer, which is not desirable. To overcome this issue, we can copy the answer before applying the inverse circuit.

To do this, we start with fresh bits, setting them to 0. Then, we perform a controlled-not operation from each answer bit to the corresponding fresh bit. This allows us to obtain a copy of the answer bits. Now, we can safely apply the inverse circuit, which erases the junk, restores the input, and preserves the copied answer.

By following this approach, we achieve the desired outcome. Starting with an input X and a bunch of zeros, the output of the circuit will be X . Importantly, there is no junk associated with the input. This quantum circuit effectively eliminates interference and performs as intended.

The implications of this approach are significant. For any given classical circuit C , we can transform it into a quantum circuit, denoted as $u \text{ sub } C$. This quantum circuit takes as input X and a series of zeros, producing an output of X and C of X , along with additional zeros. Furthermore, this transformation extends to superpositions, where the input can be a sum over X with corresponding coefficients.

Applying the $u \text{ sub } C$ circuit to a superposition input results in a superposition over X in the first register, C of X in the second register, and a series of zeros in the third register. This notation is represented as $X \otimes Y$, which is equivalent to $x \text{ tensor } Y$, indicating the state of the qubits.

This theorem establishes that any classical circuit can be converted into a corresponding quantum circuit. Initially, the interest in quantum computation arose from the question of whether quantum mechanics imposes additional constraints on what can be computed compared to classical computation. However, this theorem demonstrates that there are no such constraints and that quantum mechanics allows for the conversion of classical circuits into quantum circuits.

The concept of reversible computation in quantum information provides a solution to eliminate unwanted outputs, known as junk, while preserving the desired output. Through the application of the inverse circuit and the use of controlled-not operations, it is possible to copy the answer before removing the junk. This approach allows for the conversion of classical circuits into quantum circuits, enabling quantum computation to achieve the same results as classical computation while taking advantage of the unique properties of quantum mechanics.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: FOURIER SAMPLING**

Quantum algorithms are an important aspect of quantum computing. In this lecture, we will discuss the building blocks of quantum algorithms and how they were used in early quantum algorithms to showcase the power of quantum computing.

In the previous lecture, we talked about reversible computation in the classical setting. We considered a classical circuit that takes an input X and computes $C(X)$, where $C(X)$ is a boolean value. We introduced a classical reversible circuit that takes X , an answer bit B , and a number of work bits initialized to 0. This reversible circuit outputs X unchanged, leaves the work bits in a clean state (initialized to 0), and XORs the answer $C(X)$ with the answer bit B . The output bit becomes $B \text{ XOR } C(X)$, effectively toggling the bit B .

We then discussed how this classical reversible circuit can be implemented using quantum gates. We introduced the unitary transformation U_C , which behaves the same as the classical reversible circuit when the input bits are in classical basis states. However, since U is a unitary transformation, it can also take superposition inputs. For example, if we give it a superposition input of the form $\sum \alpha_x |x\rangle$, where α_x is the amplitude and $|x\rangle$ represents the input state, the output will be $\sum \alpha_x |x\rangle \oplus (B \text{ XOR } C(X))$. This allows us to compute in superposition.

While the ability to compute in superposition is a fundamental primitive for quantum computation, it is not sufficient on its own. We need one more ingredient, which is the Hadamard transform. The Hadamard transform is a transformation on a single qubit that maps $|0\rangle$ to $|+\rangle$ and $|1\rangle$ to $|-\rangle$, where $|+\rangle = 1/\sqrt{2}(|0\rangle + |1\rangle)$ and $|-\rangle = 1/\sqrt{2}(|0\rangle - |1\rangle)$.

We can also perform the Hadamard transform on two qubits. If we apply the Hadamard transform to the input state $|00\rangle$, it gets mapped to $|++\rangle$, which is equal to $1/2|00\rangle + 1/2|01\rangle + 1/2|10\rangle + 1/2|11\rangle$. In general, the Hadamard transform on two qubits takes the input state $|xy\rangle$ and outputs an equal superposition of all four possible input states. We can work out the transformation for other input states as well, such as $|11\rangle$, which gets mapped to $|--\rangle = 1/2|00\rangle - 1/2|01\rangle + 1/2|10\rangle - 1/2|11\rangle$.

To understand the sign pattern in the transformation, we need to realize that each Hadamard transform maps 1 to -1 when starting and ending with 1. The sign pattern is determined by the parity of the number of transitions from 1 to 1. For example, if there are an odd number of transitions, the sign is negative, and if there are an even number of transitions, the sign is positive.

To write out the unitary transformation for the Hadamard transform on two qubits, we take the tensor product of the 2×2 Hadamard matrix with itself. The resulting matrix is a 4×4 matrix with entries $1/2$, $1/2$, $-1/2$, and $-1/2$. Similarly, we can perform the Hadamard transform on three qubits, resulting in an 8×8 matrix.

Quantum algorithms utilize building blocks such as reversible computation and the Hadamard transform. The ability to compute in superposition and the sign pattern of the Hadamard transform are key ingredients in quantum algorithms.

The Hadamard transform is a fundamental concept in quantum algorithms and plays a crucial role in Fourier sampling. It starts with a classical string on the left and gives a superposition over all the possible bit strings. This is a powerful principle in quantum computation.

When working with n qubits, if each qubit is initialized in the zero state and put through the Hadamard circuit, the output is the plus state on all the qubits. This can be denoted as the sum over all n -bit strings, where each string has equal amplitude of $2^{-(N/2)}$. This is achieved by taking the tensor product of $1/\sqrt{2}(|0\rangle + |1\rangle)$ with itself n times.

If we start with an input string $u = u_1u_2\dots u_s$, where s is the number of bits in u , and perform the Hadamard transform on it, the output is still the sum of all n -bit strings, but with each string having an amplitude of $1/2^{-(N/2)}$ multiplied by either a plus or minus sign. The sign is determined by $(-1)^{(u \cdot X)}$, where $u \cdot X$ is the dot product of the input and output bits. A minus sign is obtained when the input bit is equal to the output bit being

1.

For example, if $n = 3$, $u = 111$, and $X = 101$, then $u \cdot X = 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 1 = 2$. Thus, the amplitude of X would be $(-1)^2 / 2^{3/2} = 1/2^{3/2}$.

The primitive concept in quantum computation is to start with a superposition on n qubits, apply the Hadamard transform, and then perform a measurement. The output is a new superposition with different amplitudes, and the measurement yields a result with a probability equal to the magnitude squared of the amplitude. This process is called Fourier sampling.

Fourier sampling is a powerful primitive in quantum computation because it allows us to set up any superposition, apply the Hadamard transform, and measure the resulting probability distribution.

Quantum Information: Fourier Sampling

In the field of quantum information, one fundamental concept is the idea of Fourier sampling. Fourier sampling refers to the process of extracting information from quantum circuits that perform exponential work in determining the interference of amplitudes, which ultimately leads to the desired outcomes. This process is significantly more challenging to achieve classically.

To understand the significance of Fourier sampling, let's first examine the concept of amplitudes. In quantum circuits, amplitudes, denoted as α_x , represent the probabilities associated with different quantum states. These amplitudes interfere with each other, resulting in the final outcomes of the circuit. However, determining these outcomes through classical methods can be extremely difficult.

Fourier sampling offers a way to make sense of this complex interference. By utilizing quantum circuits, we can efficiently sample from the distribution of amplitudes and observe the resulting outcomes. This allows us to gain insights into the underlying patterns and behaviors of quantum systems.

One crucial aspect of Fourier sampling is its exponential computational power. While classical methods struggle to sample from the same distribution, quantum circuits excel in this task due to their ability to perform exponential work. This exponential advantage enables us to explore and analyze complex quantum phenomena that would otherwise be computationally infeasible using classical techniques.

Fourier sampling plays a vital role in quantum information by providing a means to extract information from quantum circuits efficiently. By harnessing the exponential computational power of quantum systems, we can gain insights into the interference of amplitudes and unlock the potential of quantum algorithms.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: APPLYING FOURIER SAMPLING**

In the field of quantum information, one important concept is the Fourier sampling algorithm. To understand the power of this algorithm, let's consider a simple problem known as the parity problem. In this problem, we are given a function f that takes n bits as input and produces a single bit as output. However, we don't know the internal workings of the function and can only run it on specific inputs.

The special property of the function f is that it computes the parity of a subset of the input bits. In other words, it is of the form $u \cdot x \bmod 2$, where u is a hidden N -bit string. For example, if $n = 3$ and $u = 101$, then $f(x) = x_1 \oplus x_3$. Our goal is to determine the hidden parity mask u using the function f .

Classically, we can determine u by running the function f on different inputs. By systematically varying the inputs, we can obtain one bit of information about u with each query. Since we need to reconstruct all n bits of u , we require at least n queries.

In the quantum world, we can use the Fourier sampling algorithm to reconstruct u using fewer queries. The algorithm works by creating a superposition of all possible input bit strings x , with a phase of -1 if and only if $f(x) = 1$. This superposition is known as the phase state.

To reconstruct u , we apply the Fourier transform to the phase state. This transforms the phase state into a state that is exactly what we would get if we applied the Fourier transform to u . By running the circuit backwards, we can obtain the hidden u that we were looking for.

To set up the initial superposition, we start with n bits in the 0 state and apply the Fourier transform to them. We then set the answer bit, denoted as B , to the state $-\frac{1}{\sqrt{2}}|0\rangle - \frac{1}{\sqrt{2}}|1\rangle$. This ensures that when $f(x) = 0$, the answer bit remains unchanged, and when $f(x) = 1$, the answer bit is flipped.

By applying the Fourier transform and measuring the answer bit, we can obtain the hidden parity mask u with just one query to the circuit for computing f . This is a significant improvement compared to the classical case, where n queries are needed.

The Fourier sampling algorithm allows us to reconstruct the hidden parity mask u using fewer queries in the quantum world compared to the classical world. By creating a specific superposition and applying the Fourier transform, we can obtain the desired information with just one query.

In quantum computing, Fourier sampling is a fundamental concept used in quantum algorithms. It involves applying the Fourier transform to a superposition of states in order to extract useful information. In this didactic material, we will explore the process of applying Fourier sampling and its significance in quantum information.

To understand Fourier sampling, let's consider a simple example. Suppose we have a function, $f(X)$, where X represents the input and $f(X)$ represents the output. We can think of $f(X)$ as a black box that takes an input and produces an output. In classical computing, we would need to query the black box multiple times to determine the output for different inputs.

In quantum computing, however, we can leverage the power of superposition and perform computations on multiple inputs simultaneously. This is where Fourier sampling comes into play. By applying the Fourier transform to a superposition of inputs, we can extract information about the function $f(X)$ without individually querying each input.

To illustrate this, let's consider a circuit that computes $f(X)$ using Fourier sampling. We start by preparing the answer bit as a minus state $(-)$. Then, we run the circuit for computing $f(X)$ with the input as a superposition over all X . The output bit, which represents the answer, is set as the minus state as well.

During the computation, the phase of the superposition changes depending on the value of $f(X)$. For inputs X such that $f(X)$ equals 0 , the phase remains unchanged. However, for inputs X such that $f(X)$ equals 1 , the phase changes to minus 1 . This means that we pick up a phase of minus 1 for those inputs.

After the computation, we have a tensor product state where the first n qubits represent the desired phase state. If we perform a Hadamard transform on these qubits, we recover the original state before the Fourier sampling. Finally, by measuring the qubits, we obtain the desired output.

This algorithm can be seen as a base case for a recursive algorithm called Fourier sampling. In this recursive version, we aim to amplify the difference between classical and quantum computations. In the classical case, solving a problem of size n requires n queries, while in the quantum case, a constant number of queries is sufficient.

To understand the power of Fourier sampling, let's consider the recursion for solving the parity problem. In the classical case, the time to solve a problem of size n is at least n times the time to solve a problem of size $n/2$, plus some additional time. This leads to a super-polynomial time complexity, growing at least like n to the power of $\log n$.

In contrast, the quantum algorithm for the recursive problem satisfies a different recursion. The time to solve a problem of size n is twice the time to solve a problem of size $n/2$, plus some additional time. The solution to this recursion yields a polynomial time complexity, growing like $n \log n$. This is similar to the recursion for merge sort.

This demonstrates the power of Fourier sampling in providing a super-polynomial speedup for certain types of problems. In the next material, we will explore how Fourier sampling can be used even more dramatically.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: SIMON'S ALGORITHM**

Simon's algorithm is a quantum algorithm that provides an exponential speed-up over classical algorithms for solving a specific problem. The problem involves finding a secret string s , given a function f that maps n -bit strings to n -bit strings in a two-to-one fashion.

To understand the problem, let's consider an example. Suppose we have a function f that takes in 3-bit strings and produces 3-bit strings as output. We have a secret string s , such that for any input x , $f(x)$ is equal to $f(x + s)$, where $+$ represents bitwise addition without carrying. For example, if s is 101, then $f(000) = f(000 + 101) = f(101)$ and $f(011) = f(011 + 101) = f(110)$.

The goal is to find the secret string s . In the classical setting, we would need to try different inputs until we find two inputs that produce the same output. This would require trying $2^{(n/2)}$ inputs, which takes exponential time.

Simon's algorithm, on the other hand, can solve this problem in polynomial time using quantum computation. The algorithm works in three steps.

Step 1: Set up an appropriate superposition. We create an equal superposition over two n -bit strings, R and $R + s$, where R is a random n -bit string.

Step 2: Perform a Fourier sampling of the superposition. By applying the Hadamard transform and measuring the outcome, we obtain a random n -bit string Y that satisfies the linear equation $Y \cdot s = 0 \pmod{2}$. This equation means that the bitwise dot product of Y and s is congruent to 0 modulo 2.

Step 3: Repeat step 2 $n-1$ times. Each repetition gives us a new linear equation. By solving these linear equations, we can determine the secret string s . Since there are $n-1$ linear equations and n unknowns, we will obtain exactly two solutions. One solution will be the all-zero solution, and the other solution will give us the value of s .

By following this algorithm, we can find the secret string s in polynomial time, providing an exponential speed-up over classical algorithms.

Simon's algorithm is a quantum algorithm that solves the problem of finding a secret string s given a two-to-one function f . By setting up an appropriate superposition and performing Fourier sampling, we can obtain linear equations that allow us to determine the secret string s . This algorithm provides an exponential speed-up over classical algorithms.

In Simon's algorithm, we are given a function f that takes an input X and outputs $f(X)$. The goal is to find the hidden string s such that $f(X) = f(X \oplus s)$, where $\oplus$ denotes bitwise XOR.

To solve this problem using quantum computing, we create a quantum circuit that takes input X and zeros as initial states, and applies a Hadamard transform on n input bits. This results in a superposition over all possible values of X , with amplitudes of $1/\sqrt{2^n}$.

Next, we feed this superposition through our circuit for computing $f(X)$. The output of the circuit is stored in the second register, which contains $f(X)$. We then measure the second register to obtain the value of $f(X)$.

If we were to look at an example, we would have a superposition over all possible values of $f(X)$. For each value of X , the first register would be 0^n with an amplitude of $1/\sqrt{2^n}$. After measuring the second register, we would obtain one of these outcomes with equal probability.

Suppose we measure the second register and obtain the outcome 1 0 0. To determine the new state of the system, we cross out all parts of the superposition that are inconsistent with this outcome. In this case, we cross out all instances of 0 0 0 and obtain a new state that is a superposition over the values X and $X \oplus s$, where s is the hidden string.

When we measure the first register, we will see a superposition of a random value X and $X \oplus s$, both with the same value of $f(X)$. The state of the first register after measurement will be $1/\sqrt{2}(R + 1) \oplus s$, where R is the measured value of $f(X)$.

To understand the output of the algorithm, let's consider the state before measurement as a sum over all possible values of Y , denoted as $\beta_{[sub]y[/sub]}|y\rangle$. The amplitude $\beta_{[sub]y[/sub]}$ from R to Y is given by $(-1)^{[sup]R \cdot Y[/sup]}/2^{[sup]n/2[/sup]}$. Since we started with an amplitude of $1/\sqrt{2}$, the amplitude from $R + s$ to Y is $(-1)^{[sup](R+s) \cdot Y[/sup]}/2^{[sup](n+1)/2[/sup]}$.

By factoring out common terms, we can express the amplitudes as $(2 - (-1)^{[sup]R \cdot Y[/sup]})/2^{[sup](n+1)/2[/sup]}$ and $(2 - (-1)^{[sup](R+s) \cdot Y[/sup]})/2^{[sup](n+1)/2[/sup]}$. We then consider two cases: when Y is congruent to 1 modulo 2 and when Y is congruent to 0 modulo 2.

If Y is congruent to 1 modulo 2, the amplitude $\beta_{[sub]y[/sub]}$ is equal to 0. If Y is congruent to 0 modulo 2, the amplitude $\beta_{[sub]y[/sub]}$ is $(-1)^{[sup](R+s) \cdot Y[/sup]}/2^{[sup](n+1)/2[/sup]}$. This can be further simplified to $(-1)^{[sup]s \cdot Y[/sup]}/2^{[sup](n-1)/2[/sup]}$.

When we sample the state, we observe each value Y with a probability of $\beta_{[sub]y[/sub]}^2$. In this case, the probability of observing Y is $1/2^{[sup](n-1)/2[/sup]}$.

From this analysis, we can conclude that exactly half of the possible values of Y satisfy the condition $Y \cdot s \equiv 0 \pmod{2}$. This means that out of the $2^{[sup]n[/sup]}$ possible bit strings, half of them have this property. When we sample, we are selecting each of these strings with equal probability, resulting in a random linear equation on the hidden string s .

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: CONCLUSIONS FROM SIMON'S ALGORITHM**

To understand the process of reconstructing a secret using Simon's algorithm, let's start with an example. Suppose we are working with three qubits, and the secret s is 101. The goal is to find the secret s using Fourier sampling.

When we perform Fourier sampling, we obtain a random Y such that $Y \cdot s = 0$. In other words, $Y_1s_1 + Y_2s_2 + Y_3s_3 = 0 \pmod{2}$. Working modulo 2 means dropping all the carries. So, what are all the possible Y values that satisfy this condition when s is 101?

If we work through the possibilities, we find that the Y values that satisfy this condition are 000, 010, 101, and 111. For example, if Y is 100, then $Y_1s_1 = 1$ and the other two terms are 0, so the sum is 1 (mod 2).

To reconstruct the secret s , we sample Y multiple times and run this procedure several times. Each time we obtain a linear equation, and by sampling an appropriate number of times, we can solve these equations to figure out the secret s . It is important to obtain independent equations for effective solving.

Let's consider an example where we sample Y twice. Suppose the first sample is 101 and the second sample is 111. Now, we have two equations:

$$\begin{aligned} 1s_1 + 0s_2 + 1s_3 &= 0 \\ 1s_1 + 1s_2 + 1s_3 &= 0 \end{aligned}$$

To solve these equations, we subtract the first equation from the second equation, which gives us $s_2 = 0$. Looking at the first equation by itself, we have $s_1 + s_3 = 0$. Therefore, the solutions to these equations are $s_1 = s_3$ and $s_2 = 0$. There are two possible solutions: $s_1 = s_3 = 0$ and $s_2 = 0$, or $s_1 = s_3 = 1$ and $s_2 = 0$. Since we assume that s is nonzero, we can eliminate the first solution, and we reconstruct the secret s as 10.

Now, let's consider how we would do this in general. Suppose we are working with n bits, and we are sampling Y such that $Y \cdot s = 0 \pmod{2}$. If we sample Y $n-1$ times, we hope to obtain independent linear equations in the s variables and solve for s . Since we have n unknowns and $n-1$ equations, we will get two solutions. One solution will always be $s = 0$, which we discard since we know s is nonzero. We take the other solution.

To calculate the probability of success for this algorithm, we need to consider the probability of obtaining independent linear equations at each step. Let's analyze this step by step.

First, we sample Y once. The only way we can fail is if we get the all 0 string, which is a trivial equation. The probability of failure in this step is $1/2^{n-1}$. Therefore, the probability that we are independent at this step is $1 - 1/2^{n-1}$.

Next, we sample Y_2 . We fail if Y_2 is the all 0 string or if it is equal to Y_1 . There are two ways of failing, so the probability of failure is $2/2^{n-1}$. The probability that we are still independent at this step is $(1 - 1/2^{n-1}) * (1 - 1/2^{n-2})$.

We continue this process for the remaining steps, and each time we have one additional way of failing. Therefore, the probability of failure at the third step is $4/2^{n-1}$.

In general, the probability that the algorithm succeeds is given by the product of the probabilities that we are independent at each step. So, the probability that the third choice is independent is $(1 - 1/2^{n-1}) * (1 - 1/2^{n-2}) * \dots * (1 - 1/2^2) * (1 - 1/2^1)$.

By analyzing the probabilities at each step, we can determine the success rate of Simon's algorithm for reconstructing the secret s .

In the context of quantum information, we will now discuss Simon's algorithm and draw some conclusions from it. Simon's algorithm is a quantum algorithm that aims to solve a specific problem known as Simon's problem.

The problem involves finding a hidden string of bits, which is a secret input to a black box function.

To understand Simon's algorithm, let's first consider the concept of independence. In this algorithm, the independence of each step is crucial. We want to ensure that the choices made in each step are independent of each other. The probability of independence in each step can be calculated by analyzing the possible subsets of the first n minus two elements. The cardinality of this set is 2 to the power of N minus two, which gives us the probability of failure as 2 to the power of N minus 2 divided by 2 to the power of N minus 1 , which simplifies to $1/2$.

To determine the probability of all steps being independent, we need to calculate the product of the probabilities of independence in each step. This product can be expressed as a series, which mathematicians have evaluated to be approximately 0.2887 . However, we can also look at the worst-case scenario by considering the probability of failure at each step. By summing up these probabilities, we find that the probability of failure is at most $1/2 + 1/4 + 1/8 + \dots + 1/2$ to the power of N minus 1 , which converges to $1 - 1/2$ to the power of N minus 1 .

Therefore, the probability of success is at least $1 - 1/2$ to the power of N minus 1 . However, this probability may not be satisfactory. To improve it, we can stop one step early and calculate the probability of success in all steps except the last one. By doing so, we find that the probability of success is at least $1/2$ times $1/2$, which is $1/4$.

Simon's algorithm involves starting with input qubits in the state 0 , applying a Hadamard transform, computing the black box function, and then measuring the result. This measurement provides a linear equation that the secret string must satisfy. This process is repeated n minus 1 times, resulting in n minus 1 linear equations. By solving these equations, the secret string can be determined.

Simon's algorithm is a powerful tool in quantum computing that allows us to solve Simon's problem efficiently. By ensuring the independence of each step and analyzing the probabilities involved, we can increase the chances of success in finding the secret string.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: SIMON'S ALGORITHM IN TERMS OF THE DOUBLE SLIT EXPERIMENT**

Simon's algorithm can be understood in terms of the double slit experiment. In the double slit experiment, a source of light emits single photons that pass through two slits in a screen. When both slits are open, an interference pattern is observed on a backdrop, indicating the probability distribution of where the photons will end up.

Simon's algorithm can be viewed as a sophisticated version of the double slit experiment. Instead of photons, we consider quantum bits or qubits. We start with n qubits in a specific state, denoted as $u_1, u_2, \dots, u_n$. The middle superposition in Simon's algorithm is a superposition of all n -bit strings, where each string has an amplitude of plus or minus $1/2^{n/2}$, depending on the dot product of the string with the input state u .

To understand what happens when we apply another Hadamard transform, we compute the amplitude β_y for each possible y . β_y is the sum over all x of the amplitude of x multiplied by the amplitude of going from x to y when a Hadamard transform is applied. There are two cases to consider.

In case 1, if y is equal to u , then β_y is equal to the sum over all x of the amplitude of x squared, which simplifies to $1/2^n$. This leads to constructive interference, where all the contributions add up to 1.

In case 2, if y is not equal to u , then for exactly half the values of x , the signs of the two amplitudes are equal, and for the other half, the signs are unequal. This results in destructive interference, where the contributions cancel out, and β_y becomes 0.

Therefore, Simon's algorithm can be seen as having 2^n virtual slits, and the Hadamard transform causes constructive interference at y equal to u and destructive interference everywhere else. By changing the slit pattern in the middle based on the input to the problem, we can determine where the constructive interference occurs, which gives us the solution to the problem.

In Simon's problem, we are given a two-to-one function f , where there exists a secret string s such that $f(x) = f(x + s)$. The algorithm consists of two Hadamard transforms and a quantum circuit for computing f . The middle part of the circuit represents the superposition of all input bit strings. By measuring the qubits, we obtain a superposition where the first n qubits are in a specific state related to the secret string s .

By manipulating the slit pattern in the middle, determined by the input to the problem, we can observe where the constructive interference occurs and obtain the solution to the problem.

In the field of quantum information, there is a fascinating algorithm known as Simon's algorithm. This algorithm can be understood in terms of the double slit experiment. In this experiment, two random slits are positioned among exponentially many possibilities. However, these slits differ by exactly " s ".

When we observe the interference pattern resulting from this setup, we find that there is constructive interference on exactly half of the 2^N bitstrings. On the other half, we observe completely destructive interference. Therefore, when we perform a measurement, we randomly obtain one of the bitstrings with constructive interference.

The key insight is that if we sample any one of these bitstrings with constructive interference, denoted as " Y ", it satisfies the condition that the dot product of " Y " and " s " is zero. This condition gives us a linear equation that the secret string " s " must satisfy.

This is where Simon's algorithm comes into play. It can be seen as a virtual double slit experiment, where the slits represent the input to the problem we are trying to solve. When we measure the output, we select one of the strings with constructive interference at random. This random string yields a linear equation, which provides a constraint on the secret string " s " that we are trying to find.

By solving these linear equations, we can reconstruct the secret string " s " that was hidden in the problem. Simon's algorithm thus offers a powerful tool for solving certain types of problems in the field of quantum

information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM ALGORITHMS****TOPIC: EXTENDED CHURCH-TURING THESIS**

Quantum Information - Quantum Information Fundamentals - Quantum Algorithms - Extended Church-Turing Thesis

In the field of quantum information, the study of quantum algorithms has significant implications for the understanding of fundamental questions about computers. When we discuss the ease or difficulty of solving certain problems on a computer, such as matrix multiplication, primality testing, or factoring, it is important to consider the type of computer we are referring to.

Computer scientists and researchers have extensively analyzed this question and arrived at a remarkable conclusion known as the extended Church-Turing thesis. This thesis states that the specific details of a model of computation are not crucial. Even the most basic model of computation, such as a Turing machine, is sufficient to capture the essence of computation.

A Turing machine consists of an infinite tape divided into squares that can hold either a zero or a one. It also has a read/write head that can access and modify the tape squares, as well as an internal control mechanism that follows a set of decision rules based on its current state and the value observed on the tape. This simple model can be seen as a representation of the functions that can be computed by humans using pen and paper.

Another model that captures the concept of computation is the cellular automaton. In this model, a grid of cells is arranged, with each cell having a finite number of possible states, such as black or white. At each step, each cell examines the states of its neighboring cells, including itself, and applies a set of rules to determine its new state. This model is particularly interesting because it reflects the idea of computation in nature.

In classical physics, the behavior of physical quantities is often described by local differential equations. The cellular automaton model can be seen as a discrete version of this approach. By considering a small neighborhood around a specific point, the model predicts how the value of a physical quantity changes based on the current values in that neighborhood.

The fact that a cellular automaton can be simulated by a Turing machine with only a polynomial factor slowdown suggests that both models capture the same class of functions that can be efficiently computed. This insight indicates that whether we view humans or nature as computers, the capabilities remain equivalent.

However, quantum computers challenge the extended Church-Turing thesis. They can solve certain problems faster than classical computers, as demonstrated by Simon's problem and the problem of quantum Fourier sampling. These quantum algorithms do not adhere to the extended Church-Turing thesis, raising intriguing questions about the limits of computation in nature.

This leads to the exploration of a potential quantum Church-Turing thesis, which would propose that a quantum computer represents the ultimate computational device capable of performing any computation that nature can accomplish.

Quantum information is a fascinating field that not only allows for fast computation on quantum computers but also raises intriguing questions about the potential power of quantum mechanics itself. In this course, we will explore some of these questions, although it is important to note that some of them go beyond the scope of what we can cover here.

One fundamental concept in quantum information is the Extended Church-Turing Thesis. This thesis suggests that any physically realizable computation can be efficiently simulated by a Turing machine. However, quantum computers challenge this thesis by demonstrating that certain problems can be solved more efficiently using quantum algorithms compared to classical algorithms.

Quantum algorithms are specifically designed to harness the unique properties of quantum systems, such as superposition and entanglement, to perform computations. These algorithms exploit quantum parallelism, allowing for the simultaneous evaluation of multiple possibilities. One famous example is Shor's algorithm,

which efficiently factors large numbers, posing a significant threat to current cryptographic systems.

To better understand the power of quantum algorithms, it is essential to have a solid grasp of quantum mechanics. Quantum mechanics describes the behavior of particles at the microscopic level and provides the foundation for quantum information processing. Key principles include superposition, where a quantum system can exist in multiple states simultaneously, and entanglement, where two or more particles become correlated in such a way that the state of one particle is dependent on the state of the others.

In addition to quantum algorithms, quantum information also encompasses other topics such as quantum error correction, quantum communication, and quantum cryptography. These areas explore how to protect and transmit quantum information reliably in the presence of noise and potential eavesdroppers.

The study of quantum information offers exciting possibilities for fast computation and challenges our understanding of the limits of classical computation. By delving into quantum algorithms and exploring the power of quantum mechanics, we can gain insights into the potential of quantum information processing.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM FOURIER TRANSFORM****TOPIC: QFT OVERVIEW**

The topic of this didactic material is Quantum Fourier Transform (QFT) which is an essential aspect of Quantum Information and Quantum Algorithms. The Quantum Fourier Transform is closely related to the Hadamard Transform and is considered the workhorse of quantum algorithms. In this material, we will explore the Quantum Fourier Transform, its properties, and its applications.

To understand the Quantum Fourier Transform, let's start with a simple example. Consider a three-qubit system. If we apply the Hadamard Transform to each of the three qubits, we obtain an eight by eight matrix, which is suitably normalized. The resulting matrix has columns that are orthogonal to each other, with exactly half the entries being the same and the other half being opposite.

Now, let's move on to the Quantum Fourier Transform on three qubits. The Quantum Fourier Transform is also represented by an eight by eight matrix. However, this matrix has entries involving a primitive 8th root of unity, denoted as ω . ω is a complex number that satisfies the equation $x^8 = 1$. It has eight complex solutions, and one of them is ω . The entries of the Quantum Fourier Transform matrix are precisely these eight complex roots of unity.

The Quantum Fourier Transform has several beautiful properties. The columns of the matrix are orthogonal to each other, and the normalization factor ensures that they have unit norm. The Quantum Fourier Transform is closely related to the Hadamard Transform and can be seen as a generalization of it.

One of the applications of the Hadamard Transform is Simon's algorithm, which involves discovering a secret number based on a given function. Similarly, the Quantum Fourier Transform is used in period finding, which is a fundamental problem in quantum algorithms. Period finding aims to discover the period of a periodic function. The algorithm for period finding closely resembles Simon's algorithm, but it utilizes the Quantum Fourier Transform instead of the Hadamard Transform.

In period finding, we apply a Quantum Fourier Transform followed by the function evaluation. We then measure the output qubits, apply the Quantum Fourier Transform again, measure once more, and obtain an output value. This output value allows us to efficiently reconstruct the period of the function.

It's important to note that period finding is a crucial step in Shor's quantum algorithm for factoring, which has significant implications in cryptography and number theory.

In the rest of this material, we will delve deeper into the complex roots of unity, the general concept of the Quantum Fourier Transform, including the n -band Quantum Fourier Transform, and explore the beautiful properties of the Quantum Fourier Transform.

The Quantum Fourier Transform (QFT) is an important concept in the field of Quantum Information. In this lecture, we will explore the efficient quantum circuit for the QFT and its significance in period finding and quantum algorithms.

It is worth noting that the material we are covering now is more open-ended compared to previous topics. While we have simplified the presentation of factoring for accessibility, some of you may desire more in-depth knowledge. To address this, there are additional resources available.

Firstly, you can refer to the course notes for more details. Additionally, we have suggested reference books at the beginning of the course that can provide further insights. For those who feel they lack the necessary background, it is important to note that we have aimed to make this material as self-contained as possible. However, if you would like to brush up on your background or learn more, the course notes and reference books are excellent resources.

Furthermore, there is an additional online reference available. It is an undergraduate textbook on algorithms, and the pre-publication version can be downloaded from the instructor's website. Specifically, there are three chapters that cover the background material related to transforming the factoring problem into the Quantum

Fourier Transform.

Chapter 1 focuses on modular arithmetic, which forms the foundation for understanding the QFT. The second chapter introduces the discrete Fourier transform, which is the classical algorithmic perspective on the Fourier transform. The second half of this chapter is particularly relevant to our discussion. Finally, chapter 10 delves into quantum factoring, which may be of interest to some of you.

By exploring these resources, you can gain a deeper understanding of the background material and the transformation of the factoring problem into the Quantum Fourier Transform.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM FOURIER TRANSFORM****TOPIC: N-TH ROOTS OF UNITY**

In the context of quantum information, it is important to have a solid understanding of the n th roots of unity. Let's review some complex notation to refresh our memory. A complex number X can be expressed in the form $\cos \theta + i \sin \theta$, which can also be written as $e^{i\theta}$. Similarly, let's consider another complex number Y , which can be written as $\cos \theta_2 + i \sin \theta_2$, or $e^{i\theta_2}$.

When we multiply X and Y , we obtain the product $(\cos \theta_1 + i \sin \theta_1) * (\cos \theta_2 + i \sin \theta_2)$. By simplifying this expression, we find that the product is equal to $\cos(\theta_1 + \theta_2) + i \sin(\theta_1 + \theta_2)$, or $e^{i(\theta_1 + \theta_2)}$. This shows that when we multiply complex numbers, the angles add up.

In this discussion, we are assuming that our complex numbers lie on the unit circle, meaning they have a magnitude of 1. If we were to plot these complex numbers on the complex plane, with the real axis and imaginary axis, the unit circle represents the points where X and Y lie.

Now, let's move on to the concept of complex n th roots of unity. These are the solutions to the equation $X^N = 1$. It turns out that there are exactly N complex solutions to this equation. If we consider the complex plane again, with 1 as the reference point, we can divide the angle 2π into N equal pieces. Let's call one of these pieces Ω , where $\Omega = 2\pi/N$. If N is 12, for example, there would be 12 solutions: Ω , Ω^2 , Ω^3 , and so on.

To understand this concept visually, imagine going around the unit circle N times, starting from the point 1. After going around N times, you end up back at the point 1, which is equivalent to 2π . This means that $\Omega^N = 1$.

There are some interesting properties associated with the roots of unity. For example, if we add up all the complex n th roots of unity, the sum is equal to 0. This is because when we add complex numbers, we treat them as vectors on the complex plane. Each root of unity represents a vector, and when we add them all up, they cancel each other out completely.

This property also holds true if we consider a sum of the form $1 + \Omega^J + \Omega^{2J} + \dots + \Omega^{(N-1)J}$. As long as J is not equal to 0, this sum is equal to 0. However, if J is equal to 0, the sum becomes N . More generally, if J is a multiple of N , the sum is also equal to 0.

To prove these properties, we can use the geometric series formula. By summing the series, we find that the sum is equal to $\Omega^{N * J} - 1 / (\Omega^J - 1)$. As long as J is not equal to 0, the denominator is nonzero, resulting in a sum of 0.

Lastly, let's consider the conjugate of Ω , denoted as $\bar{\Omega}$. The conjugate of Ω is equal to $\cos(2\pi/N) - i \sin(2\pi/N)$. Interestingly, the conjugate of Ω is the same as Ω to the power of -1 or $1/\Omega$. This is because $\Omega^N = 1$.

The n th roots of unity are complex solutions to the equation $X^N = 1$. When plotted on the complex plane, they lie on the unit circle. Adding up the roots of unity results in a sum of 0, and this property holds true for certain sums involving the roots. The conjugate of Ω is equal to Ω to the power of -1 .

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM FOURIER TRANSFORM****TOPIC: DISCRETE FOURIER TRANSFORM**

The Quantum Fourier Transform (QFT) is a fundamental concept in Quantum Information. It is also known as the Discrete Fourier Transform (DFT). The QFT can be defined by the operator QFT_n , which is a normalized n by n matrix. The entries of this matrix are the n th roots of unity, denoted as ω . The formula for ω is $e^{2\pi i/n}$, where i is the imaginary unit.

To better understand the QFT, let's visualize it as a matrix. We number the rows and columns from 0 to $n-1$. Each entry in the matrix, denoted as the jk th entry, is calculated as ω to the power of j times k . The matrix is normalized by a factor of $1/\sqrt{n}$.

Let's work through some examples to illustrate the QFT. For example, let's consider QFT_2 . By calculating the fourth root of unity, we find that ω is equal to i . The matrix for QFT_4 can be obtained by applying the normalization factor and filling in the entries using the formula mentioned earlier.

During the calculations, we may encounter powers of ω that exceed the range of the roots of unity. In such cases, we can use modular arithmetic to simplify the calculations. Modular arithmetic allows us to replace a power of ω with its remainder when divided by n . This concept is crucial in understanding the QFT and will be used extensively in future lectures on factoring.

If you are interested in learning more about modular arithmetic, you can refer to online resources or consult the relevant chapters of books on algorithms.

Now, let's discuss how to apply the QFT. We apply the QFT to a state, which can be represented as a vector or in ket notation. In the case of a two-qubit system, the state can be written as a linear combination of basis states, such as $\alpha_0|00\rangle + \alpha_1|01\rangle + \alpha_2|10\rangle + \alpha_3|11\rangle$.

When we apply the QFT to a state, we obtain a new superposition of states, represented by $\beta_0, \beta_1, \beta_2$, and β_3 . To illustrate this, let's consider an example where the initial state is $|2\rangle$. Applying the QFT to this state, we find that the new state is $\beta_2 = 1$, while the other coefficients are 0.

The Quantum Fourier Transform is a fundamental tool in Quantum Information, enabling us to manipulate and analyze quantum states. Understanding its mathematical formulation and application is essential for further exploration in this field.

In the realm of quantum information, the Quantum Fourier Transform (QFT) plays a crucial role in various quantum algorithms. The QFT is a quantum analogue of the classical Discrete Fourier Transform (DFT) and is employed to convert a quantum state represented in the computational basis to its Fourier basis. This transformation is particularly useful in applications such as factoring large numbers and simulating quantum systems.

To understand the QFT, let's consider an example. Suppose we have a quantum state represented by a column vector $[a_0, a_1, a_2, a_3]$. The QFT operates on this vector to produce a new vector, where each element is a linear combination of the original elements. In this case, the transformed vector would be $[1/2 * (a_0 - a_1 + a_2 - a_3), \dots]$.

To illustrate this further, let's focus on the third column of the transformed vector. We obtain the value $1/2 * (a_0 - a_1 + a_2 - a_3)$. This value is obtained by taking a linear combination of the original elements in the third column of the initial vector.

It is important to note that the QFT is a reversible transformation, meaning that it can be reversed to obtain the original state. This property is crucial for the functioning of many quantum algorithms.

The QFT can be implemented using quantum gates such as the Hadamard gate and controlled-phase gates. These gates act on the individual qubits of the input state to perform the necessary operations for the transformation.

The Quantum Fourier Transform (QFT) is a fundamental tool in quantum information processing. It allows us to convert a quantum state from the computational basis to its Fourier basis, enabling various quantum algorithms. The QFT is a reversible transformation and can be implemented using quantum gates. Understanding the QFT is essential for delving deeper into the field of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM FOURIER TRANSFORM****TOPIC: N-TH DIMENSIONAL QUANTUM FOURIER TRANSFORM**

The quantum Fourier transform (QFT) is a fundamental operation in quantum information processing. It is analogous to the classical discrete Fourier transform (DFT) and plays a crucial role in various quantum algorithms. In this didactic material, we will discuss the efficiency of implementing the QFT and its significance in quantum computing.

The QFT matrix is equivalent to the DFT matrix used in classical computing. However, in the QFT, we often include a normalizing factor of $\frac{1}{\sqrt{N}}$ over the square root of N , where N represents the dimension of the input and output vectors. The input vector is a complex vector of dimension N , and the output vector is also a complex vector of dimension N .

To compute the product of the QFT matrix and a vector, we multiply each entry of the vector by the corresponding column of the matrix and sum the results. This requires approximately N multiplications and additions, resulting in a time complexity of order N . However, there are N^2 entries to compute, suggesting a time complexity of order N^2 .

Fortunately, the fast Fourier transform (FFT) algorithm provides a significant improvement in classical algorithms. The FFT algorithm reduces the time complexity from N^2 to approximately $N \log N$ steps. This nearly quadratic improvement is responsible for various applications in digital signal processing, such as music and video processing.

In the quantum case, we represent the input vector as the state of n qubits, where n is the logarithm base 2 of N . This exponential compression allows us to represent an exponentially large superposition. By inputting this state into a quantum circuit, we obtain the output qubits in a new state.

The complexity of the quantum circuit, measured by the number of quantum gates, can be as small as $O(N^2)$. With further optimizations, it is possible to reduce the complexity to order N . This exponential improvement in complexity is a remarkable achievement in quantum computing.

However, there is a catch in quantum computing. Unlike classical computing, where we obtain the complete output vector, in quantum computing, we can only measure a single index j with a probability proportional to the squared magnitude of the corresponding amplitude β_j . This limitation arises due to the superposition nature of quantum states.

This issue of limited access to the computed amplitudes is a significant challenge in quantum algorithms. We must find ways to utilize the powerful computations performed by nature and extract meaningful information from the limited measurements we can make. This challenge is at the heart of quantum algorithms and requires innovative techniques to overcome.

The QFT is a fundamental operation in quantum information processing, analogous to the classical DFT. The efficiency of implementing the QFT has been significantly improved by the FFT algorithm in classical computing. In quantum computing, the QFT can be implemented with a complexity of order N^2 or even reduced to order N . However, the limited measurements in quantum computing pose challenges in utilizing the full power of the computed amplitudes.

The exponential growth of technology has been a fascinating phenomenon, with various fields experiencing significant advancements over time. One such example is the observation made by Gordon Moore, the founder of Intel, who noticed that the number of transistors in a chip had been doubling every 18 months since 1965. This observation, known as Moore's Law, predicted that this trend would continue indefinitely into the future.

This exponential scaling has not only been observed in the number of transistors but also in other aspects of technological improvement. Processor speeds have increased, and the cost of computation has dropped exponentially, leading to the remarkable computer revolution we are currently experiencing.

To illustrate the impact of exponential improvement, Gordon Moore once gave a lecture using the example of

the automobile industry. He imagined a scenario where the industry had followed a similar trajectory since the late 1950s or early 1960s. In this hypothetical scenario, he suggested that by now, one could buy a Rolls Royce that would consume only a gallon of gas for approximately ten million miles of travel. This car would also be capable of traveling at 1% of the speed of light and cost less than a dime. However, a member of the audience humorously added that it would be as small as a matchbox.

The Quantum Fourier Transform (QFT) is another example of exponential improvement in performance. The QFT is a mathematical operation used in quantum computing to transform a quantum state into its frequency representation. It plays a crucial role in various quantum algorithms, such as Shor's algorithm for factoring large numbers.

The QFT offers incredible improvements in performance across multiple measures. However, the results it produces are often very small, which can make it challenging to utilize effectively. Despite this challenge, researchers and scientists continue to explore ways to harness the power of the QFT for practical applications in quantum information processing.

In the next lecture, we will delve deeper into the practical applications of the Quantum Fourier Transform and explore how this information can be effectively utilized.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: QUANTUM FOURIER TRANSFORM****TOPIC: PROPERTIES OF QUANTUM FOURIER TRANSFORM**

The Quantum Fourier Transform (QFT) is a fundamental concept in quantum information and is widely used in quantum algorithms. It possesses two important properties that make it useful in quantum computations.

The first property is the convolution multiplication property of the Fourier transform. When a QFT is applied to an input superposition state, it produces an output state that is related to the input state through a convolution operation. In other words, if we start with an input state and apply the QFT to it, we obtain a transformed state. This property is similar to the classical Fourier transform and is used in quantum sampling. After the QFT is applied, the output state can be measured, and the probability of observing a specific index J is given by the squared magnitude of the corresponding amplitude.

The second property of the QFT is its treatment of periodic functions. When the QFT is applied to a periodic function of period R , the resulting amplitudes also exhibit periodicity. The period of the transformed amplitudes is M/R , where M is the dimension of the QFT. This property is particularly useful when dealing with periodic functions in quantum computations.

To illustrate these properties, let's consider a special case. Suppose we have a function that is periodic with period R . The amplitudes of this function repeat every R indices. When we apply the QFT to this periodic function, the resulting amplitudes also exhibit periodicity. The period of the transformed amplitudes is M/R , where M is the dimension of the QFT.

In this special case, let's assume that the nonzero amplitudes of the input function are located at indices 0 , R , $2R$, and so on, up to $M/R-1$ times R . The number of nonzero amplitudes is M/R . To normalize the vector, the amplitude of each nonzero component should be the square root of R/M .

These properties of the QFT are important in quantum information processing. The convolution multiplication property allows for efficient computations of Fourier transforms in quantum algorithms. The treatment of periodic functions simplifies the analysis and manipulation of periodic quantum states.

The Quantum Fourier Transform possesses the convolution multiplication property and treats periodic functions in a special way. These properties make it a powerful tool in quantum information processing, enabling efficient computations and simplifying the analysis of periodic quantum states.

The Quantum Fourier Transform (QFT) is a fundamental operation in quantum information processing. It plays a central role in various quantum algorithms, including the quantum algorithm for factoring. In this didactic material, we will explore the properties of the Quantum Fourier Transform.

The QFT is a transformation that maps an input superposition to a new superposition. Specifically, it maps the input superposition to a new superposition with a period of M/R , where M is the original period and R is a positive integer. The new superposition consists of R non-zero amplitude states, ranging from 0 to $R-1$ times M/R . The amplitude of each state in the new superposition is $1/\sqrt{R}$.

To understand the QFT, let's represent the input superposition in vector notation. The input vector is normalized by a factor of $\sqrt{R}/M$, and its entries are initially 1 , followed by $R-1$ zeros. The distance between successive ones in the vector is exactly R .

When we perform the QFT on this input superposition, the resulting superposition is a summation of β_J , where J ranges from 0 to $M-1$. We are interested in understanding the values of β_J . Let's first consider the case when J is a multiple of M/R , denoted as $J = K$ times M/R .

For $\beta_{K \text{ times } M/R}$, we can derive an expression that involves a normalization factor, the phase factor ω raised to the power of JR times K times M/R , and the square root of R/M . Importantly, the phase factors cancel out, resulting in a uniform contribution for all components. Therefore, the amplitude of $\beta_{K \text{ times } M/R}$ is $1/\sqrt{R}$.

This observation leads us to an important insight. The QFT treats periodic functions in a special way. At the multiples of M/R , the phase factors align, resulting in constructive interference. This constructive interference occurs because the QFT hits the same point in the phase every time, leading to a uniform contribution. As a result, the amplitude of these components is $1/\sqrt{R}$.

On the other hand, for J values that are not multiples of M/R , the phases are not aligned. The QFT hits different points in the phase, leading to destructive interference. The phases are symmetrically distributed around the circle, and when we add up all these vectors, we obtain a zero vector. Therefore, the amplitude of β_J is zero whenever J is not a multiple of M/R .

The QFT treats periodic functions in a unique way. It exhibits constructive interference at the multiples of M/R , resulting in non-zero amplitudes, and destructive interference at all other points, resulting in zero amplitudes. This property of the QFT is crucial for various quantum algorithms, including the quantum algorithm for factoring.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: SHOR'S QUANTUM FACTORING ALGORITHM****TOPIC: PERIOD FINDING**

Shor's algorithm is a powerful algorithm for factoring integers. In this lecture, we will discuss the algorithm, focusing on the main building block called period finding.

Period finding involves finding the period of a periodic function. The function maps numbers from 0 to $n-1$ to a set S , and it has a period R . A periodic function repeats its pattern after a certain number of inputs. For example, if we have a function with a period of 5, it will repeat its pattern every 5 inputs. The function is also one-to-one within each period, meaning it does not repeat any values within a single period.

To make the factoring algorithm work, we need to solve period finding for a function where the period does not divide n . We also need to have a large number of repetitions of the period, which means M/R should be larger than the period itself. In other words, we need to see the function over many periods to extract the period information.

The function is given to us as a classical circuit, denoted as $C \text{ sub } F$. This circuit takes an input X and outputs $f(X)$. To solve period finding classically, we would need to randomly pick inputs and look for collisions, i.e., two different inputs that produce the same output. However, this approach would require a huge number of inputs, making it practically infeasible for large numbers.

Shor's algorithm uses quantum computing to solve period finding more efficiently. The classical circuit $C \text{ sub } F$ is converted into a quantum circuit. This quantum circuit takes inputs X in a superposition state and outputs $X f(X)$ in a superposition state as well. This means that the quantum circuit can process multiple inputs simultaneously.

The idea behind Shor's algorithm is to set up a uniform superposition of all possible inputs and run the quantum circuit. This results in a superposition of all possible outputs. By measuring the output, we can extract information about the period. The quantum algorithm significantly reduces the number of inputs needed to find the period, making it feasible for large numbers.

Shor's algorithm for factoring integers relies on period finding as its main building block. Period finding involves finding the period of a periodic function, which is done more efficiently using quantum computing. By setting up a superposition of inputs and running a quantum circuit, we can extract the period information from the output. This algorithm revolutionizes the field of integer factoring and has important implications for cryptography.

In the study of quantum information, one of the fundamental concepts is Shor's Quantum Factoring Algorithm, which involves the process of period finding. The goal of this algorithm is to find the period of a given function.

To understand how period finding works, let's consider an example. Suppose we have a function $f(x)$ and we want to find its period. We start by preparing a quantum superposition of all possible inputs, which is achieved by applying the quantum Fourier transform (QFT) to the input register.

Next, we apply the function $f(x)$ to the input register. The output of this function is stored in a second register. If we were to measure this second register, we would obtain a random value, let's say $f(a) = 4$ for some value of a . However, there are multiple values of x that give $f(x) = 4$, such as 2, 7, 12, and so on, up to 97. These values are exactly five apart from each other, forming an arithmetic progression with a common difference of 5, which is the period of the function.

To visualize this, we can draw a graph with the x-axis representing the input values and the y-axis representing the amplitudes. We observe that the amplitudes are zero for most values of x , except for a few specific values that correspond to the period.

Now, to extract the period from this superposition, we use a technique called Fourier sampling. We shift the periodic superposition so that the first non-zero amplitude is at zero. This shift also affects the rest of the amplitudes, which now become multiples of the period.

Mathematically, the shifted superposition can be represented as a sum of terms, where each term corresponds to a non-zero amplitude. The number of non-zero amplitudes is M/R , where M is the total number of possible inputs and R is the period. The amplitude of each term is given by the square root of R/M .

When we perform Fourier sampling on the first register, the output we observe is a random multiple of the period. For example, we might see an output of 60 or 80. To determine the period, we find the greatest common divisor (GCD) of these outputs. In our example, the GCD of 60 and 80 is 20, which gives us M/R . Finally, we calculate R by dividing M by the GCD, which in this case is $100/20 = 5$.

This outlines the process of solving period finding using Shor's Quantum Factoring Algorithm. The quantum circuit for implementing this algorithm is similar to the circuit for Simon's algorithm. It involves applying the quantum Fourier transform to the input register, followed by applying the function $f(x)$ and performing measurements. The principle of deferred measurement states that the measurement on the second register can be done at any point without affecting the final result, as long as there is no further communication between the qubits.

Period finding is a crucial step in Shor's Quantum Factoring Algorithm. By utilizing quantum superposition and Fourier sampling, we can efficiently determine the period of a given function, which has significant implications for factoring large numbers and cryptography.

In the context of quantum information, one of the fundamental algorithms is Shor's Quantum Factoring Algorithm. This algorithm is designed to find the period of a function, which is a crucial step in factoring large numbers. The period finding process involves a series of measurements and transformations on quantum bits, or qubits.

To understand the algorithm, let's break it down step by step. First, we start with a function $f(x)$ that takes an input x and produces an output $f(x)$. The goal is to find the period, denoted as R , of this function.

The algorithm begins by preparing a superposition of all possible inputs using quantum Fourier transform. This superposition is then measured, resulting in a measurement outcome $f(a)$, where a is a randomly chosen input. At this point, the qubits are in a periodic superposition state.

Next, another quantum Fourier transform is applied to the qubits, followed by another measurement. This measurement is called quantum Fourier sampling. It has two important properties: shifting the input does not change the output distribution, and the output of the measurement is $f(0)$ when the superposition is shifted to zero.

To find the period, the circuit is repeated several times, collecting measurement outcomes. These outcomes are used to compute the greatest common divisor (GCD) of the measurements. The GCD gives us the period R .

Now, let's consider the case where R does not divide the number we are factoring, denoted as N . In this scenario, we make the assumption that N is large enough, specifically larger than $2R^2$. This ensures that the number of periods we look at is comparable to or larger than the period itself.

In this case, the quantum circuit remains the same. We follow the same steps, applying quantum Fourier transform and quantum Fourier sampling. The output of the measurement is denoted as L . The key insight is that L divided by N is approximately equal to T divided by R , where T is an unknown integer.

To determine both T and R from L and N , we use the fact that T divided by R is the best approximation to L divided by N with a denominator as small as R . Since R is much smaller than the square root of N , we can efficiently reconstruct T divided by R using a technique called continued fractions. This can be done on a classical computer, allowing us to solve the period finding problem even when R does not divide N .

Shor's Quantum Factoring Algorithm is a powerful tool for finding the period of a function, which is crucial in factoring large numbers. By leveraging quantum Fourier transform and quantum Fourier sampling, the algorithm can efficiently determine the period, even in cases where the period does not divide the number being factored.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: SHOR'S QUANTUM FACTORING ALGORITHM****TOPIC: SHOR'S FACTORING ALGORITHM**

Shor's factoring algorithm is one of the most famous quantum algorithms. It is used to solve the factoring problem, which involves finding the prime power factorization of a given number. For example, if we have the number 60, its prime factorization is $2^2 * 3 * 5$. In general, we want to write a number n as a product of prime powers, where $p_1, p_2, \dots, p_k$ are the prime factors and $e_1, e_2, \dots, e_k$ are their respective powers.

The most interesting and difficult case of the factoring problem is when n is a product of two large primes, P and Q , which are roughly equal in length. This is the case used in the RSA cryptosystem, a widely used public key cryptographic system for secure communication. The security of the RSA cryptosystem is based on the difficulty of factoring large numbers. Scientists and mathematicians have spent decades trying to efficiently solve this problem, but the best known classical algorithms still take exponential time, with a complexity of $O(2^n)$ or $O(n^{3/2})$.

Quantum computers offer a potential solution to this problem. Shor's algorithm takes advantage of the properties of quantum mechanics to efficiently factor large numbers. To understand how it works, we need to explore some elementary modular arithmetic. In modular arithmetic, we say that a is congruent to $b \bmod n$ if the remainder of b divided by n is a . For example, $24 \bmod 21$ is 3, and $-1 \bmod 21$ is 20.

Modular arithmetic allows us to perform operations like addition and multiplication efficiently. For example, $(24 + 35) \bmod 21$ is equivalent to $(3 + 14) \bmod 21$, which is 17. Similarly, $(24 * 30) \bmod 21$ is equivalent to $(3 * 9) \bmod 21$, which is 6. This means that we can perform arithmetic operations modulo n quickly and easily.

Another important concept in number theory is the greatest common divisor (GCD) of two numbers. The GCD is the largest number that divides both numbers without leaving a remainder. Classically, we can compute the GCD efficiently using algorithms like Euclid's algorithm, which involves repeatedly dividing the larger number by the smaller number and taking the remainder until we reach a remainder of 0. The GCD is then the last nonzero remainder.

Shor's factoring algorithm is a groundbreaking quantum algorithm that can efficiently factor large numbers. It takes advantage of the properties of quantum mechanics, such as modular arithmetic and efficient computation of greatest common divisors, to solve the factoring problem. This algorithm has important implications for cryptography and the security of systems like the RSA cryptosystem.

To understand Shor's Quantum Factoring Algorithm, we need to first understand the concept of square roots modulo a number. In this algorithm, we are interested in finding non-trivial square roots of 1 modulo a given number, which in this case is 21.

To find these square roots, we start by solving the equation $x^2 \equiv 1 \pmod{21}$. Here, x represents the number we are looking for, and the congruence relation indicates that when we square x and reduce the result modulo 21, we should get 1.

One obvious solution to this equation is $x = 1$, as $1^2 \equiv 1 \pmod{21}$. Another solution is $x = -1$, which is equivalent to 20 (mod 21). When we square -1, we get 400, which is also congruent to 1 (mod 21). This shows that for any number n , $(n-1)^2 \equiv 1 \pmod{n}$.

Surprisingly, there is another number that satisfies this equation for 21, which is $x = 8$. When we square 8, we get 64, and reducing it modulo 21 gives us 1. So, 8 is a non-trivial square root of 1 modulo 21.

Now, let's see how this information helps us factorize 21. We can express $8^2 \equiv 1 \pmod{21}$ as $(8^2 - 1^2) \equiv 0 \pmod{21}$. This means that 21 divides $(8 + 1)(8 - 1)$, but it does not divide either of the factors individually.

To find the prime factors of 21, we compute the greatest common divisor (GCD) of 21 and the factors separately. The GCD of 21 and $8 + 1$ (which is 9) is 3, and the GCD of 21 and $8 - 1$ (which is 7) is 7. These are the prime factors of 21, and we have successfully factorized it.

It's important to note that 8 is not the only non-trivial square root of 1 modulo 21. Another non-trivial square root is -8, which is equivalent to 13 (mod 21). When we square 13, we get 169, which is congruent to 1 (mod 21). So, both 8 and 13 are non-trivial square roots of 1 modulo 21.

In general, if we can find a number x such that x is not congruent to $\pm 1 \pmod{n}$, but $x^2 \equiv 1 \pmod{n}$, then we can factorize n . This is because n divides $(x + 1)(x - 1)$, but not either of the factors individually. By computing the GCD of n and either of the factors, we can recover the prime factors of n .

To discover such a non-trivial factor, we can use a method called repeated squaring. We start with a random number x , and compute x^0 , x^1 , x^2 , and so on, reducing the results modulo n . We continue this process until we find a power of x that is congruent to 1 (mod n). This power of x will be a non-trivial square root of 1 modulo n , and we can use it to factorize n .

For example, if we take $n = 21$ and $x = 2$, we can create a table to calculate the powers of 2 modulo 21:

$2^0 \equiv 1 \pmod{21}$
 $2^1 \equiv 2 \pmod{21}$
 $2^2 \equiv 4 \pmod{21}$
 $2^3 \equiv 8 \pmod{21}$
 $2^4 \equiv 16 \pmod{21}$
 $2^5 \equiv 11 \pmod{21}$
 $2^6 \equiv 1 \pmod{21}$

From this table, we can see that $2^6 \equiv 1 \pmod{21}$. This implies that $(2^3)^2 \equiv 1 \pmod{21}$. Therefore, 2^3 (which is 8) is a non-trivial square root of 1 modulo 21.

By using this method, we can find non-trivial square roots of 1 modulo a given number, which can be used to factorize the number and find its prime factors.

In the field of quantum information, one of the most remarkable algorithms is Shor's Quantum Factoring Algorithm. This algorithm provides a way to efficiently factorize large numbers, which is a problem of great importance in cryptography.

The key idea behind Shor's algorithm is to exploit the quantum properties of superposition and entanglement to find the period of a periodic function. In the case of factoring, the function in question is the modular exponentiation function, which calculates the remainder when a number is raised to a power and divided by another number.

To understand how Shor's algorithm works, let's consider an example. Suppose we want to factorize the number 21. We start by picking a random number, let's say 2, and calculate its powers modulo 21. We create a table with two columns: one for the powers of 2 (from 0 to $n-1$) and another for the results of the modular exponentiation function.

As we calculate the powers of 2 modulo 21, we notice that the function values repeat after a certain number of steps. In our example, the function values repeat after 6 steps. This period, denoted as R , is the key to factorizing the number.

If we are lucky, the period is even and we can find a non-trivial square root of 1 modulo 21. This non-trivial square root leads us to the factors of 21. However, even if the period is odd, we can still find a non-trivial factor of 21 by computing the greatest common divisor of the function values and 21. If the greatest common divisor is not 1, then we have found a non-trivial factor.

Now, you may wonder how a quantum computer helps us in this process. Shor's algorithm leverages the power of quantum computing to efficiently find the period. It creates a superposition over all possible values of the exponent in the modular exponentiation function. This superposition is represented as a matrix, where each row corresponds to a different exponent value.

After creating the superposition, the quantum computer performs a measurement on a second register, collapsing the superposition to a specific value. This collapsed value corresponds to a certain row in the matrix,

which represents a specific exponent value. By performing classical computations on the collapsed value, we can determine the period of the function.

Once we have the period, we can check if we were lucky and find the factors of the number. If we were not lucky, we can repeat the process and try again. The probability of success in each trial is at least $1/2$, so we don't need to perform too many trials.

It's worth noting that the size of the period can be exponentially large compared to the input number. However, this is not a problem for a quantum computer, as it can efficiently find the period in polynomial time.

Shor's Quantum Factoring Algorithm is a groundbreaking algorithm that harnesses the power of quantum computing to efficiently factorize large numbers. By exploiting the quantum properties of superposition and entanglement, Shor's algorithm can find the period of a periodic function, which leads to the factors of the number. This algorithm has significant implications for cryptography and has the potential to break many commonly used encryption schemes.

In the field of Quantum Information, one of the most significant breakthroughs is Shor's Quantum Factoring Algorithm. This algorithm allows for efficient factorization of large numbers, which has important implications for cryptography and computational complexity.

The key idea behind Shor's algorithm is to exploit the quantum properties of superposition and entanglement to find the period of a function. By finding the period, we can then deduce the factors of a given number. The algorithm works by utilizing a quantum Fourier transform and classical post-processing.

To begin, we choose a number 'n' that we want to factorize. We don't know the period of the function, but we know that it is smaller than 'n'. We then select a value 'm' that is larger than 2 times the square of 'n'. This value of 'm' ensures that we have enough copies of the period to obtain accurate results.

The circuit for the factoring algorithm consists of the following steps. First, we initialize all qubits to the state of zero. Then, we apply the M by M quantum Fourier transform, where 'M' is the chosen value that is larger than 2 times 'n' squared. This transform prepares the qubits in a superposition of all possible states.

Next, we apply a function 'F' that performs modular exponentiation. This function can be efficiently computed using classical methods. The result of this function is stored in the second register.

After applying the function, we measure the results of the second register. This measurement provides us with a sample value. We then perform classical post-processing on the first register to check if we obtained a non-trivial square root. If we did, we have successfully factorized 'n'. If not, we repeat the process until we obtain the desired result.

Shor's Quantum Factoring Algorithm is a groundbreaking method that leverages quantum properties to efficiently factorize large numbers. By utilizing the quantum Fourier transform and classical post-processing, this algorithm has the potential to revolutionize cryptography and computational complexity.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: SHOR'S QUANTUM FACTORING ALGORITHM****TOPIC: QFT CIRCUIT**

The quantum Fourier transform (QFT) is a key component of Shor's quantum factoring algorithm, which provides an exponential advantage over classical circuits in factoring large numbers. In order to understand how the QFT works, it is important to know how the quantum circuit for the QFT is implemented.

The QFT is defined by the equation Ω^{jk} , where Ω is an n th root of unity and j and k are integers. In this lecture, we consider the case where the dimension of the QFT, denoted as M , is a power of 2 ($M = 2^n$). It is worth noting that Ω^2 is an $M/2$ th root of unity and is a primitive $M/2$ th root of unity, meaning $\Omega^2 = e^{(2\pi i/m/2)}$.

The target circuit for the QFT takes as input the state of n qubits and outputs the transformed state of these qubits. The circuit can be visualized as follows: we leave off the least significant qubit and perform a QFT on the remaining $n-1$ qubits. This QFT is implemented using a sequence of gates, including controlled phase rotations on each qubit controlled by the least significant qubit, and a Hadamard gate on the last qubit.

The size of the circuit for an M -qubit circuit, denoted as $s(M)$, can be determined using a recurrence relation. The circuit size satisfies the equation $s(M) = M + s(M-1)$, where $s(M-1)$ is the circuit size for an $(M-1)$ -qubit circuit. Solving this recurrence relation yields $s(M) = M(n+1)/2$, which is approximately $M^2/2$. Since $\log n$ is equal to $\log(M)$, the circuit size is also proportional to $\log^2(M)$.

The form of the QFT circuit can be understood by examining the classical circuit for the fast Fourier transform (FFT). The FFT circuit is a classical implementation of the Fourier transform and provides insight into how the QFT circuit is constructed. The FFT circuit divides the Fourier transform matrix into two parts vertically and divides the columns into even and odd parts. By numbering the rows and columns accordingly, the FFT circuit efficiently computes the Fourier transform.

The QFT circuit is implemented by performing a QFT on $n-1$ qubits and applying a sequence of gates, including controlled phase rotations and a Hadamard gate. The circuit size is proportional to $M^2/2$, where M is the dimension of the QFT. Understanding the classical FFT circuit provides insight into the construction of the QFT circuit.

In the context of quantum information, Shor's Quantum Factoring Algorithm is a significant breakthrough. One key component of this algorithm is the Quantum Fourier Transform (QFT) circuit. In order to understand the QFT circuit, it is important to first understand the concept of the Fourier transform.

The Fourier transform is a mathematical operation that decomposes a function into its constituent frequencies. In the context of quantum computing, the QFT circuit performs a similar function, but on quantum states instead of classical functions. It transforms a quantum state into its frequency representation.

The QFT circuit is implemented using a matrix, which we will refer to as the Fourier transform matrix. This matrix has entries that are complex numbers of the form Ω^{jk} , where Ω is a complex number and j and k are integers. The specific values of Ω and the dimensions of the matrix depend on the size of the input.

The QFT circuit can be divided into three main parts: the top submatrix, the bottom submatrix, and the phase corrections. The top submatrix is a smaller Fourier transform matrix, while the bottom submatrix is a similar matrix with a phase correction. The phase corrections are introduced to account for the multiplication of Ω to the power of j .

To apply the QFT circuit, the input vector is split into even and odd entries. The even entries are multiplied by the top submatrix, while the odd entries are multiplied by the bottom submatrix. The outputs are then combined using addition and subtraction operations, with appropriate phase corrections applied.

In terms of circuit implementation, the QFT circuit can be represented by a series of gates and operations. The input qubits are first divided into even and odd qubits. The Fourier transform circuit is then recursively applied

to the even and odd qubits. Finally, the outputs are combined using addition and subtraction gates, with appropriate phase corrections applied.

It is worth noting that the QFT circuit requires careful ordering of the input bits. The least significant bit is placed at the top, while the most significant bit is placed at the bottom.

The QFT circuit is a crucial component of Shor's Quantum Factoring Algorithm. It transforms a quantum state into its frequency representation using a Fourier transform matrix. The circuit involves dividing the input into even and odd entries, applying the Fourier transform recursively, and combining the outputs with appropriate phase corrections.

In the field of quantum information, Shor's Quantum Factoring Algorithm is a significant breakthrough. It provides a method for efficiently factoring large numbers using quantum computers. One key component of this algorithm is the Quantum Fourier Transform (QFT) circuit.

The QFT circuit is responsible for transforming the input state into a superposition of states, which is essential for the algorithm's success. It achieves this by applying a series of gates to the input qubits. The number of gates in the circuit is approximately twice the number of qubits, or $O(N)$, where N is the number of qubits.

The solution to the recurrence relation that describes the number of gates in the circuit is $O(N \log N)$. This is a significant improvement compared to classical algorithms, which would require $O(N^2)$ steps for the same task.

To understand the QFT circuit, it is helpful to consider its connection to the classical Fourier transform. The QFT circuit performs a similar transformation but in reverse order, with the least significant bit considered the most significant on the output. The circuit applies a transformation denoted as an M -dimensional transformation, where M is the number of qubits, and includes a normalization factor of $1/\sqrt{N}$.

The circuit should include gates that represent addition and multiplication by complex factors, denoted as Ω . However, in practice, the circuit can be simplified. To apply the QFT circuit to both the even and odd inputs simultaneously, we can omit the least significant qubit from the circuit. Quantum mechanics takes care of the rest, applying the QFT to the remaining qubits in superposition.

To add up the corresponding bits J and $N/M + J$, we can apply a controlled gate to the qubit left out of the QFT circuit. This gate applies a phase correction of Ω to the J th qubit if the output qubit is 0, and subtracts Ω if the output qubit is 1. The controlled gate also handles the necessary normalization.

For the lower half of the circuit, where the output qubit is 1, we want to apply a phase correction of Ω to the J th qubit. To achieve this, we write J in binary form and use a sequence of one or two-qubit gates. Each bit in the binary representation corresponds to a phase correction of Ω to a certain power.

The QFT circuit is a crucial component of Shor's Quantum Factoring Algorithm. It efficiently transforms the input state into a superposition of states, enabling the algorithm to factor large numbers. By utilizing quantum parallelism, the QFT circuit achieves significant computational savings compared to classical algorithms.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: GROVER'S QUANTUM SEARCH ALGORITHM****TOPIC: NEEDLE IN A HAYSTACK**

Grover's algorithm is an important quantum algorithm used for unstructured search. It can be thought of as searching for a needle in a haystack. In the digital equivalent, the haystack is represented by a large table with n entries. Each entry can be accessed and examined individually. The goal is to find the one marked entry, which represents the needle.

Classically, searching through the entries in random order would take an expected time of $n/2$. However, quantum mechanics offers the potential for a much faster and more efficient solution. The hope is that quantum mechanics can provide a clever and efficient way to search through the haystack using its exponential power.

Searching for a needle in a haystack is an important problem because it is related to a class of problems called NP-complete problems. These problems have significant computational implications across various disciplines, including computer science, physics, and chemistry. One example of an NP-complete problem is satisfiability, where a boolean formula on n variables needs to be satisfied by assigning values of 0 or 1 to the variables. There are 2^n possible configurations, and the goal is to find the one configuration that satisfies the formula.

Grover's algorithm can solve the satisfiability problem in more than n time, but in square root of n time. This represents a quadratic speed-up compared to classical algorithms. However, it is important to note that even with this speed-up, the algorithm still has an exponential time complexity for satisfiability problems.

Formally, in the quantum setting, we are given a boolean function from 0 to $n-1$, and the goal is to find an x such that $f(x) = 1$. The hardest case is when there is exactly one satisfying x . The function is typically given in the form of a circuit or an oracle, which takes an input x and outputs $f(x)$.

Grover's algorithm is a powerful quantum algorithm that can be used for unstructured search problems, such as finding a needle in a haystack. It offers a quadratic speed-up compared to classical algorithms for certain problems, but it still has an exponential time complexity for problems like satisfiability.

In the field of quantum information, one fundamental concept is Grover's Quantum Search Algorithm. This algorithm is designed to solve the problem of finding a specific item, also known as the "needle in a haystack" problem, in an unsorted database.

In classical computing, searching through an unsorted database requires checking each item one by one until the desired item is found. This process can be time-consuming, especially for large databases. However, with the power of quantum computing, Grover's algorithm provides a significant speedup.

The algorithm takes an input of n bits and outputs a single bit. Quantumly, we can create a quantum circuit for a function f , which takes as input x (a bunch of zeros) and outputs $f(x)$ (a bunch of zeros). The key advantage is that we can evaluate $f(x)$ in superposition, meaning we can simultaneously evaluate multiple inputs.

By applying Grover's algorithm, we can efficiently search through the database to find the desired item. The algorithm uses a combination of quantum operations, such as the Hadamard transform, the oracle function, and the Grover diffusion operator, to amplify the amplitude of the desired item and suppress the amplitudes of the other items.

The oracle function is a crucial component of Grover's algorithm. It marks the desired item by flipping the sign of its amplitude, while leaving the other items unchanged. This step is essential for the algorithm to converge towards the desired item.

The Grover diffusion operator acts as a reflection across the mean amplitude, which helps in redistributing the amplitudes and increasing the probability of finding the desired item. By repeating the oracle and diffusion steps multiple times, the algorithm gradually converges towards the solution.

It's important to note that Grover's algorithm provides a quadratic speedup compared to classical algorithms, meaning it can find the desired item in approximately $\sqrt{N}$ iterations, where N is the size of the database. This

speedup is significant for large databases, making Grover's algorithm a valuable tool in quantum information processing.

Grover's Quantum Search Algorithm is a powerful tool for efficiently searching through unsorted databases. By leveraging the principles of quantum computing, the algorithm allows for simultaneous evaluation of multiple inputs and provides a quadratic speedup compared to classical algorithms. With its applications in various fields, Grover's algorithm plays a crucial role in the advancement of quantum information processing.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: GROVER'S QUANTUM SEARCH ALGORITHM****TOPIC: GROVER'S ALGORITHM**

Grover's algorithm is a quantum search algorithm that is used to find a specific entry in an unsorted database. The algorithm consists of two main steps: phase inversion and inversion about the mean.

In the phase inversion step, the algorithm maintains a superposition over all possible entries in the database. Initially, all the amplitudes are equal to $\frac{1}{\sqrt{n}}$ over the square root of n , where n is the number of entries in the database. The goal is to find the special entry, denoted as X^* . If X is not equal to X^* , the amplitude remains unchanged. However, if X is equal to X^* , the amplitude is inverted by multiplying it by -1 .

In the inversion about the mean step, the algorithm flips the amplitudes about the mean value. The mean value, denoted as μ , is the average of all the amplitudes. The amplitudes are flipped by mapping $f(X)$ to $2\mu - f(X)$. This operation ensures that amplitudes smaller than the mean are flipped up, while amplitudes larger than the mean are flipped down.

These two steps are repeated iteratively. Each iteration increases the amplitude of the special entry, X^* , and decreases the amplitudes of the other entries. The number of iterations required is approximately equal to the square root of n .

It is important to note that both the phase inversion and inversion about the mean steps are unitary transformations, meaning they can be implemented efficiently in a quantum system. The details of how these steps are implemented will be discussed in a separate material.

Grover's algorithm is a powerful quantum search algorithm that can find a specific entry in an unsorted database. It achieves this by iteratively applying phase inversion and inversion about the mean steps, increasing the amplitude of the desired entry and decreasing the amplitudes of the other entries.

In Grover's Quantum Search Algorithm, the goal is to efficiently find a marked element in an unsorted database of size n . The algorithm achieves this by using quantum parallelism and interference.

To understand the algorithm, let's consider an example where we have a database with n elements and only one of them is marked. Initially, all elements have the same amplitude, which is $\frac{1}{\sqrt{n}}$.

The algorithm starts by applying a phase inversion operation to the marked element. This operation flips the sign of the amplitude of the marked element, effectively "inverting" it. This step increases the amplitude of the marked element to about $\frac{1}{\sqrt{2}}$, while leaving the other elements unchanged.

Next, we apply an inversion about the mean operation. This operation reflects the amplitudes across the mean amplitude. This causes the amplitudes of the other elements to become negative, while the amplitude of the marked element remains positive. As a result, the amplitude of the marked element increases further.

By repeating these two steps, the amplitude of the marked element continues to increase, while the amplitudes of the other elements decrease. After roughly square root n steps, the amplitude of the marked element becomes close to $\frac{1}{\sqrt{2}}$, and the chance of measuring the marked element becomes about $\frac{1}{2}$.

Therefore, in roughly square root n steps, Grover's algorithm can find the marked element in the database. It achieves an improvement of about $\sqrt{2}$ over square root n in each step.

To rigorously justify this improvement, let's consider the amplitude distribution of the other elements when the marked element has an amplitude of $\frac{1}{\sqrt{2}}$. In this case, the remaining elements have an amplitude of at least $\frac{1}{\sqrt{2n}}$.

When we perform an inversion about the mean operation, the amplitude of the marked element is close to $\frac{1}{\sqrt{2n}}$. Since the other elements have amplitudes of at least $\frac{1}{\sqrt{2n}}$, the improvement per step is at least square root 2 over n .

By dividing the desired improvement of 1 over square root 2 by the improvement per step, we can determine the number of steps needed to reach 1 over square root 2 . This number is bounded by the square root of n over 2 .

Grover's algorithm can find the marked element in an unsorted database of size n in roughly square root n steps. The algorithm achieves an improvement of about 2 over square root n in each step.

Please note that this explanation is slightly approximate, as it assumes certain conditions. However, the rigorous analysis justifies the overall behavior of Grover's algorithm.

In the next material, we will explore how to implement the steps of Grover's algorithm.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: GROVER'S QUANTUM SEARCH ALGORITHM****TOPIC: IMPLEMENTING GROVER'S ALGORITHM**

Grover's algorithm is a quantum search algorithm that can be used to find a specific item in an unsorted database with quadratic speedup compared to classical algorithms. In order to implement Grover's algorithm, two main steps are repeatedly performed: phase inversion and inversion about the mean.

Phase inversion is the process of inverting the phase of the marked element, which is the element where the function $f(X)$ is equal to 1. To carry out phase inversion, we start with a superposition of all possible states, represented by the sum over X of $\alpha_X X$. The goal is to map this superposition to a state where each X is multiplied by $-1/2$ times $f(X)$. To achieve this, we replace the answer bit in the circuit with a minus state. By doing so, the effect of the function f is to put the desired phase in the right place. This simple modification allows us to perform phase inversion effectively.

Inversion about the mean is the process of transforming the amplitudes of the superposition by subtracting them from twice the mean amplitude. This step helps to amplify the amplitude of the marked element and decrease the amplitude of the other elements. To implement inversion about the mean, we start with the superposition sum over X of $\alpha_X X$. We then calculate the mean of all the amplitudes. Finally, we transform each amplitude α_X by subtracting it from twice the mean. This transformation effectively amplifies the amplitude of the marked element and reduces the amplitude of the other elements.

To summarize, Grover's algorithm consists of phase inversion and inversion about the mean steps. Phase inversion is achieved by replacing the answer bit with a minus state, while inversion about the mean is achieved by subtracting each amplitude from twice the mean. These steps allow us to effectively search for a specific item in an unsorted database.

In the context of quantum information, one important algorithm is Grover's Quantum Search Algorithm. This algorithm is used to search an unsorted database in a faster and more efficient way compared to classical search algorithms.

To understand how Grover's Algorithm works, let's first discuss the concept of inversion about the mean. Inversion about the mean is an operation that can be represented by a quantum state. This state is a sum of all possible states, denoted as $|X\rangle$. The goal of the algorithm is to find the state $|X\rangle$ that satisfies a certain condition.

To carry out the inversion about the mean, we first decompose the state $|X\rangle$ into two parts: one that aligns with a uniform superposition state, denoted as $|U\rangle$, and one that is orthogonal to $|U\rangle$. The orthogonal part represents the component of $|X\rangle$ that is not aligned with $|U\rangle$. We then take the negative of the orthogonal component to obtain a new vector.

To implement the inversion about the mean, we transform the uniform superposition state $|U\rangle$ into an all-zero vector, perform a reflection about the all-zero vector, and then transform back to the original state $|U\rangle$. The transformation that moves the uniform superposition to the all-zero vector is the Hadamard transform. The reflection about the all-zero vector is achieved by leaving the all-zero vector unchanged and multiplying everything orthogonal to it by -1 . Finally, we transform back using the inverse of the Hadamard transform.

Mathematically, the transformation can be represented as follows:

1. Transform $|U\rangle$ into the all-zero vector.
2. Perform a reflection about the all-zero vector.
3. Undo the transformation to obtain the final result.

The Hadamard transform is used to move the uniform superposition state to the all-zero vector. The reflection about the all-zero vector is achieved by multiplying everything orthogonal to it by -1 . The transformation back to the original state is done using the inverse of the Hadamard transform.

Now, let's analyze the algorithm in more detail. We can express the algorithm as a matrix operation:

$H \text{ tensor } n * (1 - 2|0\rangle\langle 0|) * H \text{ tensor } n,$

where H tensor n represents the Hadamard transform applied n times, and $(1 - 2|0\rangle\langle 0|)$ represents the reflection about the all-zero vector.

By simplifying the expression, we find that the resulting matrix is a diagonal matrix with diagonal entries of $2/n - 1$ and all other entries equal to $2/n$. This matrix has a size of $n \times n$, where n is the number of qubits.

To understand the effect of this matrix, let's consider an input state $|\alpha_0\rangle$ through $|\alpha_{n-1}\rangle$. When this matrix operates on the input state, each entry is multiplied by $2/n$, except for the diagonal entries, which are multiplied by $2/n - 1$.

Grover's Algorithm uses the inversion about the mean operation to efficiently search an unsorted database. The algorithm involves transforming the uniform superposition state into an all-zero vector, performing a reflection about the all-zero vector, and then transforming back to the original state. The resulting matrix, obtained through the Hadamard transform and reflection, has diagonal entries of $2/n - 1$ and all other entries equal to $2/n$.

In Grover's algorithm, the goal is to search an unsorted database of N elements to find a specific target element. This algorithm provides a quadratic speedup compared to classical search algorithms.

The quantum circuit for Grover's algorithm consists of several steps. First, we initialize the qubits to the state $|0\rangle$. Then, we apply a Hadamard gate to create a uniform superposition over all possible n -bit strings. This allows us to explore all possible states simultaneously.

Next, we perform a phase inversion operation. This operation involves flipping the sign of the target element's amplitude, effectively amplifying its probability. This is achieved by applying a phase flip gate to the target element.

After the phase inversion, we perform an inversion about the mean operation. This operation involves reflecting the amplitudes about the mean amplitude. This amplifies the probability of the target element even further. The inversion about the mean operation is achieved by applying a combination of Hadamard gates and phase flip gates.

It is important to note that during the phase inversion and inversion about the mean operations, an extra qubit remains unchanged. Therefore, no additional operations are required for this qubit.

One iteration of the algorithm consists of the initialization, phase inversion, and inversion about the mean operations. To increase the probability of finding the target element, we repeat this iteration square root of N times. This is based on the observation that approximately square root of N iterations are required for a successful search.

Grover's algorithm is a quantum search algorithm that provides a quadratic speedup compared to classical search algorithms. It involves initializing the qubits, creating a superposition, performing a phase inversion, and applying an inversion about the mean operation. By repeating these steps a certain number of times, we can significantly increase the probability of finding the target element in an unsorted database.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: OBSERVABLES AND SCHRODINGER'S EQUATION****TOPIC: INTRODUCTION TO OBSERVABLES**

An observable in quantum information refers to a quantity that can be measured, such as energy, position, or momentum. It is represented by a matrix, specifically a K by K Hermitian matrix, where K represents the dimensionality of the system. A Hermitian matrix is a matrix that is equal to its conjugate transpose.

To understand observables further, we can consider the spectral theorem, which states that a Hermitian matrix has an orthonormal set of eigenvectors and real eigenvalues. In other words, the Hermitian matrix can be diagonalized using an orthonormal basis. The eigenvectors correspond to the possible measurement outcomes, and the eigenvalues represent the values that can be obtained from the measurement.

When measuring a quantum state, we choose a basis in which to measure it. The measurement outcome is then determined by the probability amplitudes associated with each eigenvector in the chosen basis. The new state after measurement is obtained by projecting the original state onto the eigenvector corresponding to the measured outcome.

Let's consider an example to illustrate this concept. Suppose we have a single qubit state represented by the superposition $\alpha|0\rangle + \beta|1\rangle$. We want to measure this state using the observable X , which is represented by the matrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$. It is important to note that X is both a Hermitian matrix and a unitary transformation.

To perform the measurement, we first need to find the eigenvectors and eigenvalues of X . In this case, the eigenvectors are $|+\rangle = (1/\sqrt{2})(|0\rangle + |1\rangle)$ and $|-\rangle = (1/\sqrt{2})(|0\rangle - |1\rangle)$, with corresponding eigenvalues $+1$ and -1 , respectively.

Next, we express the original state in terms of the eigenvectors of X . This can be done by taking the inner product of the original state with the eigenvectors. For our example, the state in the $|+/-\rangle$ basis is $(\alpha + \beta)/\sqrt{2}|+\rangle + (\alpha - \beta)/\sqrt{2}|-\rangle$.

Finally, the outcome of the measurement is determined by the probabilities associated with each eigenvector. The probability of obtaining the $+1$ eigenvalue is $|(\alpha + \beta)/\sqrt{2}|^2$, and the probability of obtaining the -1 eigenvalue is $|(\alpha - \beta)/\sqrt{2}|^2$. The new state after measurement is the corresponding eigenvector, either $|+\rangle$ or $|-\rangle$, depending on the outcome.

Observables in quantum information are quantities that can be measured, represented by Hermitian matrices. They have associated eigenvectors and eigenvalues, which determine the measurement outcomes and probabilities. The new state after measurement is obtained by projecting the original state onto the eigenvector corresponding to the measured outcome.

When measuring an observable in quantum mechanics, the outcome of the measurement is determined by the probabilities associated with each eigenvalue. For example, if we have a measurement with two possible outcomes, 1 and -1 , the probability of obtaining the outcome 1 is given by the magnitude squared of the sum of two complex numbers, α and β , divided by the square root of 2 . Similarly, the probability of obtaining the outcome -1 is given by the magnitude squared of the difference between α and β , divided by the square root of 2 .

After the measurement, the state of the system changes. If the outcome is 1 , the new state is represented by the "plus" symbol, and if the outcome is -1 , the new state is represented by the "minus" symbol. The expected value of the measurement can be calculated by multiplying each outcome by its corresponding probability and summing them up. This calculation is similar to finding the average value of a quantity, such as momentum.

In some cases, there may be repeated eigenvalues for an observable. In a three-dimensional system, for example, if two eigenvectors have the same eigenvalue, any linear combination of those eigenvectors will also be an eigenvector with the same eigenvalue. When applying the observable operator to a linear combination of eigenvectors, the eigenvalue can be pulled out and multiplied by the linear combination. This simplifies the calculation.

When measuring an observable with repeated eigenvalues, the outcome is determined by projecting the state onto the corresponding eigenvectors. The probability of obtaining a specific outcome is given by the square of the length of the projection onto the eigenvector. After the measurement, the new state is the projection of the original state onto the subspace spanned by the eigenvectors with the repeated eigenvalue.

To illustrate this concept, let's consider the measurement of the identity operator, represented by the matrix with ones on the diagonal and zeros elsewhere. The outcome of this measurement is always 1, regardless of the initial state. The new state after the measurement is the projection of the original state onto the subspace spanned by the eigenvectors with eigenvalue 1.

When measuring an observable in quantum mechanics, the outcome is determined by the probabilities associated with each eigenvalue. The new state after the measurement is the projection of the original state onto the subspace spanned by the eigenvectors corresponding to the measured eigenvalue. If there are repeated eigenvalues, any linear combination of the corresponding eigenvectors will also be an eigenvector with the same eigenvalue.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: OBSERVABLES AND SCHRODINGER'S EQUATION****TOPIC: OBSERVABLES PROPERTIES**

An observable for a K -level system is represented by a K by K Hermitian matrix. When measuring the system, the measurement is done in the orthonormal basis of eigenvectors of the observable. The outcomes of the measurement are real eigenvalues. This new notion of an observable does not deviate from the previous notion of a measurement, which involved an arbitrary orthonormal basis.

It is possible to design an observable, represented by a K by K Hermitian matrix, with arbitrary eigenvectors and eigenvalues. This means that the new notion of a Hermitian operator is as general as the previous notion of a measurement. To illustrate this, let's consider an example.

Suppose we want the eigenvectors to be $(1/\sqrt{2})(0 + i)(1)$ and $(1/\sqrt{2})(0 - i)(1)$, with eigenvalues $+1$ and -1 , respectively. We can design an operator A that has these eigenvectors and eigenvalues. To do this, we can use the concept of a projection matrix.

A projection matrix, denoted as P , projects a state onto another state. It can be created using the bra-ket notation as $P = |\Phi\rangle\langle\Phi|$. The inner product of the state and the projected state gives the coefficient and the vector of the projection.

In our case, we can create the matrix A as the projection onto $|\Phi_1\rangle$ with coefficient λ_1 plus the projection onto $|\Phi_2\rangle$ with coefficient -1 . By applying A to $|\Phi_1\rangle$ and $|\Phi_2\rangle$, we can verify that A has the desired eigenvectors and eigenvalues.

In general, the matrix A can be constructed by taking the projections onto the desired eigenvectors with their respective eigenvalues as coefficients. By subtracting the projections, we obtain a Hermitian matrix with the desired eigenvectors and eigenvalues.

We have shown that an observable, represented by a Hermitian matrix, can be designed to have arbitrary eigenvectors and eigenvalues. This new notion of an observable is as general as the previous notion of a measurement.

In the field of quantum information, observables play a crucial role in understanding the behavior and properties of quantum systems. An observable is a physical quantity that can be measured or observed in an experiment. In this context, we are given a set of orthonormal vectors and real numbers, and our goal is to create the corresponding observable.

To create an observable, we take the corresponding linear combination of the projectors onto these vectors. This results in a Hermitian matrix, which is a square matrix that is equal to its own conjugate transpose. The eigenvectors of this Hermitian matrix are the possible outcomes of the measurement, while the eigenvalues correspond to the probabilities of obtaining each outcome.

To better understand this concept, let's consider an example. Let's say we have a set of orthonormal vectors labeled as Φ_j , and we want to find the corresponding observable. We can express the observable as a summation of projectors onto these vectors, denoted as P_j . The projection of Φ_j onto the other vectors Φ_i is zero, as they are orthogonal to each other. Therefore, the only non-zero projection is the one onto Φ_j , resulting in the eigenvalue λ_j multiplied by the eigenvector Φ_j .

By constructing the observable in this way, we ensure that each projector is a Hermitian matrix. Moreover, the sum of Hermitian matrices is also Hermitian. Therefore, we have successfully obtained the desired observable.

The key takeaway from this discussion is that there are two equivalent ways to specify a measurement in quantum information: either by specifying an orthonormal basis or by specifying an observable. Both approaches yield the same results and provide a comprehensive understanding of the measurement process.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: OBSERVABLES AND SCHRODINGER'S EQUATION****TOPIC: SCHRODINGER'S EQUATION**

The Schrodinger's equation is one of the most important equations in quantum mechanics. It arises from the axiom of unitary evolution, which states that a quantum system evolves according to a unitary rotation of the Hilbert space. The equation provides the fundamental quantum equation of motion.

Before discussing the Schrodinger's equation, it is necessary to understand a certain observable called the energy observable. The energy observable, also known as the Hamiltonian of the system, plays a central role in quantum mechanics. It is represented by a Hermitian matrix H , with eigenvectors $|\psi_i\rangle$ and eigenvalues λ_i . The eigenvectors represent the states of the system with definite energy, while the eigenvalues represent the corresponding energies.

For example, let's consider a Hamiltonian H with two eigenvectors $|+\rangle$ and $|-\rangle$, and eigenvalues 2 and -3, respectively. This means that if the state of the system $|\psi\rangle$ is $|+\rangle$, the energy measurement will always yield -2. Similarly, if the state is $|-\rangle$, the energy measurement will always yield -3. However, if the state is a superposition of the two eigenstates, such as $|\psi\rangle = (1/\sqrt{2})|+\rangle + (1/\sqrt{2})|-\rangle$, the energy measurement will yield +2 with a probability of 1/2 and -3 with a probability of 1/2.

In the case of a hydrogen atom, the energy eigenstates are represented by the basis states $|0\rangle$, $|1\rangle$, $|2\rangle$, and so on. The Hamiltonian of the system can be written as a diagonal matrix in this basis, with the energies of the corresponding states on the diagonal.

The Schrodinger's equation is a differential equation that relates the state of the system at time $t=0$ to the state at a later time t . It is given by $i\hbar(d|\psi\rangle/dt) = H|\psi\rangle$, where i is the square root of -1 and $\hbar$ is the reduced Planck constant. The solution to the Schrodinger's equation allows us to determine the state of the system at any given time.

The Planck constant, denoted by h , plays a central role in quantum mechanics. It relates the energy of a photon to its frequency through the Planck relation $E = h\nu$. The value of the Planck constant is approximately 6.626×10^{-34} joule-seconds, while the value of the reduced Planck constant $\hbar$ is approximately 1.055×10^{-34} joule-seconds or 6.582×10^{-16} electron volt-seconds. In some units, it is convenient to set $h/2\pi$ equal to 1.

The Schrodinger's equation introduces a constant H in the equation, which represents the rate of change of the state. Solving the Schrodinger's equation allows us to determine the time evolution of the quantum system.

The Schrodinger's equation is a fundamental equation in quantum mechanics that describes the time evolution of a quantum system. It relates the state of the system at time $t=0$ to the state at a later time t , given the Hamiltonian of the system. The equation involves the Planck constant, which plays a central role in quantum mechanics.

A differential equation with an operator in it is known as Schrodinger's equation. This equation describes the evolution of a quantum state over time. In particular, it tells us how the state changes and how the phase in front of the state evolves.

If we start in an eigenstate of the Hamiltonian (H), denoted as Φ_j , where Φ_j is an eigenvector of H with eigenvalue λ_j , then we can solve Schrodinger's equation. The solution has a simple form: the state at time T is given by $e^{-i\lambda_j T/\hbar} \Phi_j$.

What this means is that the eigenstate evolves in time by just changing the phase in front of the state. The rate of change of the phase, or the rate at which it precesses, is proportional to the eigenvalue (or energy) of the state. If the energy is 0, the phase does not precess at all. If the energy is high, the phase precesses quickly.

To show this, we consider a state that is an eigenstate of H (Φ_j). The right-hand side of Schrodinger's equation simplifies to $\lambda_j \Phi_j$. This tells us that the rate of change of the state is proportional to the state itself. Therefore, even after the change, the state still points in the direction of Φ_j , albeit with a different constant in front of it.

We can conclude that the state at all times is of the form some constant (which depends on time) times the eigenstate Φ_j .

To determine how the complex number $a_j(t)$ (which depends on time) evolves, we substitute the form of the state back into Schrodinger's equation. After some algebraic manipulation, we find that the derivative of $a_j(t)$ with respect to t is equal to $a_j(t) \times \frac{E_j}{\hbar}$. Integrating both sides of this equation, we obtain $a_j(t) = a_j(0) e^{-i E_j t / \hbar}$.

Schrodinger's equation tells us that the eigenstates of H evolve in time by changing the phase in front of them. The rate of change of the phase is proportional to the energy of the state. The state at any time is given by a constant times the eigenstate. The evolution of the state over time is described by an operator, denoted as $U(t)$, which applies to the state at time 0.

In quantum information, observables play a crucial role in determining the state of a system. One way of representing observables is through the use of operators. An operator $U(t)$ can be written as $e^{-i H t / \hbar}$, where $\hbar$ is Planck's constant and t represents time. This notation allows us to express a matrix B as $B = \sum_i | \Phi_i \rangle \langle \Phi_i | B | \Phi_i \rangle \langle \Phi_i |$, where B is a Hermitian operator with eigenvalues λ_i and eigenvectors Φ_i . In this case, B is a matrix whose eigenvectors are Φ_i and the corresponding eigenvalues are λ_i .

To better understand this concept, let's consider an example. Suppose we start in the state $|0\rangle$ and our Hamiltonian is represented by the matrix X , which is $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$. We want to determine the state of the system at time T . To do this, we first need to find the eigenvectors and eigenvalues of the Hamiltonian. For matrix X , the eigenvectors are $|+\rangle$ and $|-\rangle$, with eigenvalues 1 and -1 respectively. These eigenvectors represent states of definite energy.

Next, we need to express the initial state ($|0\rangle$) in terms of the eigenstates of the Hamiltonian. In this case, the initial state can be written as $\frac{1}{\sqrt{2}}(|+\rangle + |-\rangle)$. Now, we can determine the state at time T using the equation $\Phi(t) = \frac{1}{\sqrt{2}} e^{-i E_+ t / \hbar} |+\rangle + \frac{1}{\sqrt{2}} e^{-i E_- t / \hbar} |-\rangle$. Simplifying this expression, we get $\Phi(t) = \frac{1}{2} e^{-i E_+ t / \hbar} (e^{i E_- t / \hbar} + e^{-i E_+ t / \hbar}) |+\rangle + \frac{1}{2} e^{-i E_+ t / \hbar} (e^{i E_- t / \hbar} - e^{-i E_+ t / \hbar}) |-\rangle$.

For example, if we take T equal to $\pi \hbar / 2$, we can calculate the state of the system. Plugging in the values, we get $\frac{1}{\sqrt{2}} e^{-i \pi / 2} (e^{i \pi / 2} + e^{-i \pi / 2}) |+\rangle + \frac{1}{\sqrt{2}} e^{-i \pi / 2} (e^{i \pi / 2} - e^{-i \pi / 2}) |-\rangle$. Simplifying further, we find that the state at this time is $-i(|+\rangle + |-\rangle)$. This means that in this time period, the state of the system has transitioned from $|0\rangle$ to $|1\rangle$.

Observables in quantum information can be represented using operators. The operator $U(t)$, written as $e^{-i H t / \hbar}$, allows us to express matrices in terms of eigenvalues and eigenvectors. By finding the eigenvectors and eigenvalues of a given Hamiltonian, we can determine the state of a system at a specific time. This understanding of observables and Schrodinger's equation is fundamental to the study of quantum information.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO IMPLEMENTING QUBITS****TOPIC: CONTINUOUS QUANTUM STATES**

In this lecture, we will discuss the implementation of qubits and how they can be represented using the ground and excited states of an electron in a hydrogen atom. To understand this, we will use a simple toy model for a hydrogen atom.

The main force acting on the electron is the Coulomb attraction due to the proton. We will consider this force as a radial force, with the radial distance of the electron from the proton being the main variable of interest. To simplify the problem, we will treat it as a one-dimensional problem, where the electron is confined to a certain distance from the proton.

In this one-dimensional model, we can think of the electron as being free to move along a line segment of length L . Our goal is to describe the state of the electron and determine its Hamiltonian and energy eigenstates.

Since the electron can be anywhere on this line, we need to find a way to describe its state. We have been working with quantum states that are superpositions of a finite number of possibilities. However, in this case, the electron can be anywhere within the segment between 0 and L . To describe its state, we can consider the electron as being in a superposition of states that are multiples of some small interval Δ .

We can represent the state of the electron as a superposition of states labeled by J , ranging from $-K$ to K , where K is a large number. The state is described as $\Psi = \sum (\alpha_{\text{sub}J} * J * \Delta)$, where $\alpha_{\text{sub}J}$ is a coefficient and J is the label of the state.

To ensure that the state is normalized, the sum of the squared coefficients from $-K$ to K must equal 1.

Now, let's consider what happens as we let Δ tend to 0 and K tend to infinity. In this limit, we can think of the state as a continuous function $\Psi(X)$, where X represents the position of the electron. We can think of $\Psi(X)$ as $\alpha_{\text{sub}K} * \Delta$ when $J * \Delta = X$.

To have a normalized state, we require that the integral of the magnitude squared of $\Psi(X)$ over all X is equal to 1. This can be written as the integral from $-\infty$ to ∞ of $\Psi^*(X) * \Psi(X) dX$.

We can also calculate the probability of finding the electron at a particular position. The probability of finding the electron at position $J * \Delta$ is given by the magnitude squared of $\alpha_{\text{sub}J} * \Delta$.

We have discussed the implementation of qubits using the ground and excited states of an electron in a hydrogen atom. We have considered a simplified one-dimensional model, where the electron is confined to a certain distance from the proton. We have described the state of the electron and its Hamiltonian, and we have seen how the quantization of energy levels naturally emerges. Finally, we have discussed how to implement qubits.

In the study of quantum information, one fundamental concept is the implementation of qubits, which are the basic units of quantum information. Qubits can exist in a superposition of states, allowing for the representation and manipulation of complex information.

One important aspect of implementing qubits is understanding continuous quantum states. In this context, we consider the probability of finding an electron at a particular point in space. To calculate this probability, we define two points, Y and Z , and examine how the sine of X behaves between these points.

The probability of finding the electron between points Y and Z is given by the integral from Y to Z of $\sin^2(X)$ of X DX . This integral can also be expressed as the integral from Y to Z of the magnitude of $\sin(X)$ squared DX .

This calculation allows us to determine the likelihood of finding the electron within a specific range of points in space. By understanding the behavior of continuous quantum states, we can gain insights into the distribution of quantum information and make predictions about the behavior of qubits.

The implementation of qubits involves considering continuous quantum states and calculating probabilities based on the behavior of quantum wavefunctions. This understanding is crucial for the development and application of quantum information technologies.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO IMPLEMENTING QUBITS****TOPIC: SCHRODINGER'S EQUATION FOR A 1D FREE PARTICLE**

The Schrodinger equation is a fundamental equation in quantum mechanics that describes the behavior of quantum systems. In this didactic material, we will focus on understanding the Schrodinger equation for a free particle in one dimension.

A free particle refers to a particle that is not subject to any external forces or potentials and is allowed to move freely in space. The wave function of a free particle, denoted as $\Psi(x, t)$, describes the amplitude of the particle's position at a given time. The wave function evolves with time and is represented as $\Psi(x, t) = A \sin(xt)$, where A is a constant.

The Schrodinger equation for a free particle is given by:

$$i\hbar(\partial\Psi/\partial t) = -(\hbar^2/2m)(\partial^2\Psi/\partial x^2)$$

In this equation, i represents the imaginary unit, $\hbar$ is the reduced Planck's constant, m is the mass of the particle, and $\partial/\partial t$ and $\partial^2/\partial x^2$ represent the partial derivatives with respect to time and position, respectively.

To understand the meaning of this equation, let's break it down. The term on the left-hand side, $i\hbar(\partial\Psi/\partial t)$, represents the time derivative of the wave function multiplied by the imaginary unit and the reduced Planck's constant. This term describes how the wave function changes with time.

The term on the right-hand side, $-(\hbar^2/2m)(\partial^2\Psi/\partial x^2)$, represents the second derivative of the wave function with respect to position multiplied by the negative of the reduced Planck's constant squared divided by twice the mass of the particle. This term corresponds to the kinetic energy of the particle.

By equating the two sides of the equation, we are stating that the rate of change of the wave function with respect to time is proportional to the second derivative of the wave function with respect to position.

Intuitively, we can understand the Schrodinger equation for a free particle by considering the local behavior of the wave function. The evolution of the wave function should be influenced by its neighboring values. If we imagine the particle looking at its immediate neighborhood, it compares its own value to the average of its neighbors and adjusts accordingly. This adjustment is proportional to the difference between the particle's value and the average of its neighbors.

The Schrodinger equation also arises from the Hamiltonian, which represents the observable corresponding to energy. For a free particle, the Hamiltonian is given by $-(\hbar^2/2m)(\partial^2/\partial x^2)$. The second derivative term represents the kinetic energy of the particle, as derived from classical mechanics.

In quantum mechanics, momentum is represented by the momentum operator, denoted as P . The momentum operator for a particle in one dimension is given by $P = -i\hbar(\partial/\partial x)$. The Hamiltonian for a free particle can be derived from the momentum operator by squaring it and dividing by twice the mass.

The Schrodinger equation for a free particle in one dimension describes the evolution of the particle's wave function with respect to time. It relates the time derivative of the wave function to the second derivative of the wave function with respect to position. The equation arises from considering the local behavior of the wave function and the kinetic energy of the particle.

In quantum information, one of the fundamental concepts is the implementation of qubits. To understand this, we need to delve into Schrödinger's equation for a 1D free particle.

The momentum operator, denoted as p , can be obtained from the Hamiltonian by taking the derivative of the wave function with respect to position. Mathematically, it can be represented as $p = -i\hbar(d/dx)$, where $\hbar$ is the reduced Planck's constant. This equation gives us the momentum operator in terms of the position operator.

To gain a deeper understanding of the momentum operator, let's consider the discretization of space. We can

imagine dividing the line into discrete points, allowing the particle to occupy positions such as 0, Δ , $-\Delta$, 2Δ , and so on. In this context, the discrete analog of the momentum operator, denoted as $\Delta p/\Delta x$, can be defined as the symmetric difference quotient:

$$\Delta p/\Delta x = (\psi(x+\Delta x) - \psi(x-\Delta x))/\Delta x.$$

To represent this operator as a matrix, we can consider the wave function $\psi(x)$ as a vector, where each entry corresponds to a specific position. For example, if we have K discrete points, the vector would be $\psi = [\psi(-K\Delta), \psi(-(K-1)\Delta), \dots, \psi(K\Delta)]$.

The matrix corresponding to the momentum operator would have the following form: a diagonal matrix with zeros everywhere except for -1 's below the diagonal and 1 's above the diagonal. This matrix ensures that when we multiply it with the wave function vector, the resulting vector represents the difference quotient $\Delta p/\Delta x$.

However, this operator is not Hermitian, meaning that its conjugate transpose is not equal to the original matrix. To obtain a Hermitian operator, we can multiply the matrix by the imaginary unit i . This transforms the 1 's into i 's and -1 's into $-i$'s. Now, when we take the conjugate transpose, we get back the original matrix, ensuring that the operator is Hermitian.

Schrödinger's equation for a 1D free particle can be represented by the momentum operator $p = -i\hbar(d/dx)$. To discretize space and obtain a matrix representation of the momentum operator, we use the symmetric difference quotient $\Delta p/\Delta x$. By multiplying this matrix by the imaginary unit i , we obtain a Hermitian operator.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO IMPLEMENTING QUBITS****TOPIC: PARTICLE IN A BOX**

In the study of quantum information, one fundamental concept is the implementation of qubits, which are the basic units of quantum information. One way to understand this is by considering the particle in a box model, also known as the Radio model for the hydrogen atom.

In this model, we can represent the electron in a hydrogen atom as a particle confined to a one-dimensional space. This space is a segment of length L . To describe the state of this particle, we need to understand its energy eigenstates and how they evolve in time.

There are two ways to represent the particle's state within the segment. One approach is to introduce an infinite potential barrier beyond the segment, which prevents the particle from going beyond it. Another approach is to impose boundary conditions, stating that the wavefunction is zero at the ends of the segment.

By applying the Schrödinger equation, we can describe the evolution of the particle's state. In this case, the equation becomes $\hbar \frac{d}{dt} \psi = -\hbar^2 \frac{d^2}{dx^2} \psi$, with the boundary conditions $\psi(0) = 0$ and $\psi(L) = 0$.

To solve this equation, we need to find the eigenvectors and eigenvalues of the Hamiltonian operator, which represents the energy of the system. By examining its structure, we can guess that the eigenstate is of the form e^{ikx} , where k is a constant.

By substituting this form into the Schrödinger equation, we find that the corresponding eigenvalue is $\hbar^2 k^2 / 2m$. This tells us that the eigenstates come in pairs, with e^{ikx} and e^{-ikx} having the same energy. Any linear combination of these eigenstates also has the same energy.

To determine the specific eigenstates, we need to impose the boundary conditions. By rewriting the general solution in terms of sines and cosines, we can express it as $C \sin kx + D \cos kx$. The first boundary condition, $\psi(0) = 0$, leads to $D = 0$, simplifying the solution to $C \sin kx$.

The second boundary condition, $\psi(L) = 0$, allows us to determine the values of k . This condition implies that $C \sin kL = 0$, which means that kL is equal to an integer multiple of π . Therefore, the allowed values of k are given by $k = n\pi / L$, where n is an integer.

The eigenstates of the particle in a box model for the hydrogen atom are given by $\psi_n(x) = C \sin(n\pi x / L)$, where n is an integer. These eigenstates have energies $E_n = \hbar^2 (n\pi / L)^2 / 2m$.

In the study of quantum information, one fundamental concept is the implementation of qubits. A qubit is the basic unit of quantum information, analogous to a classical bit. In this didactic material, we will explore the concept of implementing qubits using the example of a particle in a box.

To begin, let's consider a one-dimensional line segment with length L . The particle in the box is confined within this segment and its motion is governed by the Schrödinger equation. The wave function, denoted as Ψ , describes the state of the particle.

To solve the Schrödinger equation for the particle in a box, we impose boundary conditions. At the ends of the segment, the wave function must be zero, indicating that the particle cannot exist outside the box. This leads to the quantization of the wave vector, denoted as K .

The quantization condition states that K must be of the form $n\pi/L$, where n is an integer. This means that K can only take discrete values. Consequently, the energy of the particle, denoted as E , is also quantized. The energy eigenvalues are given by $E_n = (\hbar^2 K_n^2) / (2m)$, where m is the mass of the particle.

The wave function, Ψ_n , corresponding to each energy eigenvalue, is determined by the normalization condition. The wave function must be normalized such that the integral of its magnitude squared over the entire segment is equal to 1. This normalization condition allows us to determine the constant of proportionality, denoted as C .

The normalized wave function, Ψ_n , can be expressed as $\Psi_n(x) = (\sqrt{2}/L)\sin(n\pi x/L)$, where x is the position along the segment. The energy eigenvalue, E_n , is given by $(\hbar^2 n^2 \pi^2)/(2mL^2)$.

As we vary the value of n , the energy eigenvalues increase. Higher energy states correspond to more oscillations in the wave function. This can be visualized by considering the shape of the wave function for different values of n . For example, when $n=1$, the wave function is given by $\Psi_1(x) = (\sqrt{2}/L)\sin(\pi x/L)$. As n increases, the number of oscillations in the wave function also increases.

Now, let's consider the time evolution of the wave function. If we start with an arbitrary wave function at time $t=0$, we can express it as a linear combination of the energy eigenfunctions. The coefficients of this linear combination, denoted as α_n , determine the time evolution of the wave function.

The time-dependent wave function, $\Psi(t)$, can be written as $\Psi(t) = \sum_n \alpha_n e^{(-iE_n t/\hbar)} \Psi_n(x)$. The phase factor $e^{(-iE_n t/\hbar)}$ causes the wave function to precess at a rate determined by the energy eigenvalue. Higher energy states precess at a faster rate.

The implementation of qubits in quantum information involves the use of a particle in a box as a simple model. The wave function of the particle is quantized, leading to discrete energy eigenvalues. The wave function is normalized, and the time evolution is determined by the coefficients of the energy eigenfunctions.

This model provides a good approximation for the behavior of electrons in a hydrogen atom. By comparing the energy difference between the ground state and the first excited state in the hydrogen atom to the energy difference in the particle in a box, we can estimate the size of the hydrogen atom. This calculation yields a value of approximately 3.4 angstroms.

The concept of implementing qubits using the example of a particle in a box provides valuable insights into the fundamentals of quantum information. The quantization of the wave function and the time evolution of the system are key aspects to consider in understanding the behavior of qubits.

The diameter of a hydrogen atom is approximately one angstrom, which is equivalent to four nanometers. This estimation provides us with a qualitative understanding of the hydrogen atom and how energy quantization occurs. By implementing a qubit using the electron in a hydrogen atom, we can further explore this concept. In the next material, we will delve into the interpretation of this implementation using the particle in a box model.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO IMPLEMENTING QUBITS****TOPIC: IMPLEMENTING QUBITS**

In order to implement a qubit, we can use the solution to the particle in a box problem. The eigenstates of the Hamiltonian for this system are denoted as Φ_n , and they correspond to different energy levels. By considering a box with a length of 3.4 angstroms, we can use the ground state ($n=1$) and the first excited state ($n=2$) as the basis states for our qubit, representing 0 and 1, respectively.

The state of the electron in this box can be written as α times the square root of $2/L$ times sine of $\pi x/L$ for $n=1$, plus β times the square root of $2/L$ times sine of $2\pi x/L$ for $n=2$. Here, α and β are coefficients that determine the probability amplitudes of the respective states.

The time evolution of the state is given by $\Phi(T)$, which can be expressed as α_0 times $e^{-iE_1 T/\hbar}$ plus β_1 times $e^{-iE_2 T/\hbar}$. By factoring out $e^{-iE_1 T/\hbar}$, we obtain α_0 plus β_1 times $e^{-i(E_2 - E_1)T/\hbar}$. Noting that $E_2 - E_1$ is equivalent to ΔE (the energy difference between the ground and excited states of the hydrogen atom), we can rewrite this as $e^{-iE_1 T/\hbar}$ times (α_0 times the square root of $2/L$ times sine of $\pi x/L$ plus β_1 times the square root of $2/L$ times sine of $2\pi x/L$), precessing at a rate of $e^{-i\Delta E T/\hbar}$.

The important observation here is that the relative phase between the two qubit values (0 and 1) is precessing at a rate proportional to $\Delta E/\hbar$, the energy difference for the hydrogen atom. This value is approximately 10 electron volts, corresponding to a frequency of about 2.5×10^{15} Hertz. Interestingly, this frequency is close to the frequency of optical light. Atomic qubits, such as the one described here, can be controlled through interaction with light pulses.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPLEXITY THEORY****TOPIC: LIMITS OF QUANTUM COMPUTERS**

In this lecture, we will discuss quantum complexity theory and the limits of quantum computation. The main focus will be on the issue of exponential speed ups for NP-complete problems. An NP-complete problem is a problem where we are searching for a solution in a large search space. For example, in the satisfiability problem, we are given a boolean formula and we want to know if there is a satisfying assignment for the variables. Classically, solving these problems takes exponential time. Quantumly, we can express these problems as finding a marked entry in a table, and there is an algorithm that solves this problem in $O(\sqrt{N})$ steps, where N is the size of the search space. However, this is still exponential time.

There have been claims in the news about the existence of quantum computers that can solve NP-complete problems in one shot. However, a theorem states that any quantum algorithm for solving the needle in a haystack problem must take at least some constant times $\sqrt{N}$ steps. This means that the number of queries or invocations of the quantum procedure is at least some constant times $\sqrt{N}$. It is important to note that this theorem does not say that understanding how the quantum procedure works would allow us to solve the problem faster. It only guarantees a lower bound when the procedure is used as a black box.

To understand why this lower bound holds, we can consider the related problem of distinguishing between the case where there is exactly one marked item and the case where there are no marked items. Solving this problem efficiently implies solving the search problem efficiently. By performing a test run on any quantum algorithm that claims to solve the search problem, we can show that it will not succeed if it takes fewer than $\sqrt{N}$ steps.

Quantum complexity theory explores the limits of quantum computation in terms of solving NP-complete problems. While quantum algorithms can offer speedups compared to classical algorithms, there is a lower bound on the number of steps required to solve these problems. Any quantum algorithm for the needle in a haystack problem must take at least some constant times $\sqrt{N}$ steps.

Quantum Complexity Theory is a field of study that focuses on understanding the limits of quantum computers in solving computational problems efficiently. One important concept in this field is the notion of quantum algorithms and their performance.

In a quantum algorithm, the algorithm performs queries on a superposition of states to solve a specific problem. The goal is to find the state or combination of states that minimizes a certain criterion. For example, in a search problem, the goal is to find the element in a database that satisfies a certain condition.

To analyze the performance of a quantum algorithm, we can look at the total squared amplitude with which a particular state is acquired during the algorithm's execution. This can be represented as the summation of the magnitudes squared of the amplitudes for each step of the algorithm. The state with the minimum total query magnitude squared is considered to be the state to which the algorithm pays the least attention.

To formalize this idea, we define a function $f(X)$ that equals 1 when X minimizes the summation of the magnitudes squared of the amplitudes for each step of the algorithm. The goal is to show that unless the number of steps T is large, the algorithm will still answer that the desired state is not found with high probability, even when the desired state is present.

To prove this, we use a technique called the hybrid argument. We start with a test run of the algorithm where the function f is identically 0. We know that if we were to measure the output of this run, we would observe 0 with high probability, indicating that the desired state is not found.

We then gradually change the function f to the desired function f^* . In each step, we answer the queries according to f up to the last query, and then answer the last query according to f^* . The difference between the output of the algorithm for f and f^* is bounded by the absolute value of the amplitude for the last query to the desired state.

By repeating this process, we create a series of hybrid states that gradually transition from f to f^* . The

difference between each hybrid state and the previous one is bounded by the amplitude of the last query to the desired state. This allows us to analyze the difference between the output of the algorithm for f and f^* .

Using this approach, we can show that the probability of the algorithm giving the correct answer for f^* is $\text{Big } O(T^2 / N)$, where T is the number of steps and N is the size of the problem. If T is much smaller than the square root of N , the probability of being correct is very small. This proves a lower bound that states any algorithm that is correct with high probability must take about the square root of N steps.

The hybrid argument is a powerful technique in quantum complexity theory that allows us to analyze the performance of quantum algorithms and understand their limitations. It provides insights into why certain computational problems cannot be solved efficiently by quantum computers.

In quantum complexity theory, one of the fundamental concepts is the limits of quantum computers. In order to understand these limits, we need to explore the notion of distance between state vectors and how it relates to the probability of distinguishing them.

Let's consider a scenario where we query a state vector X star with an amplitude of α X star of T minus one. If the output is Φ sub two, the distance between these two vectors is given by the absolute value of that amplitude. This distance represents how far apart the state vectors are at this particular time step.

To understand why this distance is also true at subsequent time steps, we need to consider the unitary nature of the computation. The rest of the computation remains the same in both cases, and since it is unitary, it preserves the angle between the two vectors. Therefore, the distance between them remains unchanged.

We can continue this argument by stepping from F to F star one step at a time until we reach the last step where we query on F star at each step. By doing so, we can determine how much the vector has changed. Starting with the initial vector Φ naught, we change it by α of T , α T minus $1/2$ of T minus 2 , and so on, until we reach half of 1 , which gives us the final vector.

The distance between the initial and final vectors can be calculated as the sum of the absolute values of these changes. Specifically, it is at most α of 1 plus α of 2 plus α of T in absolute value. By applying Cauchy-Schwarz inequality, we know that the sum of the squares of these changes is at most T . In the worst case scenario, where we want to maximize the sum of the absolute values, we make all of them equal to 1 over square root n . This results in a distance of T over square root of n .

In order to distinguish these two vectors, the distance between them must be a constant. This distance is directly related to the probability of successfully telling them apart. Therefore, in order to have a constant probability of success, T must be large compared to the square root of n . In other words, T must be approximately equal to the square root of n .

This argument, known as a hybrid argument, demonstrates the limits of quantum computers in terms of the required time steps to achieve a constant probability of success.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPLEXITY THEORY****TOPIC: ADIABATIC QUANTUM COMPUTATION**

Adiabatic quantum computation is an alternative approach to quantum computing that differs from the circuit model. In this method, the motivation is to solve problems such as unstructured search or NP-complete problems in polynomial time. The circuit model, where the function is given as a circuit and only queries and superposition are allowed, has a limitation that any quantum algorithm must take at least the square root of the input size in steps. However, this limitation does not necessarily mean that quantum computers cannot solve NP-complete problems in polynomial time.

Adiabatic quantum optimization, introduced in a paper published in Science, offers a new framework for solving NP-complete problems. The algorithm is complex, but simulations on small instances of problems like 3-SAT showed promising results of solving them in polynomial time. Adiabatic quantum optimization works by specifying an initial Hamiltonian, H_0 , and setting up the qubits in its ground state. Then, the Hamiltonian is gradually transformed to a final Hamiltonian, H_k , using a convex combination of H_0 and H_k . The transformation is done by gradually increasing a parameter, T , from 0 to 1. The goal is to find the ground state of H_k , which represents the solution to the problem at hand.

The quantum adiabatic theorem states that if the transformation is done slowly enough, the ground state is tracked throughout the process. This means that at any given time, the system remains in the ground state of the instantaneous Hamiltonian. The total time taken for the transformation must scale as one over the square of the spectral gap, which is the difference between the two smallest eigenvalues of the Hamiltonian. The smaller the gap, the longer the transformation time required.

To illustrate how adiabatic quantum computation can be applied to a specific problem, let's consider the satisfiability problem (SAT) with three variables per clause. In SAT, we have n bits, which can be transformed into n qubits. Each clause, such as $(x_1 \text{ OR } x_2 \text{ OR } x_3)$, corresponds to a Hamiltonian acting on the three qubits. The Hamiltonian assigns a penalty of 1 if the clause is not satisfied, meaning that the only assignment that fails to satisfy the clause is $(0, 0, 0)$. Otherwise, the penalty is 0. By constructing the appropriate Hamiltonian, the satisfiability problem can be encoded for adiabatic quantum optimization.

Adiabatic quantum computation offers a different approach to quantum computing, allowing for the potential solution of NP-complete problems in polynomial time. By gradually transforming an initial Hamiltonian to a final Hamiltonian, the algorithm tracks the ground state throughout the process. The total transformation time is determined by the spectral gap of the Hamiltonian. Adiabatic quantum optimization has shown promising results in solving problems like 3-SAT in polynomial time through simulations.

Adiabatic quantum computation is a method that uses the adiabatic theorem from quantum mechanics to solve computational problems. In this method, a quantum system is prepared in the ground state of a simple Hamiltonian, and then slowly evolved to the ground state of a more complicated Hamiltonian that encodes the problem to be solved. The hope is that the system will remain in the ground state throughout the evolution, allowing us to read off the solution at the end.

One of the problems that can be solved using adiabatic quantum computation is the Boolean satisfiability problem, which asks whether a given Boolean formula can be satisfied by assigning truth values to its variables. The idea is to encode the formula as a Hamiltonian, where each term corresponds to a clause in the formula. The ground state of the Hamiltonian represents a satisfying assignment to the formula.

To see how this works, let's consider a simple example. Suppose we have a formula with three variables and two clauses: $(x_1 \text{ OR } x_2) \text{ AND } (\text{NOT } x_2 \text{ OR } x_3)$. We can encode this formula as a Hamiltonian by assigning a qubit to each variable, and introducing terms that penalize assignments that do not satisfy the clauses. In this case, the Hamiltonian would be:

$$H = -h_1 * (I - Z_1) * (I - Z_2) - h_2 * (I - X_2) * (I - Z_3)$$

Here, I is the identity operator, Z_1 and Z_2 are Pauli-Z operators acting on qubits 1 and 2 respectively, X_2 is a Pauli-X operator acting on qubit 2, and Z_3 is a Pauli-Z operator acting on qubit 3. The parameters h_1 and h_2

control the strength of the penalties.

The ground state of this Hamiltonian represents a satisfying assignment to the formula. Any arbitrary truth assignment is an eigenvector with eigenvalue equal to the number of unsatisfied clauses. So, to solve the satisfiability problem, we can simply look at the ground state and check if it satisfies all the clauses.

However, it is important to consider the time complexity of adiabatic quantum computation. The total time that this algorithm must take is inversely proportional to the square of the minimum energy gap between the ground state and the first excited state of the Hamiltonian. If the gap is small, the algorithm will take a long time to run.

While adiabatic quantum computation can provide a quadratic speed-up for certain problems, it can also encounter difficulties. It has been shown that the minimum energy gap can become exponentially small, leading to exponential running time for the algorithm. Additionally, there are challenges related to decoherence, which refers to the fragility of quantum bits (qubits) and their susceptibility to measurement by the environment.

Despite these challenges, there are potential advantages to adiabatic quantum computation. It offers a way to protect qubits from decoherence by using a Hamiltonian that holds the qubits in the ground state. This makes it easier to implement quantum computation without the need for extensive fault tolerance measures. Furthermore, while the general algorithm may not solve NP-complete problems in polynomial time, there is hope that it could provide a speed-up for typical instances of these problems in practice.

Adiabatic quantum computation is a method that uses the adiabatic theorem to solve computational problems by evolving a quantum system from a simple Hamiltonian to a more complicated Hamiltonian. While there are challenges and limitations associated with this approach, it offers potential advantages in terms of decoherence and speed-up for certain problem instances.

Quantum computing has been a topic of great interest and excitement in recent years. However, it is important to approach the subject with caution and skepticism. Many headlines and articles tend to exaggerate the capabilities of quantum computers, leading to misconceptions and misunderstandings.

One such example is the claim of creating the first practical quantum computer. While it is true that researchers have made significant progress in demonstrating quantum effects in a small number of qubits, it does not necessarily mean that they have achieved a fully functional and scalable quantum computer.

Scott Aaronson, a professor at MIT and an expert in the field, highlights the gap between demonstrating quantum effects in a few qubits and building a quantum computer that can outperform classical computers in computationally interesting tasks. He emphasizes the need for a realistic assessment of the current state of quantum computing.

It is crucial to understand that the development of a practical quantum computer is a complex and challenging task. It requires overcoming numerous technical hurdles, such as qubit stability, error correction, and scalability. These challenges are still being actively researched and addressed by scientists and engineers in the field.

While the potential of quantum computing is undeniable, it is important to separate hype from reality. The current state of quantum computers is far from being able to solve complex problems faster than classical computers. However, ongoing research and advancements in the field continue to push the boundaries of what is possible.

It is essential to approach claims about practical quantum computers with skepticism and critical thinking. While progress has been made, there is still a long way to go before we can fully harness the power of quantum computing.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO QUANTUM COMPLEXITY THEORY****TOPIC: BQP**

Quantum complexity theory is a fascinating field that explores the complexity of quantum computations. In this material, we will focus on the complexity class BQP, which stands for Quantum Polynomial Time. BQP is the quantum analog of the complexity class P, or Polynomial Time, and the class BPP, or Probabilistic Polynomial Time.

In complexity theory, we study decision problems, where we have an input X and the output is either "yes" or "no". For example, the problem of primality testing asks whether a given number is prime or not. We can restate any problem as a yes-no problem and create a language L , which consists of inputs for which the answer is "yes". For instance, the language of primality is the set of all numbers that are prime.

Now, the question is whether we can solve these problems in polynomial time using a quantum computer. We say that a language L is in BQP if there exists a sequence of quantum circuits, one for each input size, such that on inputs of size n , the quantum circuit outputs a 1 with probability at least $2/3$ if X is in L , and outputs a 0 with probability at least $2/3$ if X is not in L . Additionally, the number of gates in the quantum circuit is bounded by a polynomial in n .

If we are not satisfied with the $2/3$ probability, we can run the algorithm multiple times and take the majority answer. This allows us to increase the probability of obtaining the correct answer as close to 1 as we desire. Therefore, we can achieve an error probability of 1 over 2 to the 100 by running the algorithm a few hundred times and taking the majority answer.

One of the main questions in quantum complexity theory is the power of BQP. Can BQP solve problems that cannot be solved in classical polynomial time? We have evidence that suggests this might be the case. For example, factoring is believed to be in BQP but not in BPP. We also have evidence from blackbox results, such as the quantum algorithms for sampling and Simon's problem, which can be solved in BQP but not in classical polynomial time. However, we currently do not have a proof that BQP is strictly larger than BPP.

Another important question is whether BQP contains problems that lie outside of NP, the class of problems that can be solved in nondeterministic polynomial time. This question is still open, and if we could show that BQP is strictly larger than P, we would have solved one of the major open questions in complexity theory. The relationship between P and PSPACE, the class of problems that can be solved using polynomial space, is also an open question.

BQP is a complexity class that represents the power of quantum polynomial time. While we have evidence that BQP is more powerful than classical probabilistic polynomial time, we still lack a formal proof. The exploration of BQP and its relationship with other complexity classes continues to be an active area of research in quantum information theory.

One of the fundamental questions in quantum information is whether the class BQP, which stands for Bounded-error Quantum Polynomial time, is contained within the polynomial hierarchy. The polynomial hierarchy is a hierarchy of problems, starting with P (Polynomial time) and NP (Nondeterministic Polynomial time), and extending to higher levels such as Σ_2^P , Π_2^P , and so on. The conjecture is that there exist problems in BQP that are not in the polynomial hierarchy, meaning they are harder than NP-complete problems.

This conjecture dates back to the early days of quantum computing. It suggests that certain sampling problems in BQP are not in the polynomial hierarchy. Proving this conjecture is challenging, but there is a related conjecture proposed by Scott Aaronson a few years ago. He introduced a simpler problem called Fourier checking and conjectured that Fourier checking is not in the polynomial hierarchy.

Fourier checking is a well-defined and easily stated problem. For more details, I recommend referring to Scott Aaronson's paper on this topic. It provides a comprehensive explanation of Fourier checking and its significance in the context of quantum complexity theory.

If you are interested in exploring this topic further, I encourage you to delve into the research and study the

details of Fourier checking and its relationship to the polynomial hierarchy.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO SPIN****TOPIC: SPIN AS A QUBIT**

In these last two lectures, we will discuss the fundamental principles behind the design of a quantum computer. The main questions we will address are: how to design physical qubits, how to initialize a qubit into the state 0, how to manipulate qubits using quantum gates, and how to measure qubits.

The state of a qubit is a vector in a two-dimensional complex vector space. Physical qubits refer to the eigenstates of the physical system, which are the states 0 and 1. To apply quantum gates, we need to apply a unitary transformation to the qubit, which can be achieved by applying a Hamiltonian. The unitary transformation is given by $e^{-iHt/\hbar}$ to the iHD over H bar, where H is the Hamiltonian and D is the duration of the gate.

There are various systems that can be used to implement qubits, such as atomic qubits, photons, spins, quantum dots, and superconducting loops. Many experimentalists around the world are actively working on implementing quantum computation using these systems. However, in these lectures, we will focus on the basic principles of manipulating, initializing, and measuring qubits.

Now, let's discuss spin. Elementary particles like electrons and protons have an intrinsic angular momentum called spin. When the particle is charged, like an electron, it also has an associated intrinsic magnetic moment. The spin of an electron can point either up or down, and it is quantized into these states, which can be thought of as the basis states 0 (spin up) and 1 (spin down).

The state of the spin system can be a superposition of spin up and spin down. The magnetic moment of an electron arises from the intrinsic angular momentum, which can be visualized as a spinning charge. Although the electron is not actually rotating like a sphere, its intrinsic angular momentum is quantized and expressed as a qubit in a two-dimensional complex vector space.

To understand how an external magnetic field affects the spin of an electron, we need to relate the two-dimensional complex vector representing the spin state to real space. This involves understanding how to locate the spin of the electron in three-dimensional space and how it interacts with the external magnetic field. These concepts will be further explained in the upcoming videos.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO SPIN****TOPIC: BLOCH SPHERE**

The Bloch sphere representation is a mathematical tool used to describe the state of a qubit in quantum information. It provides a way to map the state of a qubit, which exists in a two-dimensional complex vector space, to our three-dimensional real space. This representation is based on the concept that the state of a qubit can be described using two parameters, θ and ϕ .

θ represents an angle between 0 and π , while ϕ represents an angle between 0 and 2π . The state of any qubit can then be written as $\cos(\theta/2)|0\rangle + e^{i\phi}\sin(\theta/2)|1\rangle$, where $|0\rangle$ and $|1\rangle$ are the basis states of the qubit.

To understand why this representation is valid, we start by considering the state of a qubit in terms of complex numbers. A complex number can be written in Cartesian form ($a + bi$) or in polar form ($re^{i\phi}$), where r is the magnitude of the complex number and ϕ is the angle it makes with the real axis.

By expressing the complex numbers describing the qubit state in polar coordinates, we can rewrite them as $R_0e^{i\phi_0} + R_1e^{i\phi_1}$, where R_0 and R_1 are non-negative real numbers. However, the overall phase factor $e^{i\phi_0}$ does not affect any measurements made on the system and can be disregarded.

To describe the quantum state, we have the parameters ϕ ($\phi_1 - \phi_0$) and the positive real numbers R_0 and R_1 . Since the state must be normalized, $R_0^2 + R_1^2 = 1$. By defining $R_0 = \cos(\theta/2)$ and $R_1 = \sin(\theta/2)$, we obtain the previously mentioned representation of the qubit state.

The Bloch sphere representation allows us to visualize the qubit state in three-dimensional space. We can imagine the vector representing the state as a unit vector in three dimensions, with θ and ϕ as the polar coordinates. The z-axis represents the vertical direction, while the XY plane represents the x and y axes. By specifying the length of the vector as 1 and the angles θ and ϕ , we can determine a point on the surface of the unit sphere, representing the state of the qubit.

The Bloch sphere representation provides a way to describe the state of a qubit using two real parameters, θ and ϕ , which correspond to the polar angles in three-dimensional space. This representation allows us to visualize and manipulate qubit states effectively.

In the study of quantum information, it is important to understand the concept of spin and how it is represented on the Bloch sphere. The Bloch sphere is a geometrical representation that helps us visualize the states of a qubit.

To orient ourselves in space, we need to define the positive z-axis. This corresponds to the state of the qubit when θ is equal to 0 and ϕ is equal to 0. This state is commonly referred to as the zero state. On the other hand, the state of the qubit when θ is equal to π corresponds to the negative z-axis. This state is often referred to as the minus Z state.

One interesting observation is that in the complex vector space, the zero and one states are orthogonal. However, when we represent them on the Bloch sphere, they become antipodal states, pointing in opposite directions. This is where the factor of two, represented by θ over two, comes into play in our representation.

Moving on to the x-axis, the state represented by this axis can be determined by setting θ equal to $\pi/2$ and ϕ equal to zero. This results in the state being equal to $\frac{1}{\sqrt{2}}$ times the zero state plus $\frac{1}{\sqrt{2}}$ times the one state. This state is commonly referred to as the plus state.

To determine the state pointing in the minus x direction, we can follow a similar approach. By setting θ equal to $\pi/2$ and ϕ equal to π , we find that the state is equal to $\frac{1}{\sqrt{2}}$ times the zero state minus $\frac{1}{\sqrt{2}}$ times the one state.

Moving on to the y-direction, the state pointing in this direction can be determined by setting θ equal to $\pi/2$ and ϕ equal to $\pi/2$. This results in the state being equal to $\frac{1}{\sqrt{2}}$ times the

zero state plus i times 1 over square root of 2 times the one state. Here, i represents the imaginary unit.

Similarly, the state pointing in the minus y direction can be determined by setting θ equal to π over two and ϕ equal to three times π over two. This results in the state being equal to 1 over square root of 2 times the zero state minus i times 1 over square root of 2 times the one state.

To better understand the states and their directions on the Bloch sphere, it is recommended to visually explore and manipulate the sphere. This will help solidify the understanding of which states point in which direction.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO SPIN****TOPIC: STERN-GERLACH EXPERIMENT**

In the field of quantum information, understanding the concept of spin is crucial. Spin is a fundamental property of particles, such as electrons, and it plays a significant role in various quantum phenomena. In this didactic material, we will explore the measurement of spin in the lab and the physical meaning behind the block sphere representation.

To measure the spin of an electron in the lab, we utilize the fact that the spin creates a magnetic moment. This magnetic moment allows us to manipulate and measure the spin. The approach involves using an external magnetic field, specifically a non-homogeneous field. The experimental setup consists of a unique-looking tip and a bar magnet. The resulting magnetic field is strong at one end and gradually weakens as we move along its length.

Now, imagine an electron passing through this apparatus. Depending on the spin state of the electron, its path will either bend upwards or downwards. If the electron's spin is pointing up, it takes the upper path. Conversely, if the spin is pointing down, it takes the lower path. This distinction between the two paths allows us to separate different spin states. For example, if the electron's state is described by the superposition $\alpha|0\rangle + \beta|1\rangle$, the apparatus neatly separates the two cases. By observing the electron's position, we can determine the probabilities α^2 and β^2 associated with each path, effectively measuring the spin state.

The experiment we just described, known as the Stern-Gerlach experiment, was first conducted in 1922 using silver atoms instead of electrons. The key to its success lies in the interaction between the spin and the external magnetic field. The spin, acting as a tiny magnet, aligns itself opposite to the external magnetic field to minimize its energy. This alignment creates a low energy state for the spin down configuration and a high energy state for the spin up configuration.

When the electron passes through the non-homogeneous magnetic field, the spin down state experiences a force pushing it downwards due to the field's increasing strength in that direction. On the other hand, the spin up state feels a force pushing it upwards since it has a positive energy and seeks to minimize it. This force causes the trajectory of the electron to bend accordingly. This semi-classical explanation provides insight into the workings of the Stern-Gerlach device.

Now, let's consider a different scenario. If we point the Stern-Gerlach device in a different direction, what would happen to the electron's trajectory? To answer this question, we need to understand the block sphere. The direction in which the device is pointing, denoted as u , plays a crucial role in determining the electron's behavior.

The measurement of spin in the lab involves utilizing the magnetic moment created by the spin of particles. The Stern-Gerlach experiment demonstrates how an external magnetic field can be used to manipulate and measure spin. By analyzing the trajectory of the particles passing through the non-homogeneous magnetic field, we can determine their spin states. This experiment, conducted in 1922 by Stern and Gerlach, provides valuable insights into the behavior of spin in quantum systems.

In quantum information, the concept of spin plays a crucial role. The Stern-Gerlach experiment is a fundamental experiment that helps us understand the behavior of spin in quantum systems.

The Bloch sphere is a useful tool in visualizing the direction of spin in three-dimensional real space. The direction of spin is specified by two angles, θ and ϕ , and corresponds to a quantum state, $|\psi\rangle$. This state can be represented as $\cos(\theta/2)|0\rangle + e^{i\phi}\sin(\theta/2)|1\rangle$.

When an electron is passed through a Stern-Gerlach device oriented in the direction u , it takes one of two paths. The upper path is taken if the electron is in the state $|\psi\rangle$. On the other hand, if we consider the antipodal point, $-u$, the electron takes the lower path, corresponding to the state $|\psi\rangle$. The state $|\psi\rangle$ can be calculated as $\sin(\theta/2)|0\rangle - e^{i\phi}\cos(\theta/2)|1\rangle$.

These two states, $|\psi\rangle$ and $|\psi\rangle$, are orthogonal states of the spin qubit. If the initial state of the

electron is $\alpha \psi_u + \beta \psi_{-u}$, then the probability of observing the electron in the upper path is $|\alpha|^2$, and the probability of observing it in the lower path is $|\beta|^2$.

Now, let's consider two Stern-Gerlach devices, one pointing in the direction u and the other pointing in the direction v . If we pass the electron through the first device and then through the second device, we want to know the probability of observing the electron bending upwards in both devices. Experimentally, it has been found that this probability is given by $\frac{1 + \cos(\mu)}{2}$, where μ is the angle between the directions u and v . In trigonometric terms, this probability is equal to $\cos^2(\mu/2)$.

From the perspective of the quantum state, if the electron bent upwards in the first device, it means that the spin was in the state ψ_u at that point. Now, we want to calculate the probability of the spin transitioning from ψ_u to ψ_v when measured in the basis ψ_v and ψ_{-v} . This probability is given by $\cos^2(\nu)$, where ν is related to μ by $\nu = \mu/2$.

Therefore, the experimental results and the quantum state analysis are consistent. The Bloch sphere provides a physical significance to the direction of spin. Even though spin lives in a two-dimensional complex vector space, we can think of it as living on the Bloch sphere. The direction in which the spin points on the Bloch sphere corresponds to the orientation of the magnet and determines how the spin interacts with external fields.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: INTRODUCTION TO SPIN****TOPIC: PAULI SPIN MATRICES**

In the previous material, we discussed the measurement of spin in an electron. Now, let's delve into the observable that corresponds to spin measurement. The observable depends on the direction in which we measure the spin or, in the case of the Stern-Gerlach device, how we orient it. Let's consider different scenarios.

First, let's consider the simplest case where the Stern-Gerlach device is oriented along the z-axis. In this case, we define spin up as 0 and spin down as 1. To obtain an observable that represents the spin measurement, we want eigenvalues of 1 and -1 corresponding to the spin up and spin down states, respectively. The matrix that satisfies this condition is called the Pauli spin matrix Sigma sub Z, which is represented as:

$$\begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}$$

Now, let's consider a scenario where the Stern-Gerlach device is oriented along the x-axis. The eigenvalues for the spin up and spin down states in this case are given by:

$$\begin{aligned} \text{Spin up: } & \frac{1}{\sqrt{2}} * (0 + 1) \\ \text{Spin down: } & \frac{1}{\sqrt{2}} * (0 - 1) \end{aligned}$$

To obtain an observable with eigenvalues of +1 and -1 for these states, we use the Pauli spin matrix Sigma sub X, which is represented as:

$$\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

Similarly, if the Stern-Gerlach device is oriented along the y-axis, the eigenvalues for the spin up and spin down states are given by:

$$\begin{aligned} \text{Spin up: } & \frac{1}{\sqrt{2}} * (0 + i) \\ \text{Spin down: } & \frac{1}{\sqrt{2}} * (0 - i) \end{aligned}$$

In this case, the observable is represented by the Pauli spin matrix Sigma sub Y, which is:

$$\begin{bmatrix} 0 & i \\ -i & 0 \end{bmatrix}$$

These three matrices, Sigma sub X, Sigma sub Y, and Sigma sub Z, are collectively known as the Pauli spin matrices. They are crucial in determining the spin components in the x, y, and z directions.

It is important to note that these matrices do not commute with each other. This means that if you measure the spin in one direction and then measure it in another direction, the results will not be the same as if you had measured them in the reverse order. This non-commutativity is a fundamental aspect of quantum mechanics and is related to the uncertainty principle.

Understanding the Pauli spin matrices and their non-commutativity is essential in grasping the behavior of spin in quantum systems.

Quantum Information Fundamentals - Introduction to Spin - Pauli Spin Matrices

In the realm of quantum information, the concept of spin plays a crucial role. Spin refers to an intrinsic property of particles, such as electrons, that gives rise to their magnetic moment. Understanding spin is essential for various applications in quantum computing and quantum communication.

One intriguing aspect of spin is that it can exist in multiple directions simultaneously. This property is known as superposition. When we measure the spin of a particle in one direction, it disturbs the spin in other directions,

and vice versa. This phenomenon is a consequence of the fundamental principles of quantum mechanics.

To describe spin mathematically, we use a set of matrices known as Pauli spin matrices. These matrices, named after physicist Wolfgang Pauli, provide a mathematical representation of the spin states of particles. There are three Pauli spin matrices: σ_x , σ_y , and σ_z , each corresponding to a different direction of measurement.

The Pauli spin matrices have distinct properties that make them particularly useful in quantum information. For example, they are Hermitian, meaning they are equal to their own conjugate transpose. This property ensures that the eigenvalues of these matrices are real, allowing us to extract meaningful information from spin measurements.

In addition to their Hermitian nature, the Pauli spin matrices also satisfy the Pauli algebra, a set of commutation and anticommutation relations. These relations are fundamental in quantum mechanics and play a crucial role in the manipulation and analysis of quantum systems.

By leveraging the properties of the Pauli spin matrices, researchers can perform various operations on spin states, such as rotations and transformations. These operations are vital for encoding and processing quantum information in quantum algorithms and protocols.

The concept of spin is a fundamental aspect of quantum information. It allows particles to exist in superposition and provides a basis for representing spin states using Pauli spin matrices. Understanding these matrices and their properties is essential for the development of quantum technologies.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: MANIPULATING SPIN****TOPIC: LARMOR PRECESSION**

In this lecture, we will discuss how to manipulate spin and implement quantum gates on a spin qubit. To understand this, we need to locate the spin qubit on a Bloch sphere. A quantum gate or unitary transformation on a qubit state is performed by rotating the Bloch sphere about some axis. For example, the bit flip gate (X gate) is a rotation of the two-dimensional vector space about a 45-degree axis. On the Bloch sphere, the 0 state sits on the plus z-axis, the 1 state is antipodal to it, and the plus state is on the plus x-axis. To perform the X gate, we rotate the Bloch sphere about the x-axis by 180 degrees (a π rotation), resulting in a flip where 0 goes to 1 and 1 goes to 0.

To manipulate the spin qubit, we use an external magnetic field that acts on the spin as a little magnet. By grabbing hold of this magnet with the magnetic field, we can rotate the Bloch sphere about some axis. The interaction of the spin with the external magnetic field is described by the Hamiltonian $\hat{H} = \frac{\hbar \gamma}{2} \hat{\sigma}_z B$, where $\hat{\sigma}_z$ is the Pauli spin matrix representing a phase flip. The spin observable in the z direction is $\hbar \gamma \hat{\sigma}_z$. By solving the Schrödinger equation, we can determine the state of the spin qubit as a function of time under this Hamiltonian.

The solution to the Schrödinger equation shows that the spin qubit stays at the same angle θ with the z direction but starts precessing. The angle ϕ changes as a function of time, and the spin precesses at a certain rate called the Larmor frequency (Ω_L). This evolution of the spin state is known as Larmor precession. To perform a gate that rotates about the z-axis, we would turn on the magnetic field pointing in the z direction for a specific time to rotate the Bloch sphere by the desired angle.

To manipulate spin and implement quantum gates on a spin qubit, we rotate the Bloch sphere representing the qubit state about some axis. This rotation is achieved by using an external magnetic field that interacts with the spin qubit. By understanding the Hamiltonian and solving the Schrödinger equation, we can determine the evolution of the spin state over time. Larmor precession describes the precessing motion of the spin qubit under the influence of the magnetic field.

When manipulating spin in quantum information, one important concept to understand is Larmor precession. Larmor precession refers to the rotation of a spin about a magnetic field. This rotation can be described using the Bloch sphere, which represents the possible states of a qubit.

To carry out any arbitrary single qubit gate on the qubit state, we need to understand where the Hamiltonian for Larmor precession comes from and how to solve it. The Hamiltonian is derived from classical intuition about the magnetic moment of a spinning charge.

The energy of a magnet with a magnetic moment μ in an external magnetic field B is given by $-\mu \cdot B$, meaning it wants to align with the external field. The magnetic moment usually comes from a spinning charge. If a charge moves around in a circle of radius R with velocity V , the total charge that passes by a particular point in one rotation is given by e , the charge of an electron, times the amount of charge that goes around in one rotation.

The current is then the charge over time, which is $eV/2\pi$. The magnetic moment is given by the current times the area, which is $eV/2\pi R$ times πR^2 . Simplifying this expression gives us the magnetic moment as $-eL/2m$, where L is the angular momentum.

Comparing this expression to the Hamiltonian for Larmor precession, we find that the angular momentum is related to the Hamiltonian by a factor of $\hbar/2$. This factor of 2 arises due to quantum mechanics and is known as the g factor. For electrons, the g factor is 2.

Once we have the Hamiltonian, we need to solve it to understand how the qubit state evolves. The Hamiltonian is given by $\frac{e\hbar}{2m} \hat{\sigma}_z B$, where $\hat{\sigma}_z$ is the Pauli matrix in the Z direction. The eigenvectors and eigenvalues of the Hamiltonian are simply 0 and 1, with eigenvalues of $\frac{e\hbar B}{2m}$ and $-\frac{e\hbar B}{2m}$.

Using these eigenvalues, we can compute the time evolution of the qubit state. The state at time T is given by

$\alpha = 0 + e^{i\Omega L T} \beta$, where ΩL is equal to eB_0/m . This equation describes the rotation of the qubit state about the Z axis.

Larmor precession is a fundamental concept in quantum information that involves the rotation of a spin about a magnetic field. The Hamiltonian for Larmor precession is derived from classical intuition about the magnetic moment of a spinning charge. Solving the Hamiltonian allows us to understand how the qubit state evolves over time.

When manipulating spin in quantum information, it is important to understand how the spin state evolves over time. One way to visualize this evolution is by using the Bloch sphere.

Let's consider an initial spin state represented by the vector $(\cos(\theta/2), 0, \sin(\theta/2))$ on the Bloch sphere. This state can be written as a linear combination of the basis states $|0\rangle$ and $|1\rangle$, where $|0\rangle$ represents spin up and $|1\rangle$ represents spin down. The coefficients in the linear combination are determined by the angles θ and ϕ .

As time progresses, the spin state evolves. At time T , the state can be written as $(\cos(\theta/2), 0, \sin(\theta/2)) + e^{i(\phi + \omega L T)} (0, \cos(\beta), \sin(\beta))$, where ωL is the Larmor frequency and β is a constant.

The effect of this evolution is that the phase ϕ changes to $\phi + \omega L T$. This means that the phase of the spin state processes at a rate determined by the Larmor frequency.

When manipulating spin in quantum information, the spin state evolves over time. The evolution can be visualized on the Bloch sphere, where the phase of the state processes at a rate determined by the Larmor frequency.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: MANIPULATING SPIN****TOPIC: SPIN RESONANCE**

In the previous material, we learned about using Larmor precession to implement single qubit gates on a spin. However, this method is not practical due to the large D field required and the difficulty of rapidly changing its direction. An alternative approach is to use spin resonance, which provides finer control for implementing quantum gates.

To understand spin resonance, let's consider it as a two-step process. In the first step, we turn on a large DC field, B_0 , pointing in the Z direction. This splits the energy levels of spin-up and spin-down states, corresponding to 0 and 1, respectively. The energy splitting between these states is $\hbar \Omega_{\text{L}}$, where Ω_{L} is the Larmor frequency, given by $\Omega_{\text{L}} = e B_0 / m$.

In the second step, we turn on a small AC field, $B_1 \cos(\Omega_{\text{L}} T)$, which oscillates at the Larmor frequency. This field points in the X direction, causing the field to oscillate back and forth. The effect of this AC field is to induce spin flips or a controlled mixing between the 0 and 1 states.

Quantitatively, the Larmor frequency is $\Omega_{\text{L}} = e B_0 / m$, and the Rabi frequency is $\Omega_1 = e B_1 / (2m)$. If the initial state of the qubit is given by $\cos(\theta/2) |0\rangle + e^{i\phi} \sin(\theta/2) |1\rangle$, the evolution of the qubit under the DC field and the small AC field is described by the equation:

$$\Phi(T) = \cos(\theta + \Omega_1 T / 2) + e^{i\phi} \sin(\theta + \Omega_1 T / 2)$$

This equation shows that the qubit's state evolves as it spirals down on the Bloch sphere. The angle θ changes due to the Rabi oscillations at the frequency Ω_1 , while the change in Φ is determined by the Larmor frequency Ω_{L} . Typically, Ω_{L} is much larger than Ω_1 , with Ω_{L} in the gigahertz range and Ω_1 in the kilohertz range.

Spin resonance provides a more practical way of implementing single qubit quantum gates compared to Larmor precession. By using a combination of a large DC field and a small oscillating AC field, we can induce controlled mixing between the spin-up and spin-down states, allowing for precise manipulation of quantum information.

To manipulate spin in quantum information, we can use spin resonance techniques. One way to achieve spin manipulation is by applying a small alternating current (AC) field on top of a constant direct current (DC) field. By turning on the AC field for a specific duration, we can flip the spin.

To perform a spin flip, we need to choose a pulse duration that satisfies the condition $\Omega_1 \Delta T = \pi$, where Ω_1 is the frequency of the AC field and ΔT is the pulse duration. This condition ensures that the spin is flipped to the opposite direction.

To understand the effect of the AC field on the spin, we can think about it intuitively. When we turn on the DC field, the spin starts precessing at the Larmor frequency. However, if we watch the spin in a rotating frame that rotates at the Larmor frequency, the spin appears stationary to us. In this rotating frame, the effect of the DC field ceases to exist.

Now, let's consider the effect of the AC field in the rotating frame. The AC field, represented by $B_1 \cos(\Omega_{\text{L}} T)$, oscillates between -1 and 1. We can achieve this effect by considering two counter-rotating fields, each with a magnitude of $B_1/2$. One of these components rotates along with us in the rotating frame, while the other component rotates relative to us at twice the frequency. The component rotating along with us appears stationary in the rotating frame, while the other component cancels out due to its rotation.

In the rotating frame, the AC field appears as an effective magnetic field pointing in the X direction, with a magnitude of $B_1/2$. This effective magnetic field causes the spin to precess at a frequency of Ω_1 , which is given by $\Omega_1 = (e * B_1/2) / m$, where e is the charge of the particle and m is its mass.

The net effect of both the DC and AC fields is a combination of precessions around the X and Z axes. This results in Rabi oscillations, which occur at a frequency of $\Omega_1/2$. The factor of two arises because only half of the

oscillating field rotates along with us in the rotating frame.

By applying an AC field on top of a DC field and watching the spin in a rotating frame, we can manipulate the spin and induce Rabi oscillations. This technique is essential for controlling and manipulating spin in quantum information processing.

In quantum information, manipulating spin is a crucial aspect. To understand how two-qubit gates are implemented, we need to create an interaction between two spins. This is achieved by placing two electrons next to each other, allowing them to feel each other's presence. Each electron has a spin pointing in a certain direction, and because of the associated magnetic moment of the spin, they interact with each other.

The ground state of these two spins, when they are close to each other and interacting, is known as the Bell State or the singlet state. This state is represented by the equation $\frac{1}{\sqrt{2}}(|01\rangle - |10\rangle)$, where $|0\rangle$ represents the spin-up state and $|1\rangle$ represents the spin-down state. The singlet state has a total spin of zero and is rotationally invariant.

To implement two-qubit gates, such as those used in quantum computing, the creation of entanglement between spins is necessary. By utilizing the interaction between the spins of two electrons, we can achieve this entanglement. Single-qubit gates alone are not sufficient for quantum computing, but when combined with the ability to implement any two-qubit gate, we obtain a universal family of gates.

By placing two electrons next to each other and allowing them to interact, we can create an entangled state known as the Bell State or singlet state. This state is crucial for implementing two-qubit gates and achieving the universal family of gates necessary for quantum computing.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: MANIPULATING SPIN****TOPIC: CLASSICAL CONTROL**

A quantum computer is a complex system that relies on the principles of quantum mechanics to perform computations. In this didactic material, we will explore the fundamentals of quantum information, specifically focusing on manipulating spin and classical control.

The model of a quantum computer consists of qubits, which are the basic units of quantum information. These qubits are controlled by a classical computer, which issues commands to manipulate them. The classical computer interacts with the qubits through external means, such as lasers or other equipment. This external control allows for flexibility in implementing operations on the qubits, as the sequence of gates can be determined by the programmer.

However, there is a challenge in achieving this control while also isolating the qubits. When the qubits interact with the external world, there is a risk of decoherence, which refers to the inadvertent measurement of the quantum system by the environment. Decoherence is a significant obstacle in implementing a quantum computer, as it disrupts the delicate quantum states of the qubits.

The contradiction arises from the fact that we want to control the qubits from the outside, but at the same time, we need to prevent them from being measured by the environment. The question then becomes, how can we interact with the qubits without measuring them?

To illustrate this concept, let's consider a quantum computer with a qubit in the state $\alpha|0\rangle + \beta|1\rangle$. Assume that the qubit interacts with the environment, represented by another qubit, and subsequently gets lost. Later, when we measure the output of the quantum computer, the environment qubit is also measured. According to the principle of deferred measurement, if the environment qubit has not interacted with the rest of the system for a long time, the measurement can be moved back in time without affecting the outcome.

This principle implies that any interaction with the environment can be regarded as a measurement. Therefore, it becomes challenging to simultaneously achieve the goals of external control and qubit isolation. However, the design of quantum computers relies on a classical computer that controls the qubits, based on this understanding.

Now, let's delve into the principles behind manipulating spin qubits. Consider a spin qubit in the state $\alpha|\text{up}\rangle + \beta|\text{down}\rangle$. To perform a bit flip operation, we want the qubit to be in the state $\alpha|\text{down}\rangle + \beta|\text{up}\rangle$. Due to the energy difference between the spin up and spin down states, a bit flip operation can be achieved.

It is important to note that this didactic material provides an overview of the topic of quantum information, specifically focusing on manipulating spin and classical control. Further exploration and understanding of this topic require in-depth study and application of mathematical formalism and quantum mechanics principles.

In the context of manipulating spin in quantum information, it is important to understand the concept of classical control. Classical control refers to the ability to control and manipulate quantum systems using classical methods and techniques. In this didactic material, we will explore how classical control can be achieved in the context of manipulating spin.

To begin, let's consider the process of flipping the spin of a system from the spin-up state to the spin-down state. In order for this flip to occur, the system must emit a photon. Similarly, to go from the spin-down state to the spin-up state, there must be an absorption of a photon. At first glance, it may seem that this process involves a measurement of the spin state. However, this is not the case.

To understand why this is not a measurement, let's consider the use of a linearly polarized field to carry out the spin flip. Let's assume that this field contains K photons. In this scenario, there are two possibilities: either a photon is emitted or a photon is absorbed. If a photon is emitted, the system ends up with $K+1$ photons. If a photon is absorbed, the system ends up with $K-1$ photons. In principle, we can differentiate between these two cases by observing the state of our environment. However, this does not constitute a measurement of the spin state.

The key insight here is that the state of the environment provides information about the number of photons, not the actual spin state. This means that while we are trying to control and manipulate the qubit to perform a spin flip, we are not actually measuring the qubit itself. This is important because in quantum computing, measurements can disrupt the quantum state and introduce errors.

So, how can we achieve classical control without measuring the qubit? The answer lies in the nature of the field we are using. If we consider using a laser, for example, the state of the field can be described as a superposition of different photon numbers. This superposition can be represented as a Gaussian distribution, where each term corresponds to a different number of photons.

When we apply this field to the qubit, the result depends on whether a photon is emitted or absorbed. In the case of emission, the Gaussian distribution shifts to the left by one position. In the case of absorption, the Gaussian distribution shifts to the right by one position. However, due to the large width of the Gaussian distribution, these shifts are not noticeable.

In other words, the probability of distinguishing between the two scenarios (emission or absorption) is inversely proportional to the width of the Gaussian distribution. Since the width is typically very large, the probability of distinguishing between the two scenarios is very low. This implies that there is almost no trace in the environment indicating whether a single photon was emitted or absorbed by the qubit.

This property allows us to carry out classical control of the quantum system without introducing measurement errors. By carefully manipulating the field and taking advantage of the large width of the Gaussian distribution, we can effectively control the qubit without disrupting its quantum state. This is a crucial aspect of implementing quantum computers, where precise control is essential for performing quantum operations.

Classical control plays a vital role in manipulating spin in quantum information. By using classical methods and techniques, we can control and manipulate quantum systems without directly measuring their states. This allows for precise control of qubits, enabling the implementation of quantum computers and other quantum technologies.

EITC/QI/QIF QUANTUM INFORMATION FUNDAMENTALS DIDACTIC MATERIALS**LESSON: SUMMARY****TOPIC: SUMMARY**

Quantum Information Fundamentals - Summary

In this course on quantum information, we covered various important topics, but there are still some areas that we didn't have the opportunity to explore. One of these topics is quantum error correcting codes, which were a significant discovery in the field of quantum computing. Initially, it was believed that quantum systems couldn't be protected from environmental decoherence. However, quantum error correcting codes allow us to transform a quantum state into a longer state on more qubits, enabling the correction of errors caused by environmental noise. By applying an error correcting circuit to fresh qubits, the errors get wiped out, and the information about these errors is stored in the clean qubits.

Another crucial aspect we didn't cover extensively is fault-tolerant quantum computation. This theory focuses on performing computations on qubits while decoherence is occurring, which is essential for implementing quantum computers. Additionally, we didn't delve into more advanced algorithms such as phase estimation and quantum walk-based algorithms. Quantum cryptography, which utilizes the properties of quantum mechanics to implement secure cryptographic systems, also remained unexplored.

Furthermore, we didn't have enough time to discuss the current state of experimental realization in quantum information. The field is rapidly advancing, and it is important to stay updated on the latest developments. Lastly, there is a growing interaction between quantum computation and various branches of physics, particularly in the area known as quantum Hamiltonian complexity. This field explores the connection between quantum complexity theory and condensed matter physics.

Although we couldn't cover all these topics in detail, it is important to mention their relevance to the concepts we did cover. For example, teleportation, which we discussed earlier in the course, can be utilized to carry out quantum gates and quantum computation. This idea plays a significant role in efficient fault-tolerant quantum computing. Additionally, the CH SH game, which we used to test quantum mechanics, has applications in building random number generators that can be certified as truly random.

This course provided a solid foundation in quantum information fundamentals. However, there are still many exciting areas to explore within the field, such as quantum error correcting codes, fault-tolerant quantum computation, advanced algorithms, quantum cryptography, experimental realization, and the interaction between quantum computation and different branches of physics.

In order to test the validity and functionality of a claimed quantum computer, it is important to have a method of verification. This is especially relevant when someone claims to have developed a quantum computer that can solve problems that are extremely complex to solve using classical methods. One way to carry out this verification is through the use of the CH SH game.

The CH SH game is closely related to the testing of a quantum computer's output. By understanding the principles behind the CH SH game, one can gain insights into how to effectively test the behavior of a quantum computer. For those who are interested in delving deeper into this topic, further reading and research is recommended.

In a previous survey, feedback was collected to help improve the course on quantum information. The feedback received was highly valuable and provided several key takeaways. One important point was the need to assist students in dealing with the mathematical background required for the course. To address this, it was suggested that background material on basic linear algebra be made available a few weeks before the start of the course. This would allow students to assess their own readiness and brush up on the necessary mathematical concepts.

Another aspect that emerged from the feedback was the potential for utilizing the discussion forum in a more meaningful way. The discussion forum plays a vital role in the course, and there are opportunities to enhance its usage. It was noted that the EDX platform, on which the course is hosted, is undergoing significant development. This presents an exciting prospect for improving the discussion forum and exploring new and

interesting ways to utilize it.

Expanding the range of topics covered in the course was also a suggestion put forth by many students. While some examples were provided on a previous slide, it was acknowledged that there is room for further expansion. Additionally, students expressed an interest in being provided with pointers to research questions. This suggestion is highly valuable and will be considered for future iterations of the course.

Based on the feedback received and ongoing course development, a new survey will be distributed. This survey will build upon the existing questions and include additional inquiries. The input from students is invaluable in shaping the direction of the course and will guide future improvements.

It is important to acknowledge the significant effort and dedication put into creating and maintaining this course. The main contributor, Spangle, is a graduate student who has worked tirelessly in collaboration with the professor to design the course and develop the content. Spangle's commitment stems from the belief that making this course accessible worldwide is of utmost importance.

Furthermore, the contributions of Gasps and Simon Stevenson, who invested considerable effort in managing the discussion forums, should be recognized. The discussion forums are considered the lifeblood of the course, and their active involvement greatly enhances the learning experience for all participants.

As we conclude this journey, it is important to note that feedback and thoughts from students are highly valued. The discussion forum provides a platform to express opinions and share experiences related to the course. We sincerely hope that you have enjoyed the course and invite you to continue engaging in the discussion forum to provide further insights and reflections.