

Before starting QED I want to use the ideas we developed on spacetime diagrams to give some hints about what is to come in the next few weeks as we discussed elementary particles.

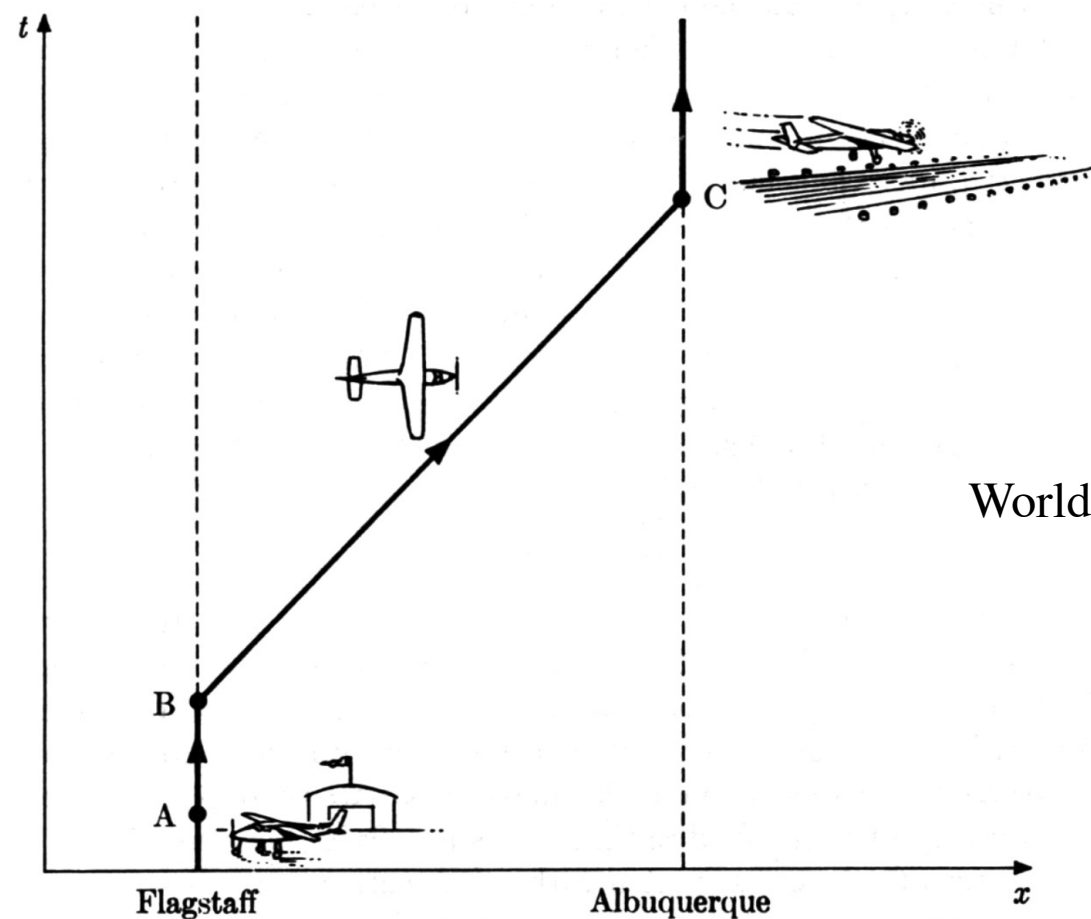
Some thoughts based on Spacetime Diagrams to be expanded in later discussions

In picturing the interactions of particles it is **important** to think of the when as well as the where, hence to think of paths through time as well as paths through space.

Such spacetime paths are called **worldlines** as we have discussed.

A full spacetime map would require four dimensions, but often two suffice, one for a particular direction in space, one for time as we have discussed.

As a **reminder** of the appearance of world-line diagrams, the figure below shows a simplified spacetime map for an airplane trip from Flagstaff to Albuquerque (due east).



World-line diagram: a trajectory in a spacetime map.

The axes are labeled x (distance to the east) and t (time).

Recall that there is no such thing as standing still in a spacetime map.

Something that remains at rest - the city of Flagstaff - **still** moves through time.

The world line of Flagstaff is shown by a vertical dashed line in the figure, and another **vertical** dashed line farther east is the world line of Albuquerque.

Now what about the airplane?

While parked on the field in Flagstaff, it also moves only through time, not through space, and also traces out a vertical world line (segment AB in the diagram).

Then the plane takes off and flies east, moving both in space and time, and tracing out the world line BC.

Having landed at Albuquerque, it once again moves only in time, and its world line continues upward.

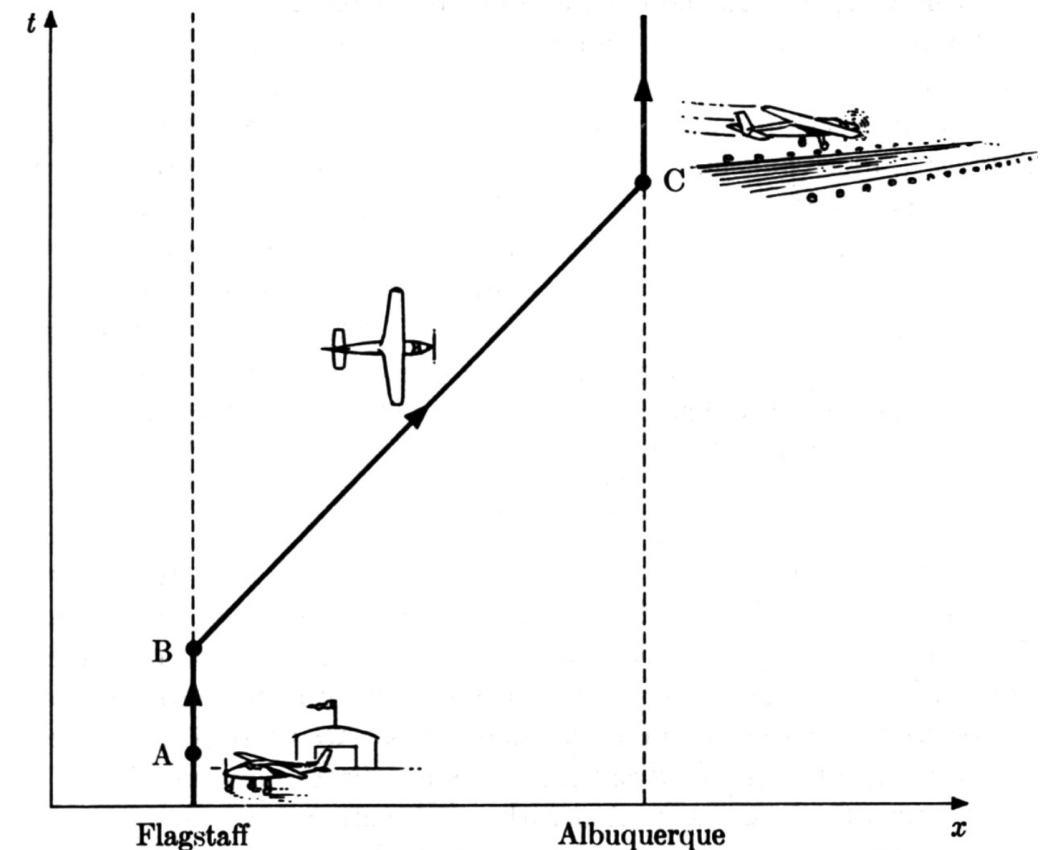
I have attached arrows to the world line of the airplane.

They may seem redundant, since there is, after all, only one possible direction to move through time - forward.

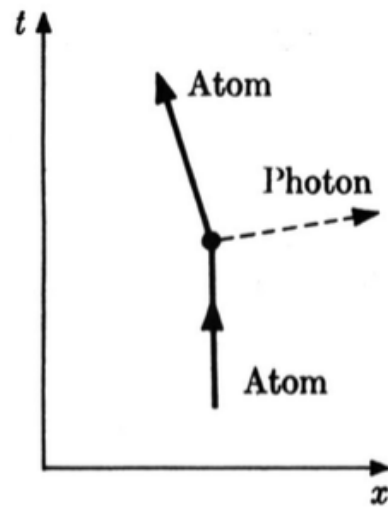
But the arrows do no harm and will serve a useful **purpose** later in the fundamental-particle world.

This simple spacetime map describes motion only along a single straight line, but that will be sufficient for understanding particle interactions.

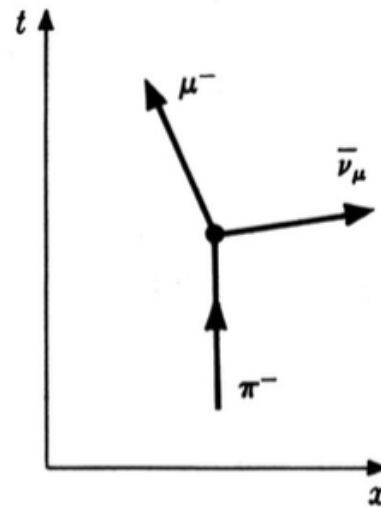
We may imagine more complicated maps with two or even three space dimensions and one time dimension; but it is better to restrict the diagrams to the form illustrated here.



Turning now to the small-scale world, I present world lines of some simple occurrences in the figures below.



(a)



(b)

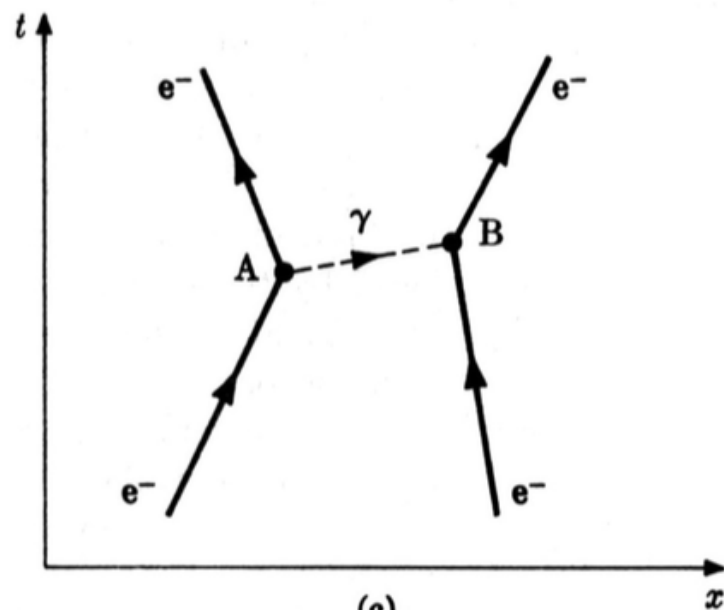
World lines of some occurrences in the particle world.

(a) Emission of photon by an atom.

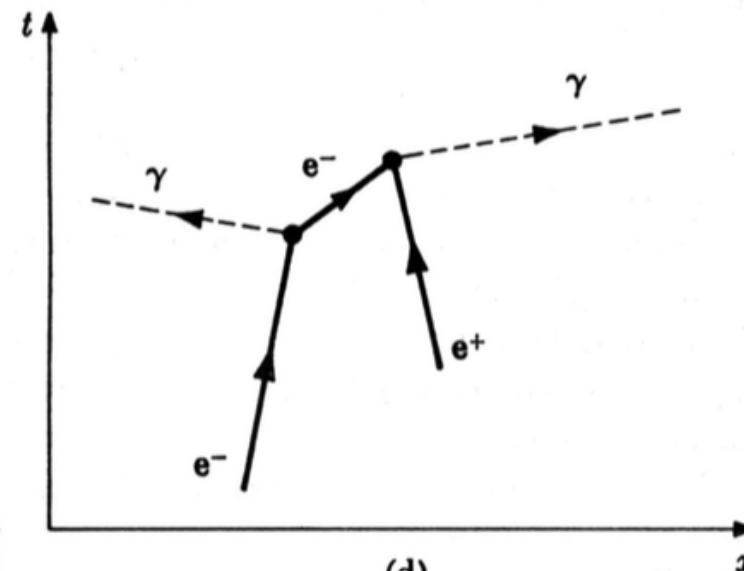
(b) Pion decay.

(c) Scattering of one electron by another.

(d) annihilation of an electron and a positron.



(c)



(d)

Diagram (a) shows the process of photon emission by an atom.

Initially (starting at the bottom of the diagram) an atom is at rest and, like the city of Flagstaff, traces out a straight vertical world line.

It then emits a photon that flies off to the right; the atom itself, having fired the photon, recoils and moves off more slowly to the left.

Note that the more slowly a particle moves, the more **nearly** vertical is its world line - no motion at all is exactly vertical.

Conversely, the faster the particle moves, the more its world line is tipped over toward the horizontal. But it can never become exactly horizontal; for that, the particle would have to travel from one place to another in no time at all.

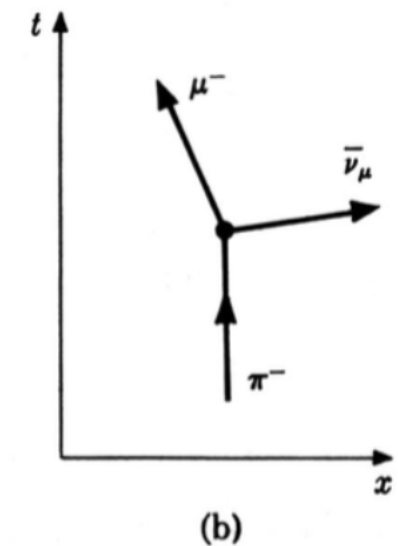
The photon line is tipped the maximum amount; the inverse of its slope is the speed of light.

Diagram (b) illustrates pion decay: $\pi^- \rightarrow \mu^- + \bar{\nu}_\mu$

At some point, indicated by the black dot, the negative pion **ceases** to exist.

It is **annihilated** and its world line terminates.

But at that same place and same time (that is, at that same point in spacetime) the negative muon and the antineutrino are **born** and fly apart, the antineutrino line being tipped almost as far as a photon line (since it moves at nearly the speed of light).



The black dot locates an event, an occurrence at a single point in space and time.

In each of the other diagrams, there are also important events. In the world of particles, every significant event is **marked** by the creation and/or annihilation of particles.

Diagram (c) illustrates one of the processes that contributes to the “scattering” (deflection) of one electron by another.

At point A, an electron emits a photon - a **virtual** photon - that is absorbed by another electron at point B.

The photon here plays the role of an “**exchange** particle.”

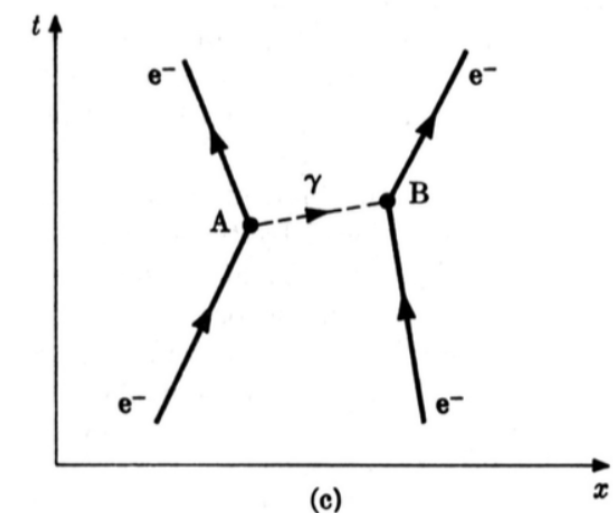
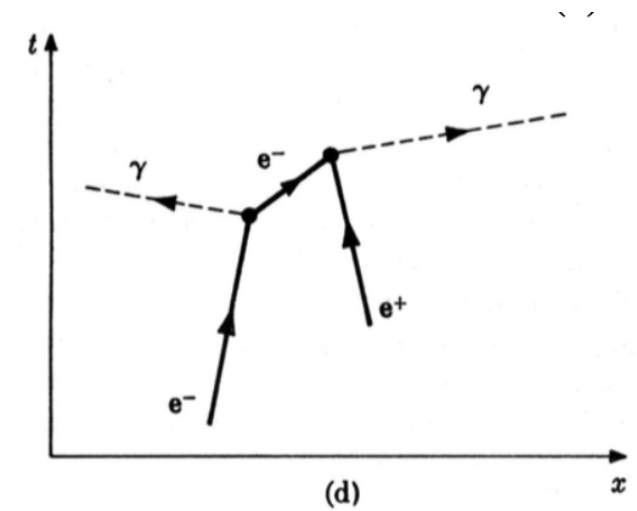


Diagram (d) shows the process of positron-electron annihilation,

This time the photons are **real**, not virtual.

The diagram illustrates a feature of all events of the particle world when examined most closely (with the highest-power microscope, so to speak).



Each basic event is a “**three-prong vertex**,” at which two fermion world lines (the electrons, in this case) and one boson line (the photon, in this case) meet.

Diagrams (b), (c), and (d) above represent, so far as we know, what is really going on at the smallest-scale level.

Not so, diagram (a).

An **atom** is not itself a fundamental particle, so its world line is really a bundle of world lines.

This means that the black dot in diagram (a) represents more activity than the diagram shows.

It is “**cover**” for finer-grained mesh of events.

We must be prepared for the possibility that the future will reveal an inner structure for other apparently simple events, even for the possibility that the apparently catastrophic acts of sudden annihilation and creation really result from a smooth, continuous flow of events over tinier regions of space and time than it has been possible to investigate so far.

This is pure speculation.

Down to the shortest distances ($\sim 10^{-17}$ m) and shortest times ($\sim 10^{-25}$ s) to which scientists have so far probed, the basic events of the particle world still **seem** to be the catastrophic events of sudden creation and destruction of the packets of field energy we call particles.

The accumulated **evidence** from experiments probing the world of the very small, together with the confirming support of the quantum theory of fields, **leads** to the following important general conclusion:

All of the interactions in nature arise from acts of **annihilation and creation** of particles at definite points in space and time.

There are actually two important ideas here.

First, all interactions involve the creation and annihilation of particles; second, these creations and annihilations do not take place over a region of space or over a span of time, but are instantaneous and localized at points.

By an “interaction” I mean simply the influence of anything on anything else.

Thus, all ordinary forces - pushes or pulls of one thing on another - are interactions.

Also, the decay of an unstable particle is the manifestation of an interaction.

The final particles are “influenced” by the initial particle - they come into existence only because the initial particle was there.

Contrast this new view of interactions with the classical view.

The Sun and the Earth “interact,” for the Earth is pulled by the Sun.

Nothing is apparently being created or destroyed, and nothing seems to be happening instantaneously at isolated spacetime points.

But, according to the new view, gravitons are constantly being emitted and absorbed, both by the Sun and by the Earth. Each act of emission or absorption occurs at an instant of time and a point in space.

The “force” the earth experiences is nothing but the accumulated effect of all of these graviton interactions.

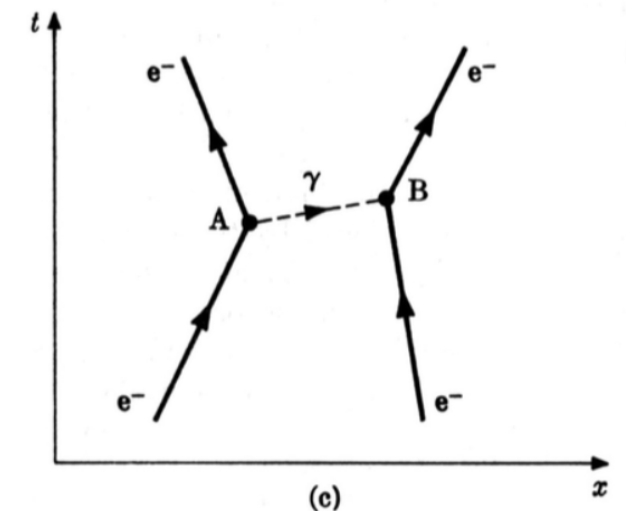
To turn to an example in the particle world that is more surely understood, consider the scattering of two electrons.

According to the **old** view, the electrons approach each other, feel a mutual repulsive force, and are deflected.

The **new** view, illustrated by diagram (c), goes deeper; it explains “why” the electrons exert forces on each other.

In that diagram you see two electrons approaching one another.

At point A, the electron on the left emits a photon and changes its speed.



At point B, the electron on the right absorbs the photon and changes its speed.

The two electrons have interacted, or exerted forces on each other, since their motion is altered. It was the **exchange** of a photon that brought about this apparent interaction.

Strictly speaking, the basic interaction was not between the two electrons at all, but between each electron and a photon.

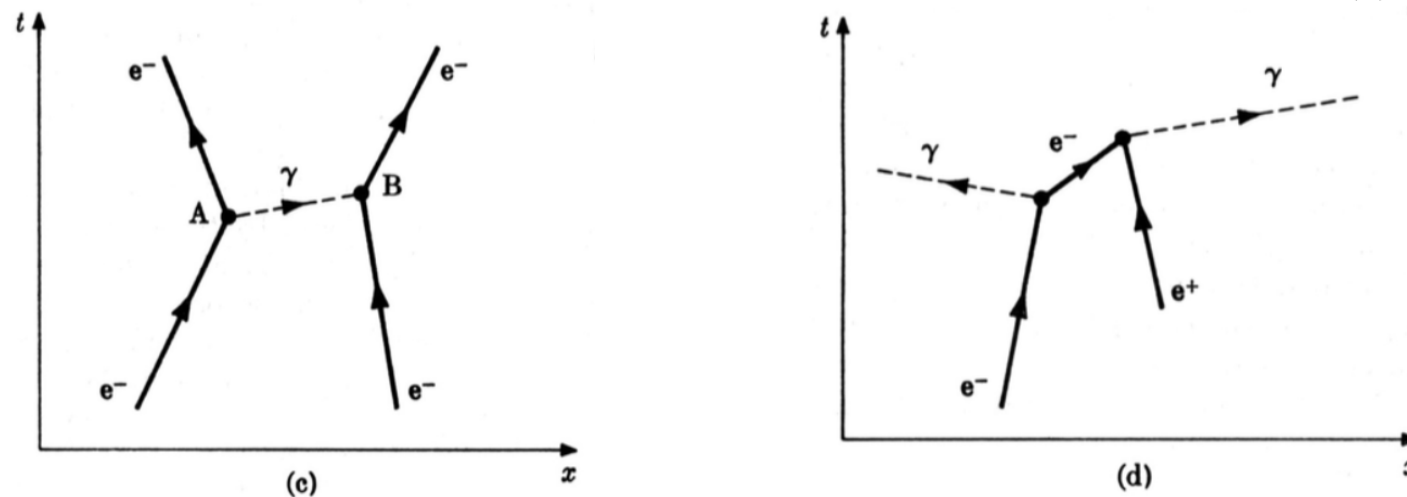
The second electron is only **indirectly** aware of the presence of the first.

The old idea of action at a distance, of a force “reaching out” from one to the other, is **completely** abandoned.

It is replaced by the idea of a “**local** interaction,” of each electron interacting locally—that is, at its own location—with a photon.

Of course, that particular diagram is only one of many; the others involve more complicated exchanges between the electrons.

The **net** effect of all the possible exchanges is the deflection of the electrons according to the ordinary repulsive electric force, but with motion that is a series of hops rather than a smooth flow.



According to the present theory of the interactions of electrons and photons, diagrams (c) and (d) are pictures of what “**really**” happens on the submicroscopic scale.

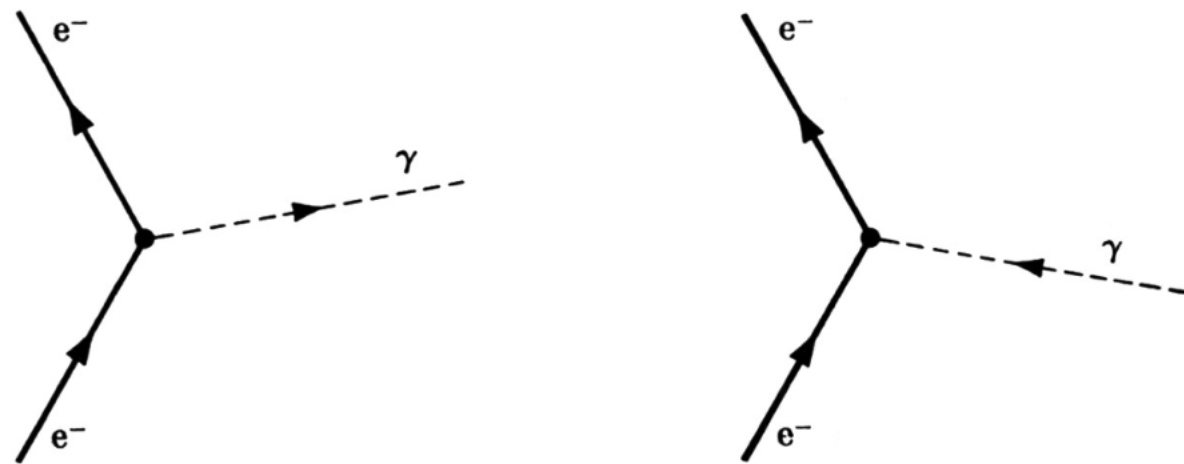
Such diagrams, as we said, are called **Feynman** diagrams, after Richard Feynman, who showed in 1949 that such pictures have an exact correspondence to mathematical expressions in the field theory of electrons and photons.

These diagrams, therefore, portray what is “**really**” happening, and provide a convenient way to catalogue the various possible processes of creation, annihilation, and exchange.

The key points of a Feynman diagram are the “**vertices**,” representing those spacetime points at which (in this example) photons are created or absorbed.

All processes involving photons, and therefore all of the interactions associated with electromagnetism, arise from elementary events of photon creation or photon annihilation.

These **fundamental** interaction events can be pictured by a single kind of vertex that looks like either of the diagrams in the figure below.



Basic electron-photon interaction vertices.

The **solid** lines represent charged particles and the **dashed** line represents a photon.

Vertices like these can be recognized in diagrams (c) and (d) above.

If the solid lines represent the world lines of electrons, for example, it appears that the fundamental interaction event may be regarded as an event in which a photon is **created or absorbed**, and an electron simultaneously changes its state of motion.

There is, however, a more general and more fruitful interpretation.

The vertex may be taken to **represent** a point where the world line of one electron terminates and the world line of another begins.

According to this view, the vertex represents a truly **catastrophic** event.

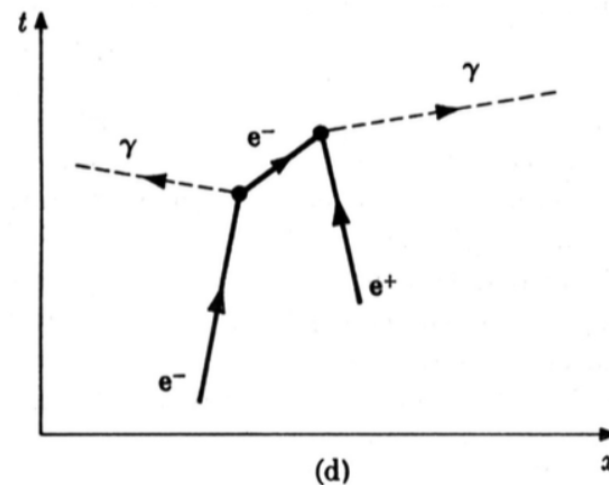
Nothing survives it.

Instead of thinking of a single electron being changed at the vertex, one can think of one electron being destroyed and another electron being created.

Since all electrons are indistinguishable, it has no real meaning to say that the outgoing electron is the same as or different from the incoming electron.

To think of the outgoing electron as a new and different electron, however, **corresponds** more closely to the mathematical theory of the fundamental interaction.

The creation-destruction interpretation also leads to a **simple, unified description** of particle events and antiparticle events.



The vertex on the right in diagram (d) **seems** to be different.

Instead of being a point where one electron world line ends and another begins, this vertex is a point where **both** an electron world line and a positron world line end.

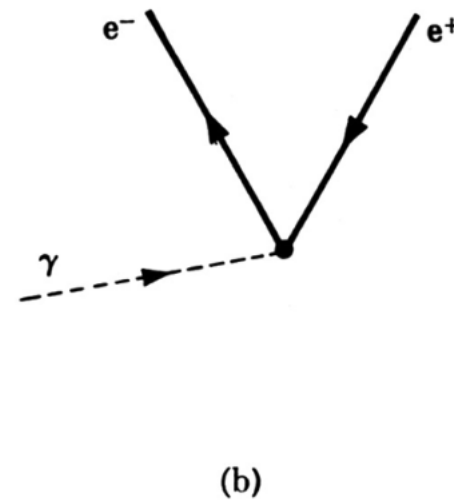
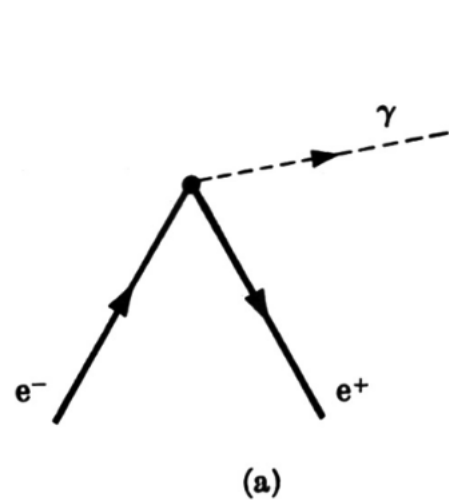
There is a simple artistic **trick** by which we can change the picture significantly.

Suppose we turn around the **arrow** on the positron line.

The arrowhead was, after all, **redundant**, since all particles move forward in time.

We can use it instead as a label to **distinguish** particles from antiparticles.

An arrowhead pointing in the “**correct**” direction will indicate a particle (for example, an electron); an arrowhead pointing in the “**wrong**” direction will indicate an antiparticle (a positron).



Additional fundamental electron-photon interaction vertices that include antiparticles (positrons).

Using this revised notation, we **show** in the diagram above an electron- positron annihilation vertex and an electron-positron creation vertex.

These vertex diagrams involving positrons look like twisted versions of the fundamental vertex diagrams illustrated above.

The generalized **conclusion** is that the fundamental electron-photon vertex with its limbs twisted in all possible directions in spacetime represents all the possible basic interactions among electrons, positrons, and photons.

This provides a **magnificently** simple and general view of the underlying basis of all electromagnetic phenomena.

What Feynman showed when he discussed the connection between such world-line diagrams and the structure of the mathematical theory of the electron-positron-photon interaction was that the device of the reversed arrowhead is **much more** than an artistic trick.

According to the field theory of electrons, the creation of a positron is “equivalent” to the annihilation of an electron (they are not identical processes, but the theory says that whenever one can happen, the other must also be able to happen).

Moreover, the mathematical description of a positron field propagating forward in time is identical with the description of an electron field propagating **backward** in time.

It is perfectly **possible and consistent** to think of particles moving backward in time as well as forward.

This circumstance **need not** lead one to deep philosophical conclusions, although philosophical implications are hard to escape.

The positron **may** be described as an electron moving backward in time, but it does not have to be so described.

An **alternative** description is equally possible in which the positron is a normal particle moving forward in time.

Nevertheless, this picture of backward motion in time is a **tantalizing** one that simplifies the view of elementary interactions and provides a “natural explanation” for the existence of antimatter.

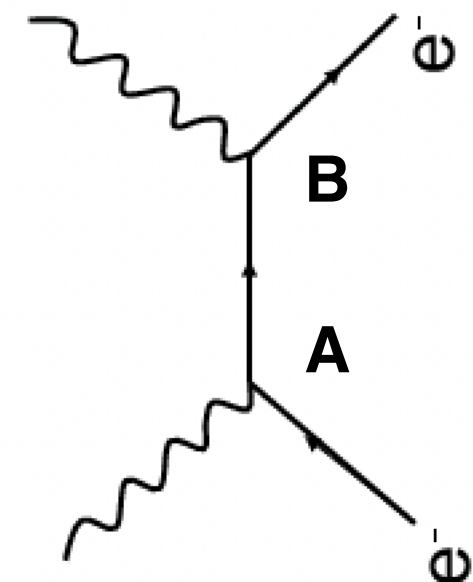
Consider the Feynman diagram illustrated below, for example.

According to the **normal** view of time unrolling in one direction, you start at the bottom of the diagram and read up. (Put a horizontal straight edge over the bottom of the diagram and move it slowly upward, noticing what lines it touches.)

First an electron and a photon are approaching each other.

At vertex A, the photon is absorbed by the electron.

At vertex B, that electron emits a photon.



I will have much more to say as we proceed!!

Feynman diagram for the process of photon- electron scattering.

QED: The Strange Theory of Light and Matter (ala Feynman)

Introduction

Many adults are very curious about physics and often ask me for explanations.

I generally do all right, but eventually I fail at most interesting point - we always get hung up on crazy ideas of quantum mechanics.

I can't explain these ideas in an hour or an evening.

It requires a much longer time in general.

Over the years I have developed a class for adults that teaches quantum mechanics - some of you have taken that class - it presents formal theoretical ideas and requires some math development.

For first part of this class we are going to try, however, following Feynman, to present quantum mechanics in arena of photons and electrons without the need for much math development.

This will allow beginners to understand the way quantum mechanics works
and

supplement ideas for those of you that have taken my other class.

First we study the parts of physics that are known, rather than parts that are unknown.

People always asking for latest developments in unification of this theory with that theory,
don't ever learn about any of theories that we know pretty well.

Always want to know about things that we don't know well.

So, I will tell you about a subject that is very thoroughly analyzed.

I will describe as accurately as I can the strange theory of light and matter - interaction of light and electrons.

Called quantum electrodynamics, or QED for short.

Physics has history of synthesizing many phenomena into a few theories.

For instance, in early days there were phenomena of motion and phenomena of heat; there were phenomena of sound, of light, and of gravity.

But soon discovered, after Sir Isaac Newton explained laws of motion, that some of apparently different things were aspects of same thing.

For example, phenomena of sound could be completely understood as motion of atoms in air.

So sound no longer considered something in addition to motion.

Also discovered that heat phenomena easily understandable from laws of motion -> large areas of physics theory synthesized into simplified theory.

Theory of gravitation, on other hand, not understandable from laws of motion, and even today stands isolated from other theories.

i.e., gravitation is, so far, not understandable in terms of other phenomena.

After synthesis of phenomena of motion, sound, and heat, was discovery of number of phenomena that call electrical and magnetic.

In 1873 these phenomena were synthesized with phenomena of light and optics into single theory by James clerk Maxwell

—-> proposed light is an electromagnetic wave.

So at that stage, there were laws of motion, laws of electricity and magnetism, and laws of gravity.

Around 1900 a theory was developed to explain what matter was.

Called the electron theory of matter, and said there were little charged particles inside of atoms.

Theory evolved gradually to include heavy nucleus with electrons “going around it”.

Attempts were made to understand motion of electrons going around nucleus using mechanical laws

analogous to way Newton used laws of motion to figure out how earth went around sun

and they were a real failure: all predictions came out wrong!

Incidentally, theory of relativity - a great revolution in physics - was also developed at about that time.

But compared to discovery that Newton’s laws of motion quite wrong in atoms, theory of relativity was only minor modification.

Working out another system to replace Newton’s laws took a long time because phenomena at atomic level were quite strange

One had to lose common sense to perceive what happening at atomic level.

Finally, in 1926, an “uncommon-sensy” theory was developed to explain “new type of behavior” of electrons in matter.

Looked cockeyed, but it was not ::: called theory of quantum mechanics.

Word ‘quantum’ refers to peculiar aspect of nature that goes against common sense.

It is this aspect that is subject of this discussion.

Theory of quantum mechanics also explained details like

why oxygen atom combines with 2 hydrogen atoms to make water, and so on.

Quantum mechanics supplied theory behind chemistry.

So, fundamental theoretical chemistry is really physics.

Because theory of quantum mechanics could explain all chemistry and various properties of substances, it was tremendous success.

But there still existed the problem of interaction of light and matter.

Maxwell’s theory of electricity and magnetism had to be changed to work with
the new principles of quantum mechanics.

So a new theory, quantum theory of interaction of light and matter, called by horrible name
“quantum electrodynamics”,

finally developed by physicists in 1929.

But theory was very troubled.

If one calculated something approximately, it always gives a reasonable answer.

But if tried to compute more accurately - adding corrections that thought were going to be small, found they were very large - in fact, gave **infinity**!

So turned out couldn't really compute anything beyond certain accuracy.

An approximation to complete theory gave good results but an exact calculation using the theory **always** gave infinity.

Paul Dirac, using theory of relativity, developed a relativistic theory of electrons, but did not completely take into account all effects of electron's interaction with light.

Dirac's theory said electron had magnetic moment - something like a little magnet - that had strength of exactly 1 in certain units.

Then in 1948 it was discovered in experiments that actual number was closer to 1.00118 (with uncertainty of 3 on last digit).

It was known that electrons interact with light, so some small correction was expected.

Also expected that correction would be understandable from new theory of quantum electrodynamics.

But when was calculated, instead of 1.00118 result was **infinity** - which is wrong - experimentally!

Problem of how to calculate things in quantum electrodynamics straightened out by Schwinger, Tomonaga, and Feynman in 1948 (19 years later!!).

Schwinger was first to calculate this correction using new "shell game"; his theoretical value was 1.00116

-> close enough to experimental number to show that they were on right track.

At last, had quantum theory of electricity and magnetism with which one could calculate!

This is theory I am going to describe.

Theory has been put through wringer - some recent numbers:

experiments says $1.0011596521873 \pm 0.00000000000028$

theory says $1.0011596521865 \pm 0.00000000000035$

Accuracy of numbers - if measure distance from Los Angeles to New York to this accuracy,
would be exact to thickness of a single human hair.

That's how delicately quantum electrodynamics has, in past 90 years, been checked
- both theoretically and experimentally.

There are other things in quantum electrodynamics that have been measured with comparable accuracy,
which also agree very well.

Things have been checked at distance scales that range from one hundred times size of earth down to
1/100 size of atomic nucleus.

These numbers are meant to make you believe theory is probably not too far off!

Before we are through, I will describe how calculations are made.

To impress you with vast range of phenomena that theory of quantum electrodynamics describes

However, it is easier to say it backwards:

theory describes all phenomena of physical world except gravitational effect,
thing that holds you in your seats (actually, that's combination of gravity and politeness),
and radioactive phenomena, which involve nuclei shifting in their energy levels.

So if leave out gravity and radioactivity (nuclear physics), what is left?

Gasoline burning in automobiles, foam and bubbles, hardness of salt or copper, stiffness of steel, etc.

In fact, biologists are trying to interpret as much as can about life in terms of chemistry, and as already explained,

theory behind all of chemistry is quantum electrodynamics.

Must clarify something:

When say that all phenomena of physical world can be explained by theory, don't really know that.

Most phenomena familiar with involve such tremendous numbers of electrons that hard to follow that complexity.

In such situations, can use the theory to figure roughly what ought to happen and that is what happens, roughly, in those circumstances.

But if arrange in laboratory experiment involving just few electrons in simple circumstances, then can calculate what might happen very accurately, and can measure very accurately, too.

Whenever do such experiments, theory of quantum electrodynamics works very well.

Physicists always checking to see if something the matter with theory.

That's the game, because if something wrong,

it's interesting!

But so far, have found nothing wrong with theory of quantum electrodynamics

jewel of physics

proudest possession.

Theory of quantum electrodynamics also model for new theories that attempt to explain nuclear phenomena, things that go on inside nuclei of atoms.

If think of physical world as stage,

then actors would be not only electrons,

which are outside nucleus in atoms,

but also quarks and gluons and so forth

- dozens of kinds of particles - inside nucleus.

And though “actors” appear quite different from one another,

all act in certain style - strange and peculiar style - “quantum” style.

Will deal with particle physics later.

In meantime, going to tell you about photons - particles of light - and electrons, to keep it simple.

Because it's way they act that is important, and the way they act is very interesting.

Now know what lectures about -

will you completely understand it?

Generally, taught to physics students in 3rd/4th year of graduate school.

So, NO, you are not going to completely understand it.

Why, are you doing this?

Why are you here - when won't completely understand this stuff?

I say do not turn away because you won't completely understand stuff

Physics students don't completely understand it either

I don't completely understand it — — - Nobody does.

What is meant by “understanding”?

When have a lecture —> many reasons why might not understand speaker

language bad

doesn't say what means to say

or says it upside down - and hard to understand.

That's rather trivial matter, and that will not be a problem in this class.

Another possibility, especially if lecturer is physicist,

is that they use ordinary words in funny way.

Physicists often use ordinary words such as “work” or “action” or “energy” or even, as you shall see, “light” for some technical purpose.

Thus, when I talk about “work” in physics, I don’t mean same thing as when I talk about “work” on street.

During lecture might use one of those words without noticing that being used in unusual way - try my best to catch myself - that’s my job - but is error easy to make.

Next reason that might not understand what I am telling you is,

while am describing how Nature works, won’t understand **why** Nature works that way.

But you see, **nobody** understands that. Can’t explain WHY Nature behaves in peculiar ways.

Digress: Physics/Philosophy

Finally, is this possibility: after tell you something,

just can’t believe it - can’t accept it - don’t like it.

Little screen comes down and you don’t listen anymore.

Going to describe how Nature is - and if don’t like it - going to get in way of your understanding it.

Problem that physicists learned to deal with: learned to realize that whether like a theory or don’t like theory is not essential question.

It is, whether or not theory gives predictions/explanations that agree with experiment.

It is not question of whether theory philosophically delightful, or easy to understand, or perfectly reasonable from point of view of common sense.

Theory of quantum electrodynamics describes Nature as absurd from point of view of common sense.

And agrees fully with experiment.

So hope you can accept Nature as She is - absurd.

Going to have fun talking about absurdity, because I find it delightful.

Please don't turn off because can't believe Nature so strange.

Just hear all out, and I hope you'll be as delighted as I am when we're through.

How am I going to explain things that we don't explain to students until 3rd-year graduate students?

An Analogy involving Mayans

Mayan Indians interested in rising and setting of Venus as morning "star" and as evening "star" -
were very interested in when it would appear.

After years of observation, noted that 5 cycles of Venus very nearly equal to 8 of "nominal years" of 365 days (were aware that true year of seasons was different and made calculations of that also).

To make calculations, Maya invented system of bars and dots to represent numbers (including zero), and had rules by which to calculate and predict not only risings and settings of Venus, but other celestial phenomena, such as lunar eclipses.

In those days, only few Maya priests could do such elaborate calculations.

Now, suppose we were to ask one of them how to do just one step in process of predicting when Venus will next rise as morning star - subtracting two numbers.

And let's assume that, unlike today, had not gone to school and did not know how to subtract.

How would priest explain to us what subtraction is?

Could either teach us numbers represented by bars and dots and rules for “subtracting” them,
or could tell us what was really doing:

“Suppose want to subtract 236 from 584.

1st, count out 584 beans and put in pot.

Then take out 236 beans and put to one side.

Finally, count beans left in pot

number is result of subtracting 236 from 584”.

Might say,

“My Quetzalcoatl!

What tedium - counting beans, putting them in, taking them out - what a job!”

To which priest would reply,

“That’s why have rules for bars and dots.

Rules tricky, but are much more efficient way of getting answer than counting beans.

Important thing

makes no difference as far as answer concerned

can predict appearance of Venus by counting beans (slow, easy to understand) or by using tricky rules (much faster, but must spend years in school to learn them)”.



To understand how some mathematics works - as long as don't have to actually carry out - is not so difficult.

That's my position:

I'm going to explain what physicists are doing when predicting how Nature will behave, but not going to teach you any tricks so can do it efficiently.

Will discover that in order to make any reasonable predictions with new scheme of quantum electrodynamics,

will have to draw awful lot of little arrows on piece of paper.

Takes 7 years - 4 undergraduate and 3 graduate - to train physics students to do that in tricky, efficient way(using mathematics).

That's how I will skip seven years of education in physics:

By explaining quantum electrodynamics to you in terms of what they we are really doing, I hope you will be able to understand it better than some of the advanced students!

Taking example of Maya one step further,

could ask priest why 5 cycles of Venus nearly = 2,920 days, or 8 years.

Would be all kinds of theories about why, such as,

“20 is an important number in our counting system, and if you divide 2,920 by 20, you get 146, which is one more than a number that can be represented by the sum of two squares in two different ways”, and so forth.

But that theory would have nothing to do with Venus, really.

In modern times, have found that theories of this kind not useful.

So again, not going to deal with fact that Nature behaves in peculiar way —- She does;
no good theories to explain that.

What I have done so far is get you into right mood to listen to me.

Otherwise, I have no chance. So now we're off, ready to go!

Begin with light.

When Newton started looking at light,

1st thing he found was that white light is mixture of colors.

Separated white light with prism into various colors,

but when put light of one color - red, for instance - through another prism,

found could not be separated further.

So Newton found that white light is mixture of different colors, each of which is pure in sense that can't be separated further.

(In fact, particular color of light can be split one more time in different way, according to so-called "polarization".

This aspect of light not vital to understanding character of quantum electrodynamics, so for simplicity will leave it out - at expense of not giving you absolutely complete description of theory.

This slight simplification will not remove, any real understanding of what we will be talking about.

Still, must be careful to mention all of things I leave out.)

When say “light”, don’t mean simply light you can see, from red to blue.

Turns out that visible light just part of long scale that’s analogous to musical scale in which there are notes higher than you can hear and other notes lower than you can hear.

Scale of light described by number - called frequency - and as number get higher, light goes from red to blue to violet to ultraviolet.

Can’t see ultraviolet light, but it can affect photographic plates.

It’s still light - only number(frequency) different.

Shouldn’t be so provincial: what one can detect directly with own instrument - eye, isn’t only thing in world!

If continue simply to change number, go out into X-rays, gamma rays, and so on.

If change the number in other direction, go from blue to red to infrared (heat) waves, then television waves, and radio waves -> all of that is “light”.

Going to use just red light for most of examples, but theory of quantum electrodynamics extends over entire range described, and is theory behind all these various phenomena.

Newton thought light made up of particles - called them “corpuscles”

- and was correct (but reasoning that he used to come to decision was erroneous).

Know now that light made of particles because we can take very sensitive instrument that makes clicks when light shines on it, and if light gets dimmer, clicks remain just as loud - just fewer of them.

Thus light something like raindrops - each little lump of light called a **photon** - and if light is all one color, all “raindrops” are “same size” or “same frequency”.

Human eye is very good instrument:

takes only about 5 or 6 photons to activate nerve cell and send message to brain.

If we were evolved little further so could see 10 times more sensitively, wouldn't need this discussion

- would all have seen very dim light of one color as series of intermittent little flashes of equal intensity —> would fully understand QED

Might wonder how it is possible to detect a single photon.

One instrument that can do this = photomultiplier, and now describe briefly how it works:

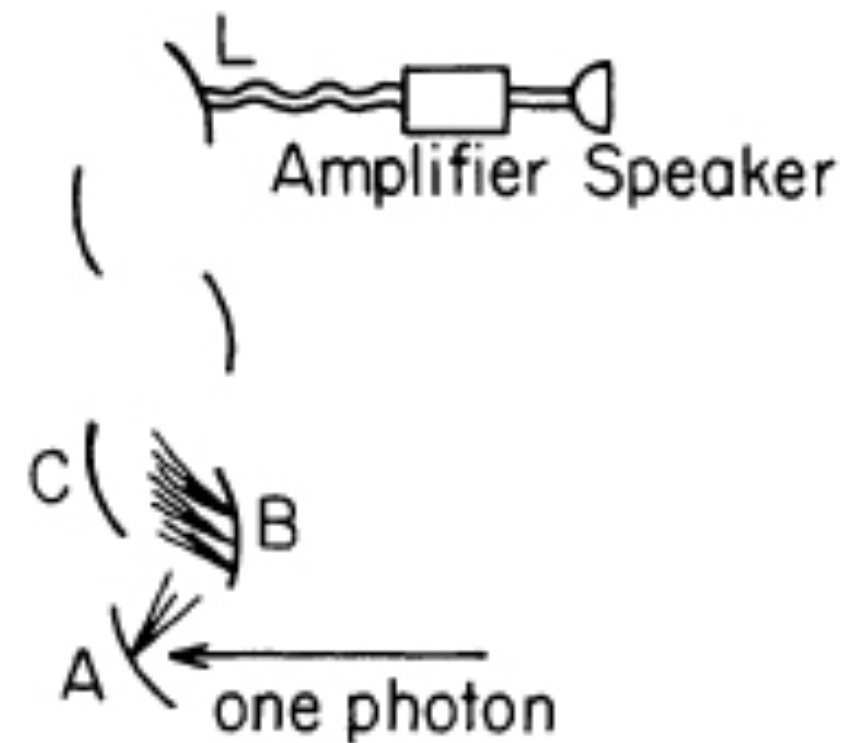
When photon hits metal plate A at bottom (Figure right), causes electron to break loose from one of atoms in plate.

Free electron strongly attracted to plate B (has positive charge on it)

and hits with enough force to break loose 3 or 4 electrons.

Each of electrons knocked out of plate B attracted to plate C (also charged),

and collision with plate C knocks loose even more electrons.



Process repeated 10 or 12 times, until billions of electrons,
enough to make sizable electric current, hit last plate, L.

Current amplified by regular amplifier and sent through speaker to make audible clicks.

Each time photon of given color hits photomultiplier,
click of uniform loudness heard.

If put whole lot of photomultipliers around

and let some very dim light shine in various directions,
light goes into one multiplier or another
and makes click of full intensity.

It is **all or nothing**:

if one photomultiplier goes off at given moment,

none of others goes off at same moment

(except in rare instance that 2 photons happened to leave light source at same time).

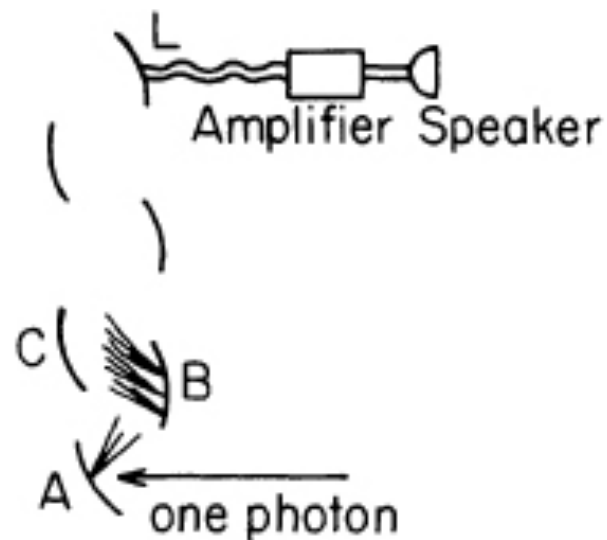
There is no splitting of light into “half particles” that go different places.

Want to emphasize that light comes in this form - particles.

Very important to know that light behaves like particles,

especially for those of you who have gone to school,

where you were probably told something about light behaving like waves.



I'm telling you way it does behave - like particles. **Waves are only when very large number together!!**

Might say that it's just photomultiplier that detects light as particles,

but no, every instrument that has been designed to be sensitive enough to detect weak light
has always ended up discovering same thing: light is made of particles.

AS I mentioned above, Intense beams of light(huge numbers of photons/second) behave like wave so interference experiments are done with intense beams.

Going to assume familiar with properties of light in everyday circumstances

- things like, light goes in straight lines;
- bends when goes into water;
- when reflected from surface like mirror,
angle at which light hits surface equal to angle at which leaves surface;
- light can be separated into colors;
can see beautiful colors on mud puddle when little bit of oil on it;
- lens focuses light, and so on.

Going to use these phenomena that are familiar with in order to illustrate truly strange behavior of light;

Going to explain familiar phenomena in terms of theory of quantum electrodynamics.

Showed photomultiplier to illustrate essential phenomenon that you may not have been familiar with
that light made of particles - but now are familiar with that, too!

Now, you are all familiar with phenomenon that light partly reflected from some surfaces, such as water.
There are many romantic paintings of moonlight reflecting from lake.

When you look down into water you can see what's below surface (especially in daytime),
but can also see a reflection from surface.

Glass is another example:

if have lamp on in room and looking out through window during daytime,
can see things outside through glass as well as dim reflection of lamp in room.

So light partially reflected from surface of glass.

Before I go on, I want you to be aware of simplification we are going to make that will correct later on:

When I talk about partial reflection of light by glass,

I am going to pretend that light reflected only by the surface of glass.

In reality, a piece of glass is terrible monster of complexity

where huge numbers of electrons are jiggling about.

When photon comes through, it interacts with electrons throughout glass, not just on surface.

Photon and electrons do kind of dance,

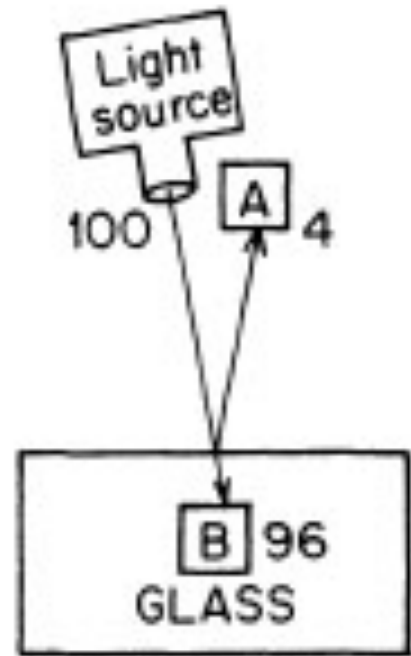
the net result of which is same as if photon hit only surface.

So I will make this simplification for while.

Later on, I will show you what actually happens inside glass so you can understand why result is same.

Now describe an experiment, and discuss its surprising results.

In this experiment some photons of same color - red light - are emitted from light source (Figure below) down toward a block of glass.



Photomultiplier placed at A,

above glass, to catch any photons reflected by front surface.

To measure how many photons get past front surface,

another photomultiplier placed at B, inside glass.

Never mind obvious difficulties of putting photomultiplier inside block of glass - remember I am a theoretical physicist!!

What are results of experiment?

For every 100 photons that go straight down toward glass at 90° angle of incidence, an average of 4 arrive at A and 96 arrive at B.

So “partial reflection” means that 4% of photons are reflected by front surface of glass, while the other 96% are transmitted.

Already we are in great difficulty: how can light be partly reflected?

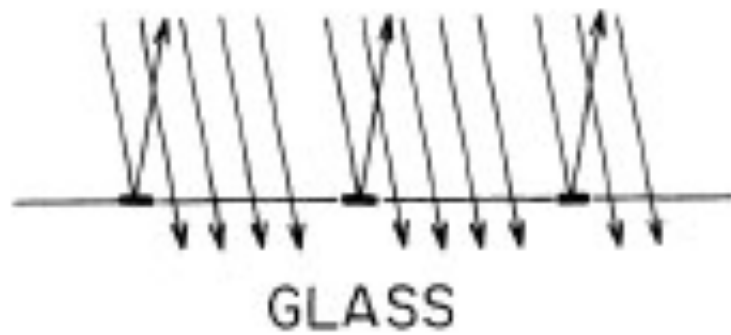
Each photon ends up at A or B - how does photon “make up its mind” whether should go to A or B?

May sound like joke, but can’t just laugh;

we are going to have to explain how in terms of theory!

Partial reflection is already a deep mystery, and was a very difficult problem for Newton.

Are several possible theories that could make up to account for partial reflection of light by glass.
One of them is that 96% of surface of glass is “holes” that let light through,
while other 4% of surface is covered by small “spots” of reflective material (Figure below).



Newton realized this is not possible explanation.

Soon we will encounter strange feature of partial reflection
that will drive you crazy
if you try to stick to theory of “holes and spots”
or to any other such “reasonable” theory!

Another possible theory is that photons have some kind of internal mechanism
“wheels” and “gears” inside turning in some way
so that when photon “aimed” just right, it goes through glass,
and when not aimed just right, it reflects.

Can check theory by trying to filter out photons that are not aimed right
by putting a few extra layers of glass between source and first layer of glass.

After going through filters,
photons reaching glass should all be aimed right, and none should reflect.

Trouble with theory is, it doesn’t agree with experiment:

even after going through many layers of glass, 4% of photons reaching given surface reflect off it.

Can try to invent reasonable theory that can explain how photon “makes up its mind”
whether to go through glass or bounce back.

But turns out that it is impossible to predict which way any particular photon will go.

Philosophers have said that if **same** circumstances don't always produce **same** results,
predictions are impossible and science will collapse.

Here is such a circumstance - **identical** photons are always coming down
in same direction to the same piece of glass and that produces different results.

Cannot predict whether given photon will arrive at A or B.

All we can predict is that out of 100 photons that come down,
average of 4 will be reflected by front surface.

Does this mean that physics, the science of great exactitude,
is reduced to calculating only the probability of an event,
and not predicting exactly what will happen(was only an assumption anyway)?

Yes -> seems like a retreat, but that's the way nature is:

Nature often permits us to calculate only probabilities. Yet science has not collapsed.

While partial reflection by single surface is a deep mystery and a difficult problem,
partial reflection by two or more surfaces is absolutely mind-boggling.

Let me show you why.

Do 2nd experiment,

measure partial reflection of light by two surfaces.

Replace block of glass with very thin sheet of glass -

two surfaces exactly parallel to each other -

and place photomultiplier below sheet of glass, in line with light source.

Photons can reflect from either front surface or back surface to end up at A;

all others will end up at B (Figure below).

Might expect front surface to reflect 4% of light

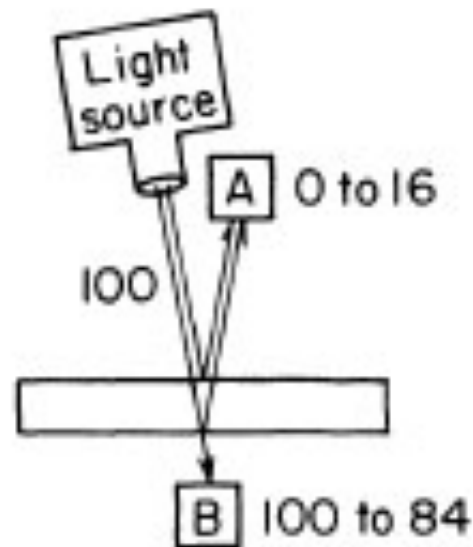
and back surface to reflect 4% of remaining 96%,

making total of about 8%.

So should find that out of every 100 photons that leave light source,

about 8 arrive at A.

What actually happens under these carefully controlled experimental conditions is, number of photons arriving at A is rarely 8 out of 100.



With some sheets of glass, consistently get reading of 15 or 16 photons - twice expected result!

With other sheets of glass, consistently get only 1 or 2 photons.

Other sheets of glass have partial reflection of 10%;

Some eliminate partial reflection altogether!

What can account for such crazy results?

After checking various sheets of glass for quality and uniformity,
we discover that they differ only slightly in thickness.

To test idea that amount of light reflected by 2 surfaces depends on thickness of glass,
we do a series of experiments:

Starting out with thinnest possible layer of glass,
count how many photons hit photomultiplier at A each time 100 photons leave light source.
Then replace layer of glass with slightly thicker one and make new counts.
After repeating process few dozen times, what are results?

With thinnest possible layer of glass,
find number of photons arriving at A nearly always 0 - sometimes 1.

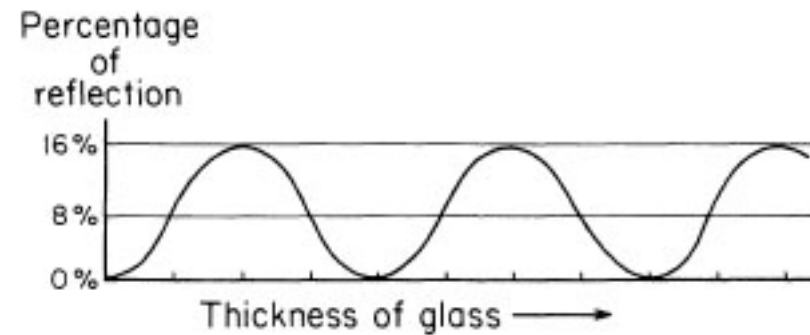
When replace thinnest layer with slightly thicker one,
find amount of light reflected higher - closer to expected 8%.

After few more replacements count of photons arriving at A increases past 8% mark.

As continue to substitute still “thicker” layers of glass - up to about 5 millionths of inch now -
amount of light reflected by 2 surfaces reaches maximum of 16%, and then goes down,
through 8%, back to 0

if layer of glass just right thickness, there is no reflection at all. (Try to do that with holes/spots theory!)

With gradually thicker and thicker layers of glass,
partial reflection again increases to 16% and returns to 0
cycle that repeats itself again and again (Figure below).



Newton discovered these oscillations and did experiment that could be correctly interpreted only if oscillations continued for at least 34,000 cycles! He was using a very crude non-pure light source.

Today, with lasers (which produce very pure, monochromatic light - single wavelength),
can see this cycle still going strong after more than 100,000,000 repetitions
which corresponds to glass more than 50 meters thick.

We don't see this phenomenon every day because light sources are normally not monochromatic.

Turns out that prediction of 8% is right as **overall average**

(since actual amount varies in regular pattern from 0 to 16%),

but exactly right only twice each cycle

Like a stopped clock (which is right twice a day).

How can we explain strange feature of partial reflection that depends on thickness of glass?

How can front surface reflect 4% of light (confirmed in 1st experiment)

and then, by putting 2nd surface at just right distance below,

can somehow “turn off” reflection?

And placing 2nd surface at slightly different depth,
can “amplify” reflection up to 16%!

Can it be that back surface exerts some kind of influence
or effect on ability of front surface to reflect light?

What if put in 3rd surface?

With 3rd surface, or any number of subsequent surfaces,
amount of partial reflection again changed.

Find ourselves chasing down through surface after surface with this theory,
wondering if have finally reached last surface.

Does photon have to do that in order to “decide” whether to reflect off front surface?

Does photon have to “know” somehow how many surfaces are coming up before getting to them?

Newton made ingenious arguments concerning problem, but realized, in end, that he had not developed a satisfactory theory.

For many years after Newton, partial reflection by 2 surfaces was happily explained by theory of waves(constructive/destructive interference),

but when experiments were made with very weak light hitting photomultipliers

i.e, switch from beams with many photons per sec to few per second, the wave theory collapsed:

As light got dimmer and dimmer, the photomultipliers kept making full-sized clicks just fewer of them.

Light was behaving as particles - as we said earlier!!

Situation today

We haven't got good model to explain partial reflection by 2 surfaces;

but we have set of rules(QED) —> we can calculate

probability that particular photomultiplier will be hit by photon reflected from sheet of glass.

I have chosen this calculation as 1st example of methods provided by theory of quantum electrodynamics on purpose.

I am going to show “how we count the beans”

i.e., what physicists do to get right answer from the theory.

I am not going to explain how photons actually “decide” whether to bounce back or go through - that is not known.

Probably, the question has no meaning - take my QM class to understand that.

Will only show how to calculate correct probability

that light will be reflected from glass of given thickness,

because that is the only thing physicists know how to do!!!! We only answer HOW not WHY!!!

What we do to get answer to problem is analogous to things have to do to get answer to every other problem explained by quantum electrodynamics.

You have to brace yourselves for this -

Not because difficult to understand, but because it seems like an absolutely ridiculous procedure:

All we do is draw little arrows on piece of paper - that's all!

Now, what does arrow have to do with chance that particular event will happen?

According to rules of “how we count the beans”

probability of event = square of length of some arrow.

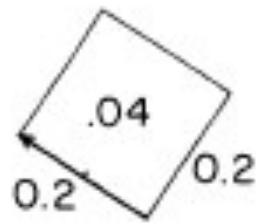
Let me illustrate. (This stuff replaces a bunch of complicated mathematics that requires QM to understand).

For example,

in 1st experiment (when measuring partial reflection by front surface only),

probability that photon would arrive at photomultiplier at A was 4%.

That corresponds to arrow whose length is 0.2, because 0.2 squared is 0.04 (see figures below).



In 2nd experiment (when replacing thin sheets of glass with slightly thicker ones),

photons bouncing off either front surface or back surface arrived at A.

How do we draw arrow to represent this situation?

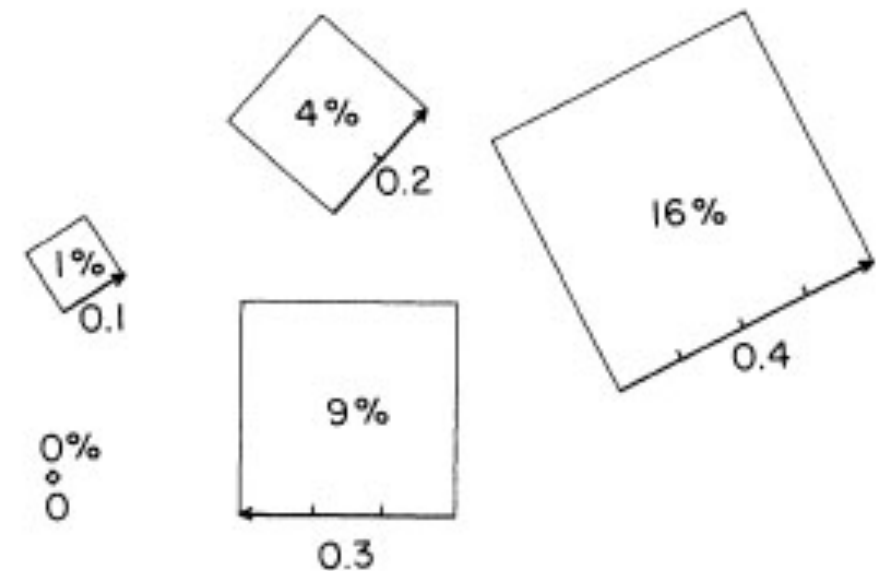
Length of arrow

must range from 0 to 0.4

to represent probabilities of 0 to 16%,

depending on thickness of glass (Figure right).

Start by considering various ways that photon could get from source to photomultiplier at A.



Since making simplification that light only bounces off either front surface or back surface

—> 2 possible ways photon could get to A.

In this case draw 2 arrows

one for each way event can happen

and then combine them into “final arrow”

whose square represents probability of event.

In QM, lengths are “amplitudes” and
their squares are “probabilities”

If had been 3 different ways event could have happened,

would have drawn 3 separate arrows before combining them.

Now, we show how combine arrows.

Say want to combine arrow x with arrow y (Figure right).



All have to do is put head of x against tail of y

(without changing directions),

and draw final arrow from tail of x to head of y.

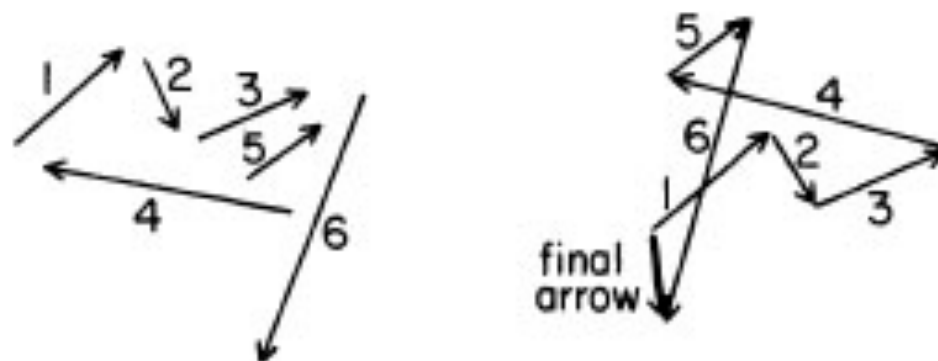
That’s all there is to it!! (that is just **ordinary vector addition!!!!**)

Can combine any number of arrows in this manner (technically, called “adding arrows”).

Each arrow tells you how far, and in what direction, to move in a dance.

“vector addition”

Final arrow tells you what single move to make to end up in same place (Figure below).



Now,

what are specific rules that determine length and direction of each arrow
that combine in order to make final arrow?

In first case, will be combining two arrows

one representing reflection from front surface of glass,
and other representing reflection from back surface.

Take length first.

As saw in 1st experiment (where we put photomultiplier inside glass),
front surface reflects about 4% of photons that come down.

Means “front reflection” arrow has length of 0.2.

Back surface of glass also reflects 4%,

so “back reflection” arrow’s length also 0.2.

To determine direction of each arrow, (process is called “time evolution” in QM)

imagine we have a stopwatch that can time photon as it moves.

Imaginary stopwatch has single hand that turns around very, very rapidly.

When photon leaves source, start stopwatch.

As long as photon moves, stopwatch hand turns (about 36,000 times per inch for red light);
when photon ends up at photo-multiplier, stop watch.

Hand ends up pointing in certain direction $\longrightarrow$ direction will draw arrow.

Need one more rule in order to compute answer correctly:

When considering path of photon bouncing off front surface of glass (air to glass),

reverse direction of arrow,

i.e., draw back reflection arrow pointing in same direction as stopwatch hand (glass to air)

but draw front reflection arrow in opposite direction (what is different?).

think interface

Now, draw the arrows for case of light reflecting from extremely thin layer of glass.

To draw front reflection arrow,

imagine photon leaving light source (stopwatch hand starts turning),

bouncing off front surface,

and arriving at A (stopwatch hand stops).

Draw little arrow of length 0.2 in direction opposite that of stopwatch hand (Figure below).

To draw back reflection arrow,

imagine photon leaving light source (stopwatch hand starts turning),

going through front surface and bouncing off back surface,

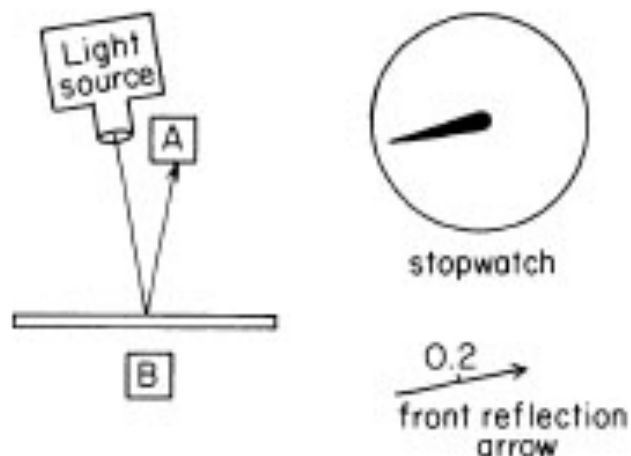
and arriving at A (stopwatch hand stops).

Stopwatch hand is pointing in almost same direction,

because photon bouncing off back surface of glass

takes only slightly longer to get to A

since it goes through **extremely thin** layer of glass twice.

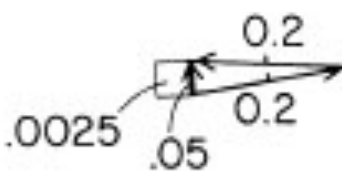
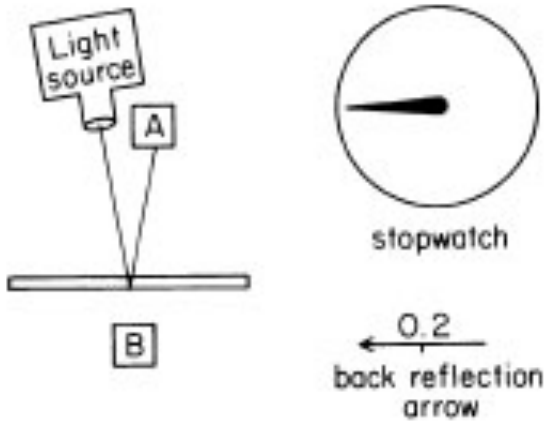


Now draw little arrow of length 0.2 in same direction that stopwatch hand is pointing (Figure below).

Now combine 2 arrows.

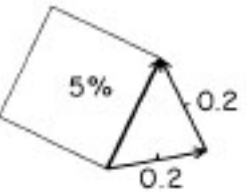
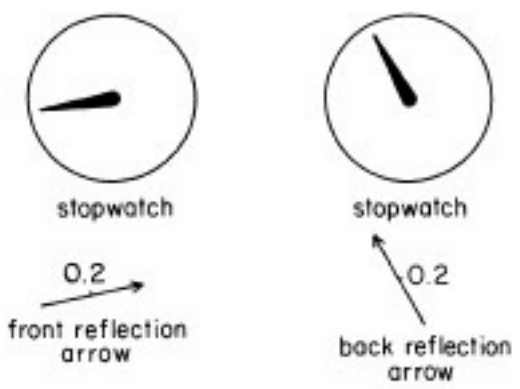
Since both have same length but are pointing in opposite directions,
final arrow has length of nearly 0,
and square even closer to 0.

Thus, probability of light reflecting from an infinitesimally thin layer of glass is essentially 0 (Figure below).



That is it - that is the way to calculate!!!! Let us continue to more difficult cases.

When replace thinnest layer of glass with slightly thicker one,
photon bouncing off back surface takes bit longer
to get to A then in 1st example;
stopwatch hand therefore turns little bit more before stops,
and back reflection arrow ends up
in slightly greater angle relative to front reflection arrow.



Final arrow is little bit longer,
and square is correspondingly larger (Figure right).

Another example, look at case
where glass just thick enough that stopwatch makes extra half turn
as it times photon bouncing off back surface.

This time, back reflection arrow ends up pointing exactly in same
direction as front reflection arrow.

When combine 2 arrows, get final arrow whose length is 0.4,
and whose square is 0.16, representing probability of 16% (Figure right).

If increase thickness just enough

so that stopwatch hand timing back surface path makes extra full turn,
2 arrows end up pointing in opposite directions again, and final arrow will be 0 (Figure below).

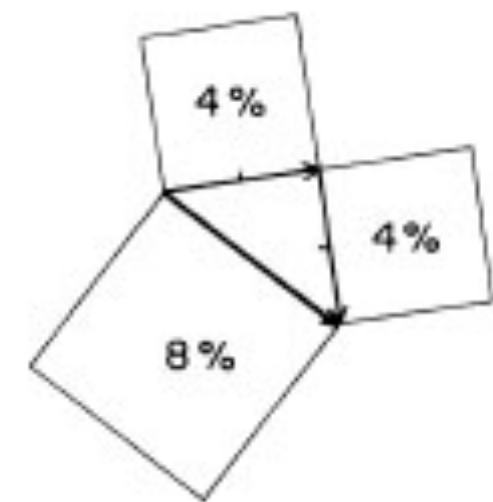
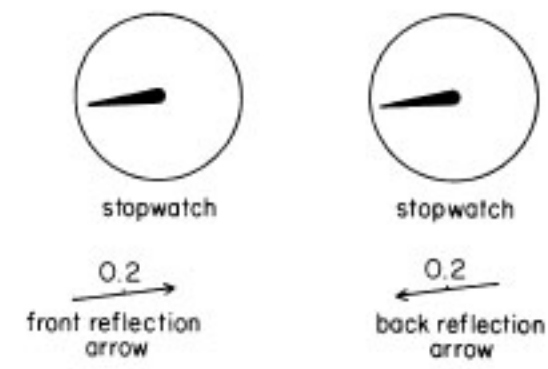
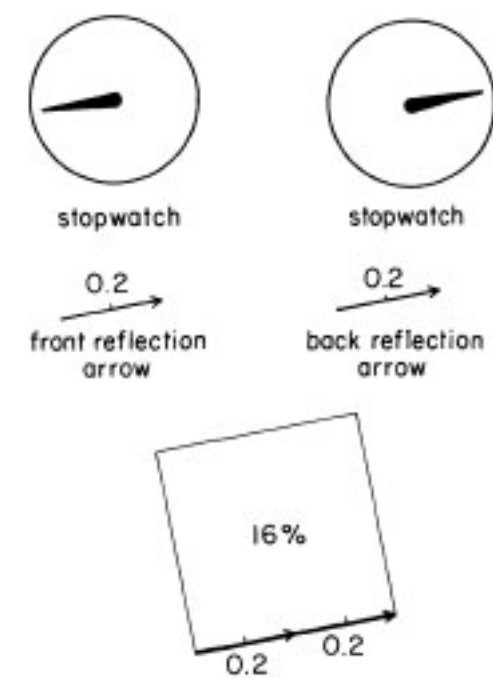
Situation occurs over and over,

when thickness of glass just enough to let stopwatch hand timing
back surface reflection make another full turn.

If thickness of glass just enough to let stopwatch hand timing back reflection
make an extra 1/4 or 3/4 turn, —> 2 arrows will end up at right angles.

Final arrow = hypotenuse of right triangle, and according to Pythagoras,
square on hypotenuse = sum of squares of other two sides.

Here is value that is right “twice a day” : 4% + 4% makes 8% (Figure right).



Notice as gradually increase thickness of glass,

front reflection arrow always points in same direction,

whereas back reflection arrow gradually changes its direction.

Change in relative direction of 2 arrows

makes final arrow go through repeating cycle of length 0 to 0.4;

thus square of final arrow goes through repeating cycle of 0 to 16% that observed in experiments (Figure below).

Just shown

how strange feature of partial reflection

can be accurately calculated

by drawing some little arrows on piece of paper.

Technical word for arrows is “probability amplitudes”,

and feel more dignified

when we say we are

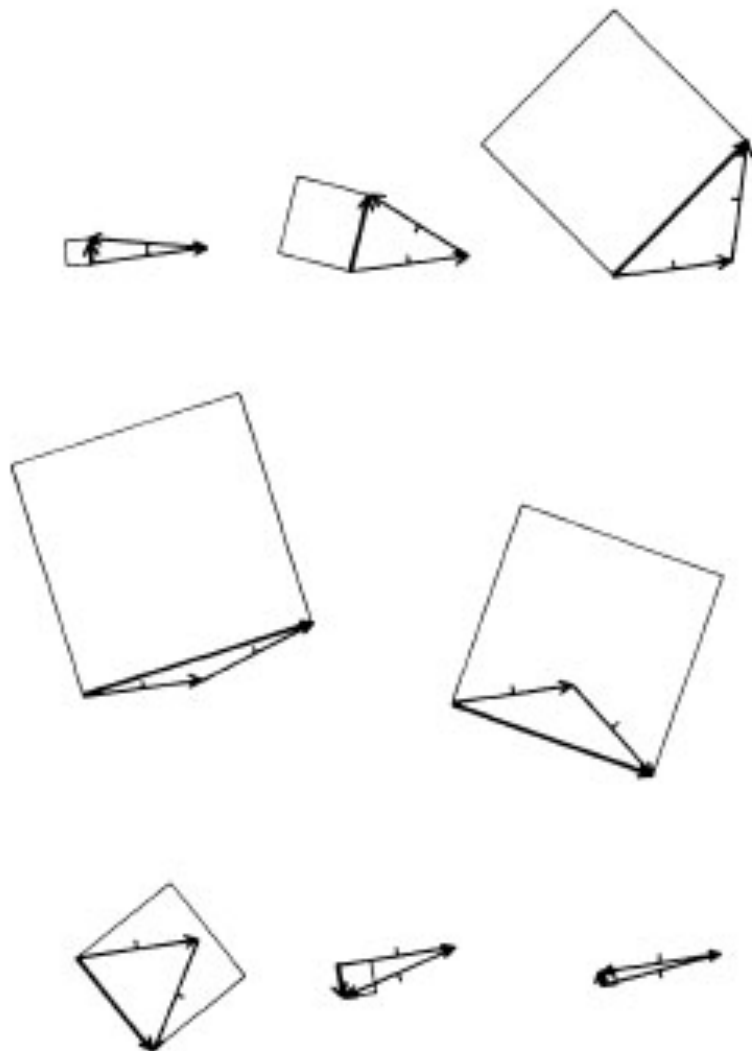
“computing the probability amplitude for an event”.

Prefer, though, to be more honest,

and say that we are trying to find arrow

whose square represents

the probability of something happening.



Before discussing photons,

let us discuss the colors you see on soap bubbles.

Or better, let a car leak oil in mud puddle.

When look at brownish oil in the dirty mud puddle,
you see beautiful colors on surface.

A thin film of oil floating on mud puddle acts something like very thin sheet of glass

it reflects light of one color from 0 to maximum, depending on oil thickness.

If shine pure red light on film of oil,

see splotches of red light separated by narrow bands of black

(where there was no reflection) because oil film's thickness not exactly uniform.

If shine pure blue light on oil film,

see splotches(different) of blue light separated by narrow bands of black.

If shine both red and blue light onto oil,

see areas that have just right thickness to strongly reflect only red light,

other areas of right thickness to reflect only blue light;

still other areas have thickness that strongly reflect both red and blue light (see as violet),

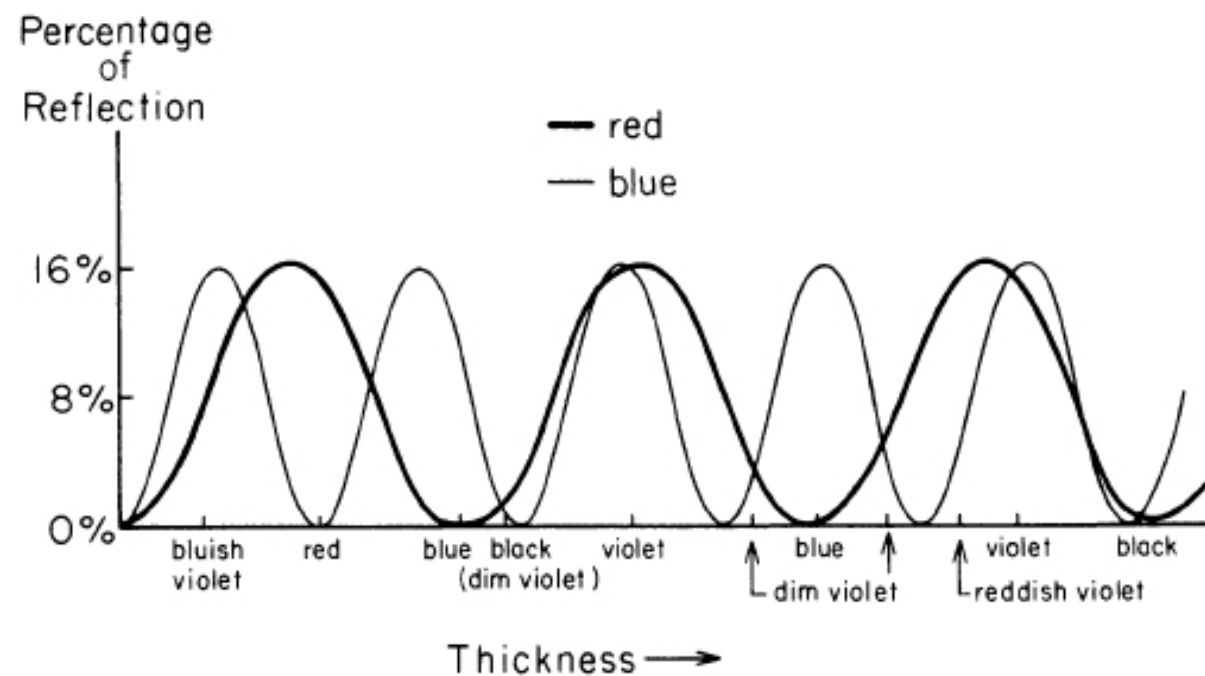
while other areas have exact thickness to cancel out all reflection,

and appear black.

To understand all this better,

need to know that cycle of 0 to 16% partial reflection by 2 surfaces repeats more quickly for blue light than for red light.

Thus at certain thicknesses, one or other or both colors are strongly reflected, while at other thicknesses, reflection of both colors is cancelled out (Figure below).



Cycles of reflection repeat at different rates
because stopwatch hand turns faster
when it times a blue photon
then does when timing a red photon.

In fact,

that's only difference between red photon and blue photon
(or photon of any other color, including radio wave, X-rays, and so on)

i.e., the speed of stopwatch hand(the frequency).

When shine red and blue light on film of oil,

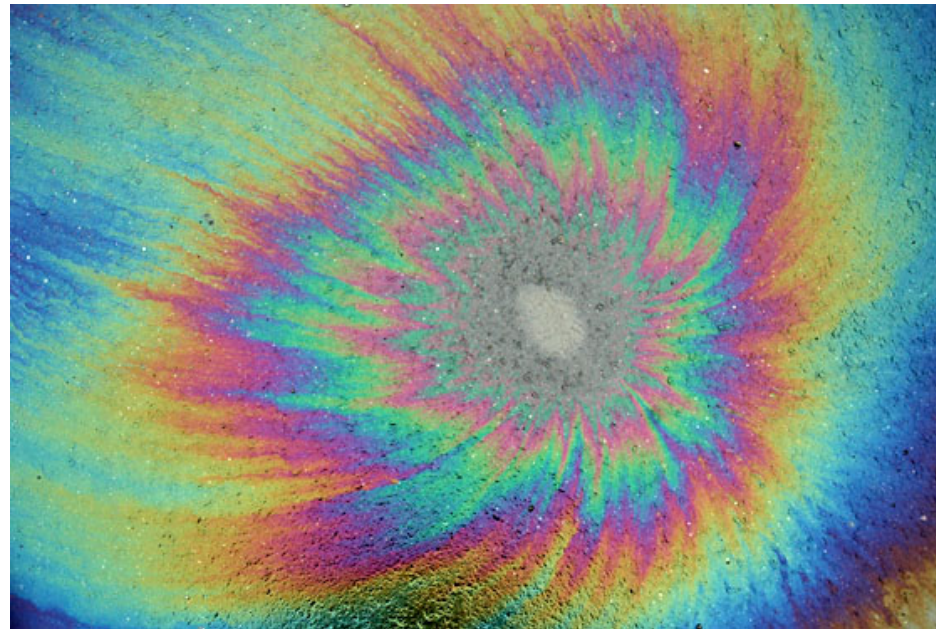
patterns of red and blue, and violet appear, separated by borders of black.

When sunlight,

which contains red, yellow, green, and blue light,

shines on film of oil, areas that strongly reflect each color overlap

and produce all kinds of combinations which our eyes see as different colors (see image below).



As oil film on puddle spreads out and moves over surface of water,

changing its thickness in various locations,

patterns of color constantly change.

(If, on other hand, were to look at same mud puddle at night

with one of those sodium streetlights shining on it,

would see only yellowish bands separated by black -

because those particular streetlights emit light of only one frequency or color.)

Phenomenon of colors produced by partial reflection

of white light by 2 surfaces called iridescence - it is found in many places.

Perhaps you have wondered how brilliant colors of hummingbirds and peacocks are produced.

Now you know.

How those brilliant colors evolved is also an interesting question.

When you admire a male peacock,

you should give credit to generations of lackluster females for being selective about their mates.

(humans got into act later and streamlined the selection process in peacocks.)

In coming discussion

we will show how this seemingly absurd process of combining little arrows

computes the correct answer for those other phenomena we are familiar with, i.e.,

light travels in straight lines;

reflects off mirror at same angle that came in

(“the angle of incidence is equal to the angle of reflection”);

lens focuses light, and so on.

The new framework will describe everything you know about light. **That is what a theory does!!**

Photons: Particles of Light

We are talking about light.

The 1st important feature about light —> appears to be particles

When very weak monochromatic light (light of one color) hits a detector,
the detector makes equally(same) loud clicks less and less often as the light gets dimmer.
The other important feature about light is partial reflection of monochromatic light.
Average of 4% of photons hitting single surface of glass reflected.
Already contains a deep mystery,
since it is impossible to predict which photons will bounce back and which will go through.

With 2nd surface,
results were strange:
instead of expected reflection of 8% by 2 surfaces,
partial reflection can be amplified as high as 16% or turned off,
depending on thickness of the glass.

Strange phenomenon of partial reflection by 2 surfaces
can be explained for intense light by the theory of waves and interference,
but wave theory **cannot** explain how detector makes equally loud clicks as light gets dimmer.

Quantum electrodynamics “resolves” this so-called “wave-particle duality” by saying
light is made of particles (as Newton originally thought),
but the price of this greatest achievement of science
is the retreat by physics to a position of being able to only calculate the probability
that a photon will hit detector, without offering good model of how it actually happens.

We have described

how physicists calculate the probability that particular event will happen.

We simply draw some arrows on piece of paper according to some rules,

which are summarized as follows:

GRAND PRINCIPLE:

Probability of event = square of length of arrow

called “probability amplitude”.

Arrow of length 0.4 $\rightarrow$ probability of 0.16, or 16%.

GENERAL RULE

for drawing arrows if event can happen in alternative ways:

Draw arrow for each way, then combine arrows (“add” them)

by hooking head of one to tail of next.

A “final arrow” then drawn from tail of first to head of last.

Final arrow is one whose square gives probability of entire event.

We also found some specific rules for drawing arrows in case of partial reflection by glass .

Now we show how this new model of the world,
that is so utterly different from anything you've ever seen before,
can explain all simple properties of light that you know, i.e.,
when light reflects off mirror,
angle of incidence is equal to angle of reflection;
light bends when goes from air into water;
light goes in straight lines;
light can be focused by lens, and so on.

Theory also describes many other properties of light that you are probably not familiar with.
Also many other phenomena that will not discuss.

However, we can guarantee every phenomenon about light that has been observed in detail can be explained by this theory of quantum electrodynamics, even though going to describe only simplest and most common phenomena.

Let us start with a mirror, and problem of determining how light reflected from it
(Figure next page).

At S have source that emits light of one color
at very low intensity (use red light again).

Source emits one photon at a time.

At P, place photomultiplier to detect photons.

Let's put at same height as source (not like (b))

drawing arrows easier if everything symmetrical.

Want to calculate chance

that detector will make click

after photon has been emitted by source.

Since possible that photon could go straight across to detector, place screen at Q to prevent that.

Now,

would expect that all light that reaches detector

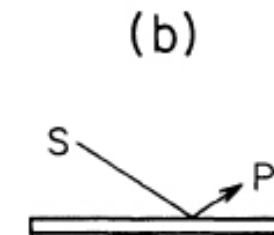
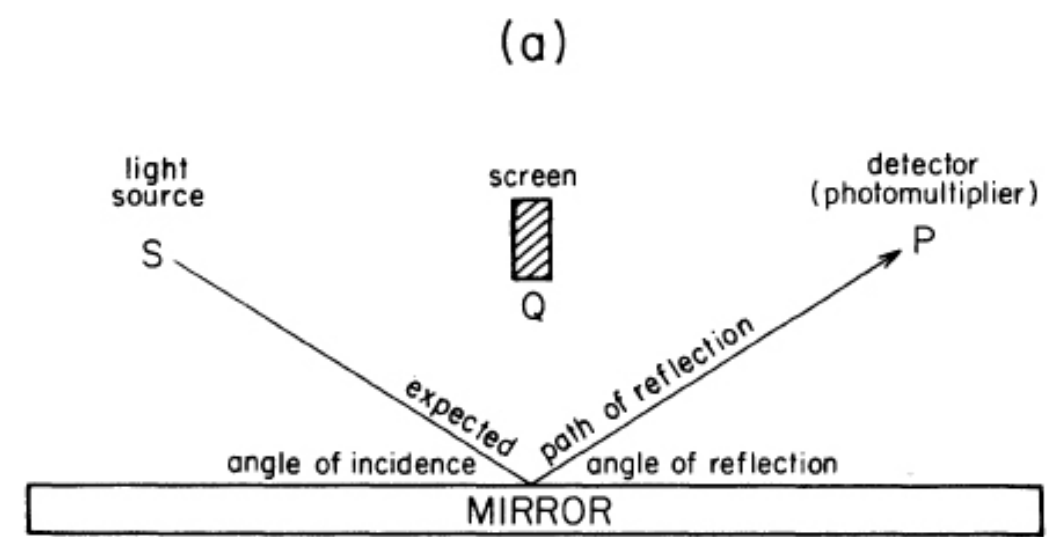
reflects off middle of mirror,

because that's place where angle of incidence equals angle of reflection.

And seems fairly obvious

that parts of mirror near two ends

do not have as much to do with reflection as with price of cheese, correct?



Although might think that parts of mirror near two ends

have nothing to do with reflection of light that goes from source to detector -

let us see what does quantum theory say.

Rule: probability that particular event occurs

is square of final arrow that is found

by drawing arrow for each way event **could** happen,

and then combining (“adding”) arrows.

In experiment measuring partial reflection of light by two surfaces,

there were just 2 ways photon could get from source to detector.

In this experiment,

there are “millions” of ways photon could go:

could go down to left-hand part of mirror at A or B (for example) and bounce up to detector (Figure right);

it could bounce off part where you think it should, at G;

or, could go down to right-hand part of mirror at K or M

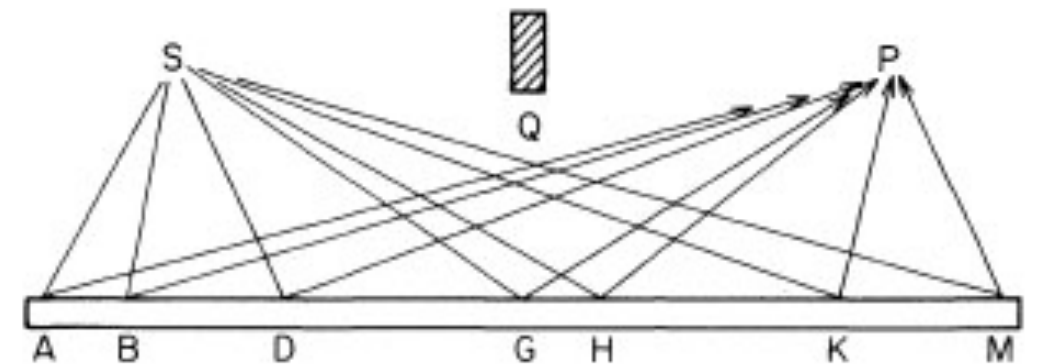
and bounce up to detector.

You might think that I’m crazy, because for most of ways I said the photon could reflect off mirror,

the angles aren’t equal.

But I’m not crazy, because that’s way light really goes!

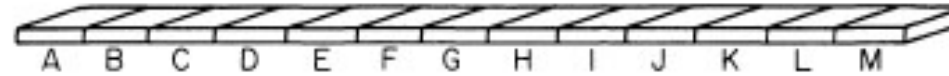
You ask — How can that be?



To make problem easier to understand,

suppose that mirror consists of only long strip from left to right

forget for moment, that mirror also sticks out from paper (Figure below).



While there are millions of places where light could reflect from strip of mirror,

we make an approximation by temporarily dividing mirror into definite number of little squares,

and consider only one path for each square -

calculation gets more accurate (but harder to do) as make squares smaller and consider more paths.

Now draw little arrow for each way light could go in this situation.

Each little arrow has certain length and certain direction.

Consider length first.

Might think that arrow you draw to represent path that goes to middle of mirror, at G,

is by far longest

(since seems to be very high probability that any photon that gets to detector must go that way),

and arrows for paths at ends of mirror must be very short.

No, No, No

should not make such an arbitrary rule - don't guess - investigate!

Eventually, find correct rule -

What actually happens - much simpler:

photon that reaches detector has **nearly equal chance** of going on any path,
so all little arrows have **nearly same length**.

(are some slight variations in length due to various angles and distances involved,
but so minor that we can ignore them.)

So say that each little arrow draw will have arbitrary standard length

make the standard length very short

because there are so many of these arrows representing many ways light could go (Figure below).

Although safe to assume that length of all arrows nearly the same,

directions clearly differ because timing **different**

Remember, direction of particular arrow

determined by final position of imaginary stopwatch

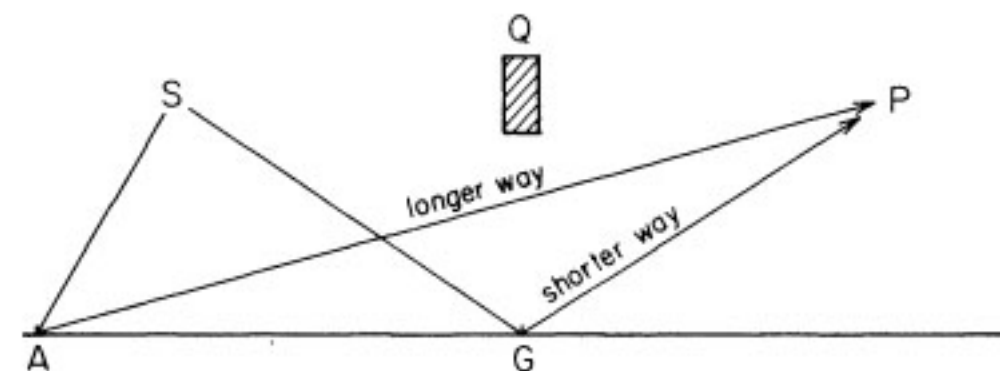
that times photon as moves along particular path.

When photon goes way off to left end of mirror, at A,

and then up to detector,

clearly **takes more time**

than photon that gets to detector by reflecting in middle
of mirror, at G (Figure right).



Or, imagine for moment that you were in hurry

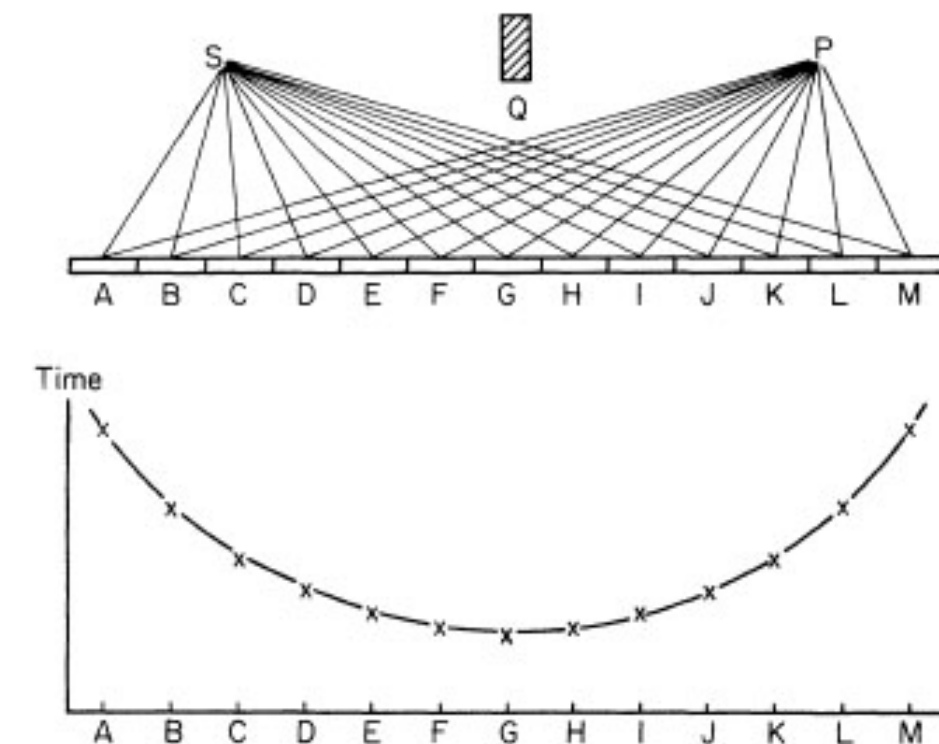
and had to run from source over to mirror and then to detector.

Know that certainly isn't good idea to go way over to A and then all way up to detector;

would be much faster to touch mirror somewhere in middle.

To help us calculate direction of each arrow,

going to draw graph right underneath sketch of mirror (Figure below).



Directly below each place on mirror where light could reflect, going to show, vertically how much time would take if light went that way.

More time takes, higher point will be on graph.

Starting at left, time takes photon to go on path

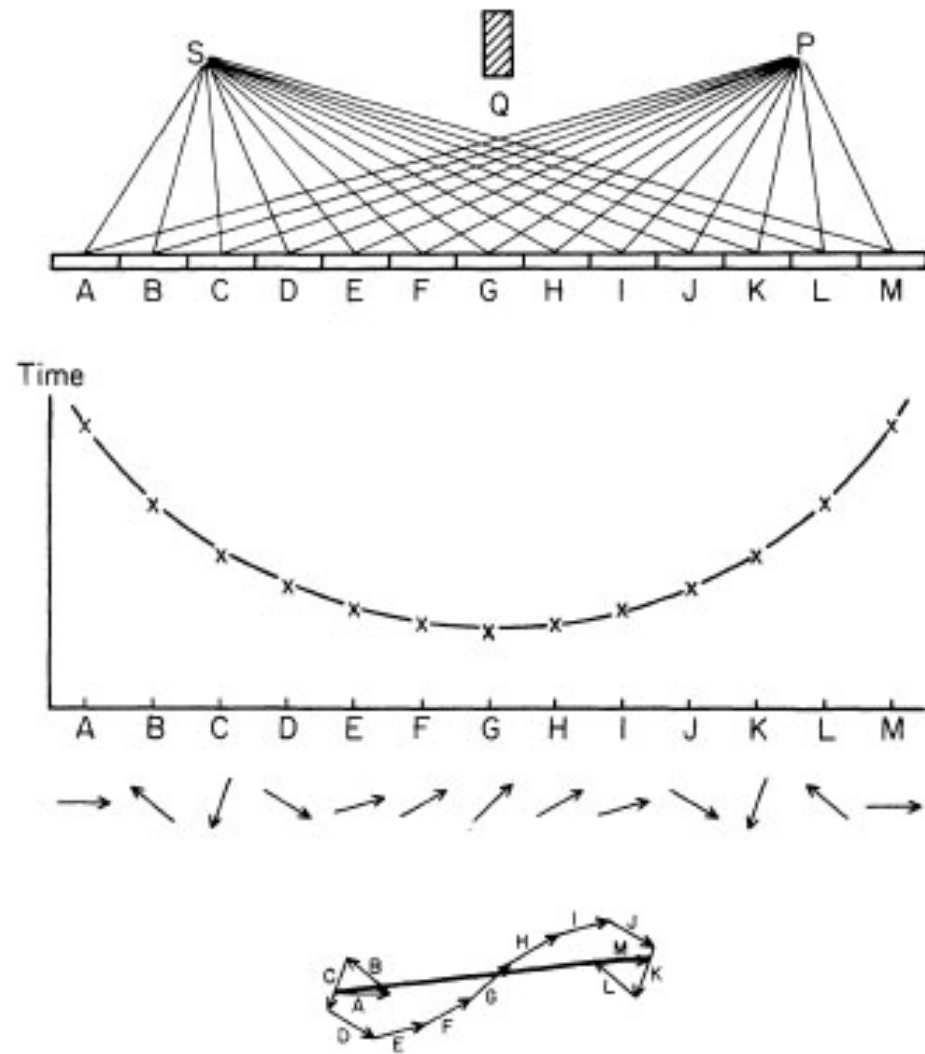
that reflects at A pretty long,

so plot point high up on graph.

As move toward center of mirror,

time takes for photon to go particular way goes down,

so plot each successive point lower than previous one.



After pass center of mirror, time takes photon to go on each successive path gets longer and longer,

so plot points correspondingly higher and higher.

To aid your visualization, we connect points:

they form a symmetrical curve that

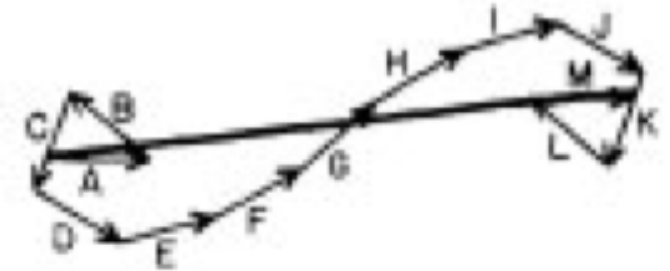
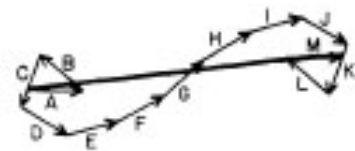
starts high, goes down, and then goes high again.

Now, what does that mean for direction of little arrows?

Direction of particular arrow corresponds

to amount of time would take photon to get

from source to detector following that particular path.



Draw arrow, starting at left.

Path A takes most time; its arrow points in some direction (Figure right).

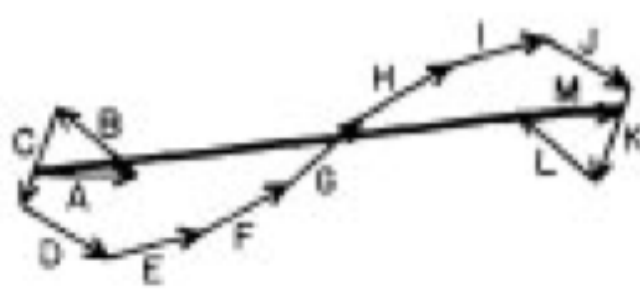
Arrow for path B points in different direction because time different.

At middle of mirror, arrows F, G, and H point in nearly same direction because times nearly the same.

After passing center of mirror, see that each path on right side of mirror corresponds to path on left side of mirror whose time is exactly same

(consequence of putting source and detector at same height, and path G exactly in middle).

Thus arrow for path J, for example has same direction as arrow for path D.



Now let's add little arrows (Figure above).

Starting with arrow A, hook arrows to each other, head to tail.

Now if we were to take walk using each little arrow as step, wouldn't get very far at beginning, because direction from one step to next so different (changes quickly).

But after while arrows begin to point in generally same direction, and make some progress.

Finally, near end of walk, direction from one step to next again quite different, so stagger about some more.

All we have to do now is draw a final arrow.

Simply connect tail of 1st arrow to head of last one, and see how much progress made on walk (Figure above).

And behold - get sizable final arrow!

Theory of quantum electrodynamics predicts that light does, indeed reflect off mirror! Wow!

Now investigate. What determines how long final arrow is?

Notice a number of things.

1st - ends of mirror not important: since there, little arrows wander around and don't get anywhere.

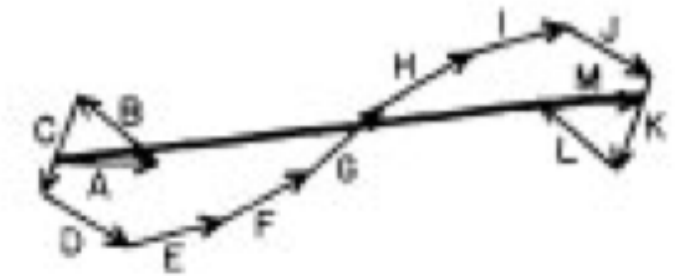
If chopped off ends of mirror - parts that instinctively knew was wasting time fiddling around with - would hardly affect length of final arrow.

So where is part of mirror that gives final arrow substantial length?

It's part where arrows all pointing in nearly same direction - because time almost same.

If look at graph showing time for each path (Figure right),

see that time nearly same from one path to next at bottom of curve,
where **time is least**.



To summarize, where time **least** is also where time for nearby paths nearly **same**;

that's where little arrows point in nearly same direction and add up to substantial length;

that's where probability of photon reflecting off mirror determined.

And that's why, in this approximation,

we **can get away** with crude picture of world that says that light **only** goes where time is least
and easy to prove that where time least, i.e., **angle of incidence = angle of reflection**.

So theory of quantum electrodynamics gives right answer

middle of mirror is important part for reflection -

but correct result came out at expense of believing that light reflects all over mirror,

and having to add bunch of little arrows together whose sole purpose was to cancel out.

All that might seem to you to be waste of time - some silly game for mathematicians only.

After all, doesn't seem like "real physics" to have something there that only cancels out!

Let's test idea that reflection **really** is going on all over mirror by doing another experiment.

1st, chop off most of mirror, and leave about 1/4, over on left.

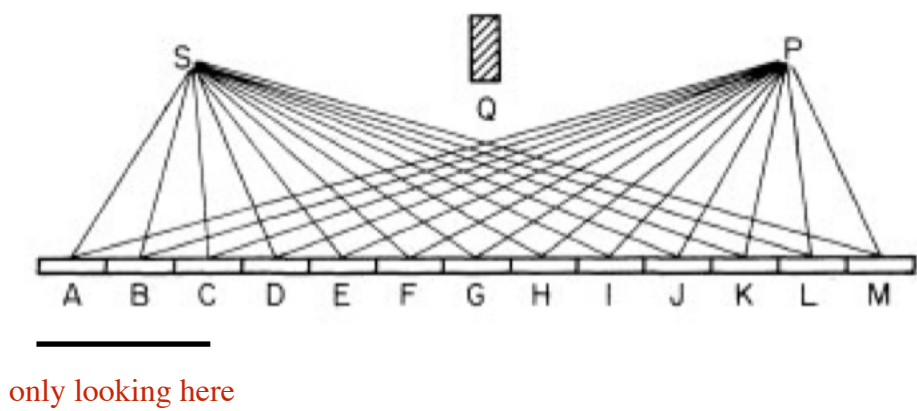
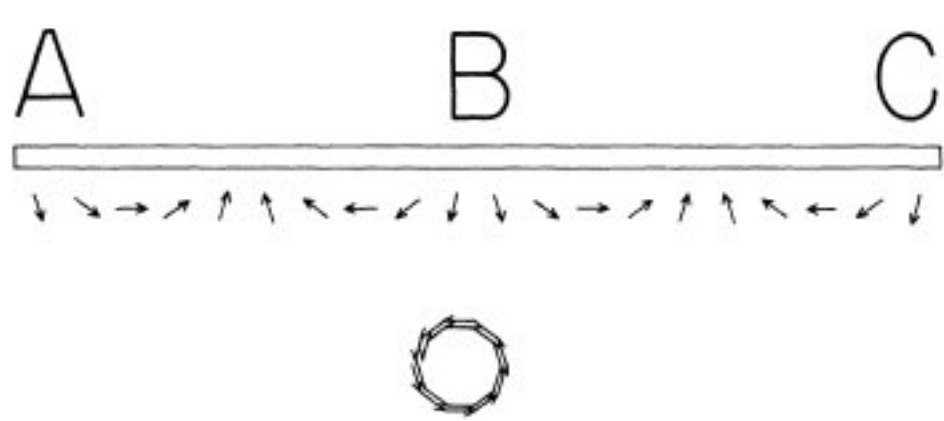
Still have a pretty big piece of mirror, but in wrong place.

In previous experiment arrows on left side of mirror were pointing in directions very different from one another because of large difference in time between neighboring paths.

In this experiment going to make more detailed calculation

by taking intervals on that left-hand part of mirror that much closer together

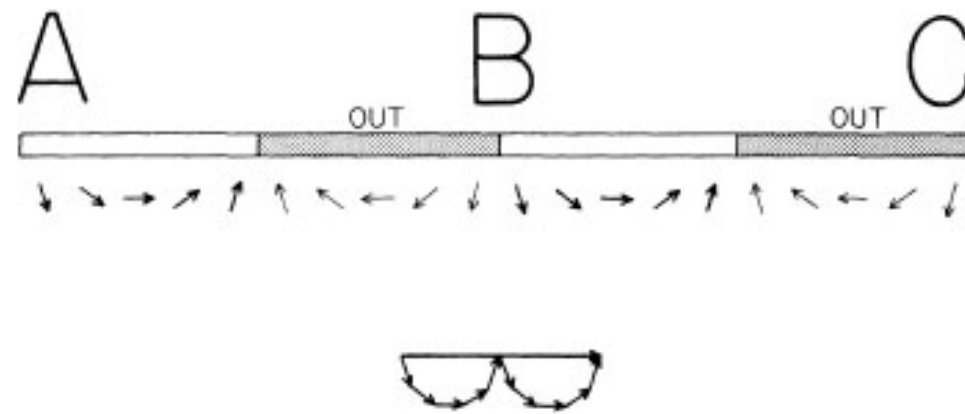
fine enough that there is not much difference in time between neighboring paths (Figure below).



With more detailed picture, see that some of arrows point more or less to right;
others point more or less to left.

If add all arrows together, have bunch of arrows going around in what is essentially a circle,
getting nowhere.

But suppose carefully **scrape** mirror away in areas whose arrows have bias in one direction - (left) -
-> only those places whose arrows point generally other way remain (Figure below)



When add up only arrows that point more or less to right,
get series of dips and substantial final arrow - according to theory,
should now have strong reflection!!

And indeed experiment shows that is what happens —> theory is correct!

Such a “mirror” is called a “diffraction grating”, and it works like charm.

Isn't it wonderful - can take piece of mirror where you didn't expect any reflection,
scrape away part of it, and it reflects!

Particular grating we just showed was tailor-made for red light.

Wouldn't work for blue light;

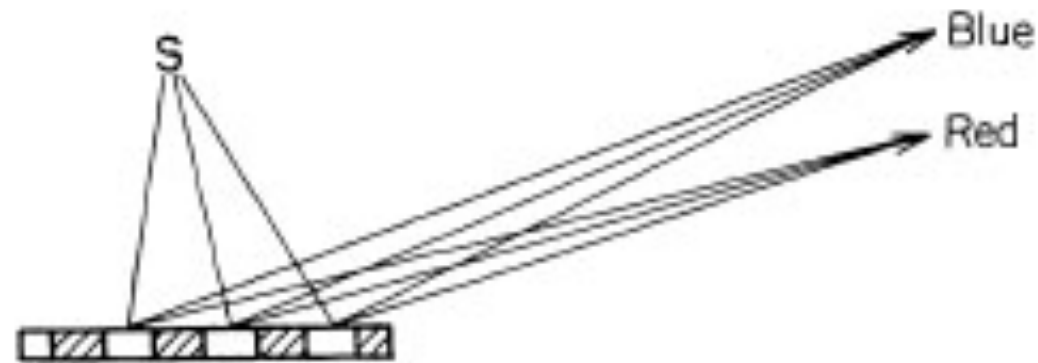
would have to make new grating with cut-away strips spaced closer together because,
as said earlier, stopwatch hand turns around faster when it times blue photon
compared to red photon.

So cuts especially designed for “red” rate of turning don't fall in right places for blue light;
arrows get kinked up and grating doesn't work very well.

But as matter of fact,

it happens that if we move photomultiplier down to somewhat different angle,
grating made for red light now works for blue light.

Just lucky accident, consequence of geometry involved (Figure below).



If shine white light down on grating, red light comes out at one place, orange light comes out slightly above it, followed by yellow, green, and blue light - all colors of rainbow.

Where there is series of grooves close together, can often see colors - for example, when hold LP phonograph record (or better, CD) - under bright light at correct angles.

Perhaps have seen those wonderful silvery signs on cars:

when car moves, see very bright colors changing from red to blue.

Now know where colors come from:

looking at grating - mirror that's been scratched in just right places.

Sun is light source, and eyes are detector.

So grating shows

that can't ignore parts of mirror that don't seem to be reflecting;

If do some clever things to mirror,

can demonstrate reality of reflections from all parts of mirror
and produce striking optical phenomena.

More importantly,

demonstrating reality of reflection from all parts of mirror
shows that there must be amplitude - an arrow -
for **every way** an event can happen.

And in order to calculate correctly probability of event in different circumstances,
have to add arrows for **every way** that event could happen
not just ways we think are important ones!

Now, talk about something more familiar than gratings.

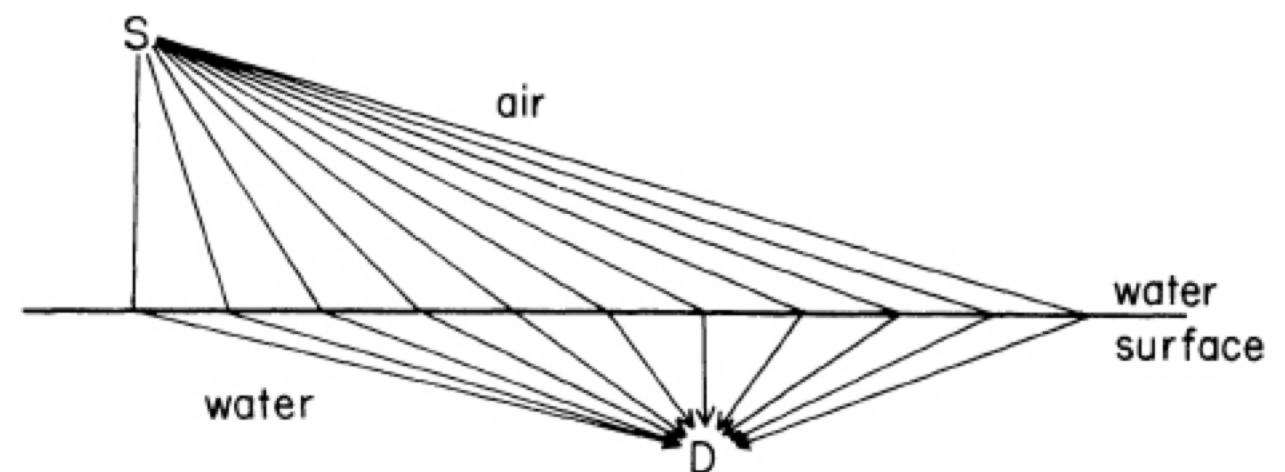
Let us talk about light going from air into water.

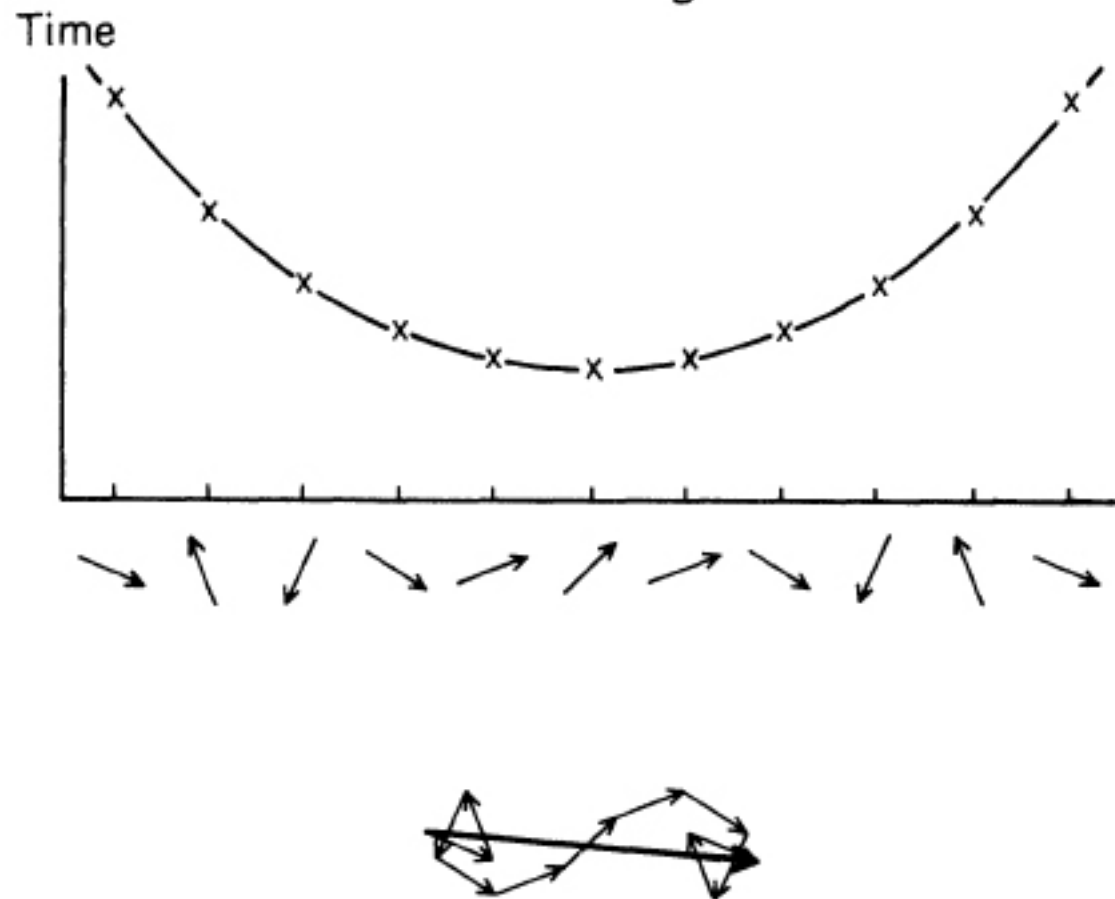
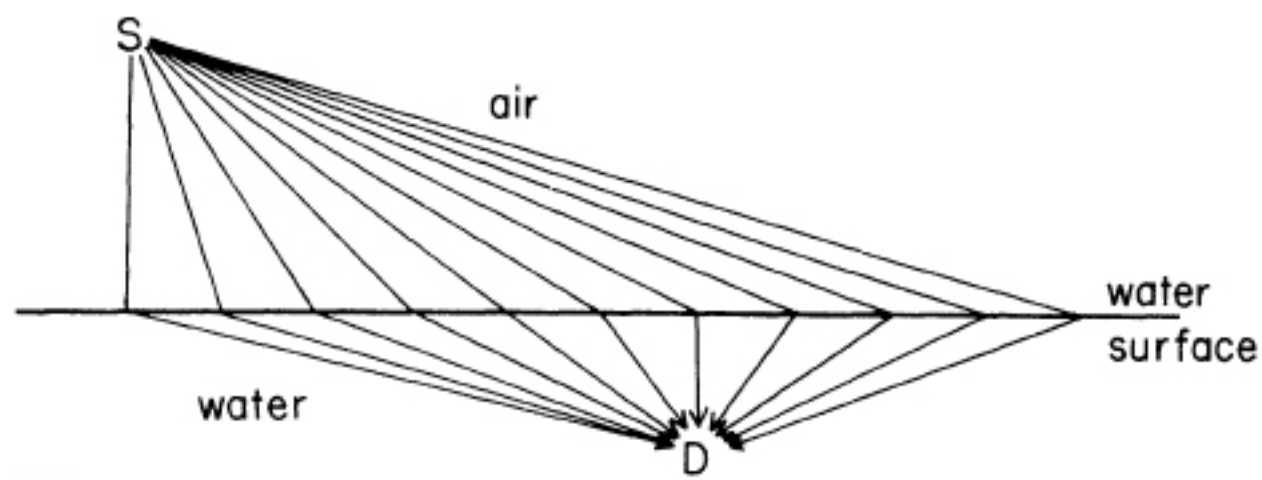
This time, we put the photomultiplier underwater.

Suppose experimenter can arrange that!

Source of light is in air at S,

and detector is underwater, at D (Figure right).





Once again, want to calculate probability that photon will get from light source to detector.

To make calculation, should consider **all** ways light could go.

Each way light could go contributes little arrow and, as in previous example,

all little arrows have nearly same length.

Can again make graph of time takes photon to go on each possible path.

Graph will be curve very similar to one made for light reflecting off mirror: starts up high, goes down, and then back up again;

most important contributions come from places where arrows point in nearly same direction (where time nearly same from one path to next), which is at bottom of curve.

That is also a place where time is least, so all have to do is find out where time least.

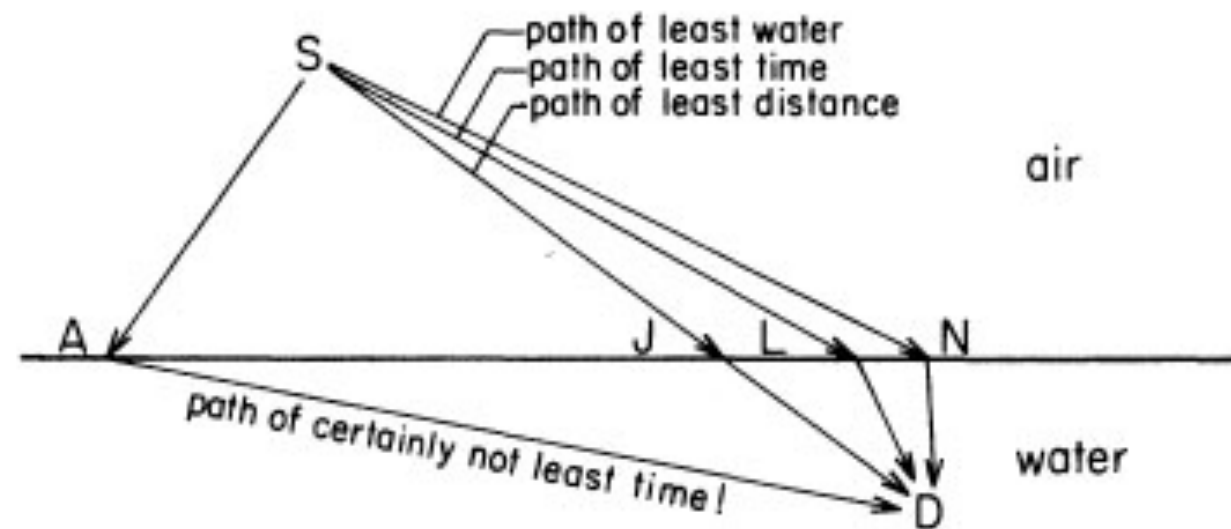
Turns out that light seems to go slower in water than does in air (explain why later),

which makes distance through water more “costly”,

so to speak, then distance through air.

Not hard to figure out which path takes least time:

suppose you're lifeguard, sitting at S, and someone drowning, at D (Figure right).



Can run on land faster than can swim in water.

Problem is, where do we enter water in order to reach drowning victim fastest?

Run down to water at A, and swim like hell?

Of course not.

But running directly toward victim and entering water at J not fastest route, either.

While would be foolish for lifeguard to analyze and calculate under circumstances,

there is a computable position at which time minimum:

it's a compromise between taking direct path, though J,

and taking path with least water, though N.

And so it is with light - path of least time enters water between J and N, such as L.

Light “bends” at interface; why underwater objects seem displaced!

This page is extra for you to read

Another phenomenon of light that would like to mention briefly is mirage.

When driving along road that is very hot, can sometimes see what looks like water on road.

What you really are seeing is the sky.

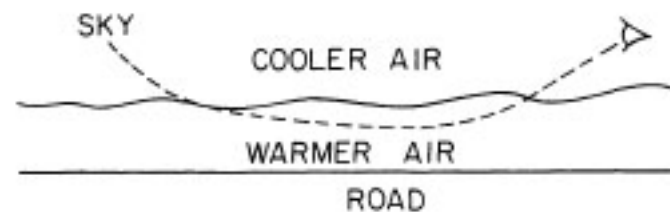
We normally see sky on road, because road has puddles of water on it

that is just partial reflection of light by a single surface.

But how can you see sky on road when no water there?

What we need to know is that light goes slower through cooler air than through warmer air,
and for mirage to be seen,

observer must be in cooler air that is above hot air next to road surface (Figure below).



How it is possible to look down and see sky

can be understood by finding path of least time.

I'll let you play with that one at home - it's fun to think about, and pretty easy to figure out.

In examples so far

showed light reflecting off mirror and light going through air and then water,
making the approximation: for simplicity,
we drew various ways that light could go as double straight lines -
two straight lines that form an angle.

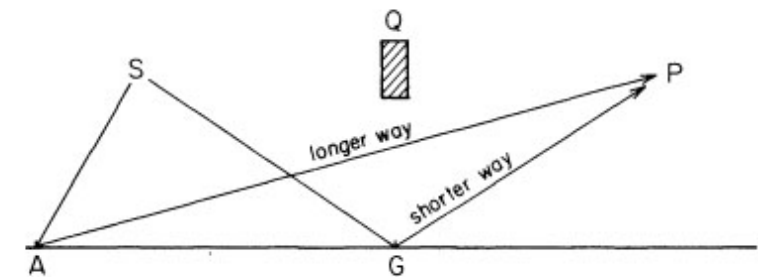
But don't have to assume that light goes in straight lines

when it is in uniform material like air or water;

even that is explainable

by the general principles of quantum theory:

probability of event found by adding arrows for all ways event could happen.

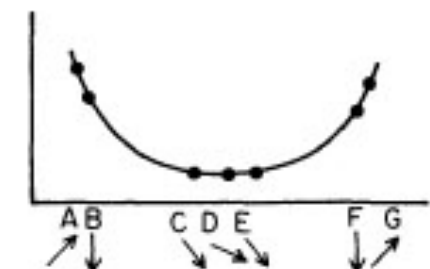
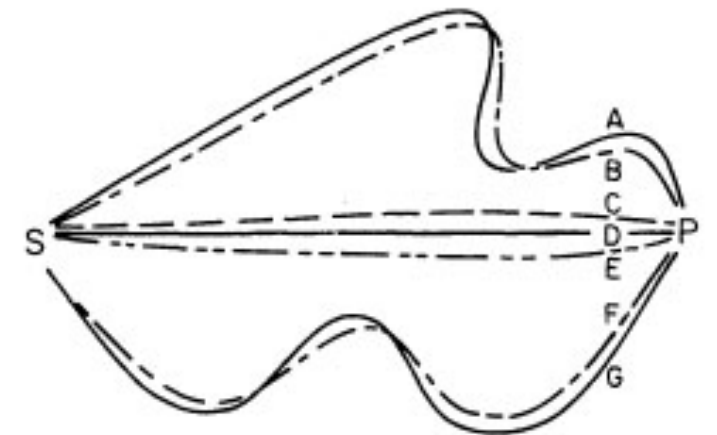


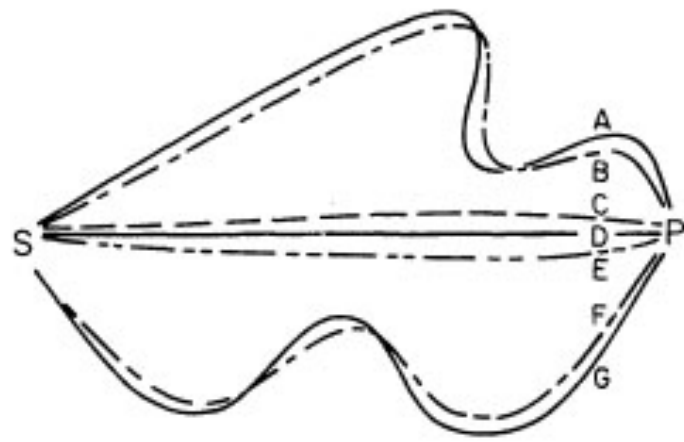
Next example, let us see how

We show how, by adding little arrows,
it can appear that light goes in a straight line.

Put source and photomultiplier at S and P, respectively (Figure right),
and look at all ways light could go - in all sorts of crooked paths -
to get from source to detector.

Then draw little arrow for each path, and we're now experts!

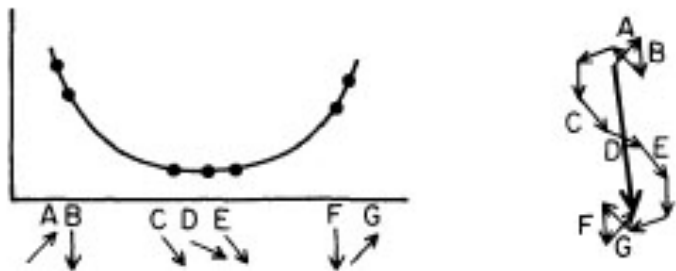




For each crooked path, such as path A,
there's nearby path that's little bit straighter
and distinctly shorter -
that is, takes **much less** time —> cancel out.

But where paths become nearly straight - at C, for example
nearby, straighter path has nearly same time.

That's where arrows add up rather cancel out;
that's where light goes.



Important to note that single arrow that represents straight-line path, through D (Figure above),
not enough to account for probability that light gets from source to detector.

Nearby, nearly straight paths - through C and E, for example - also make important contributions.

So light doesn't really travel only in straight line;

it “smells” or “samples” neighboring paths around it, and uses small core of nearby space.

In the same way,

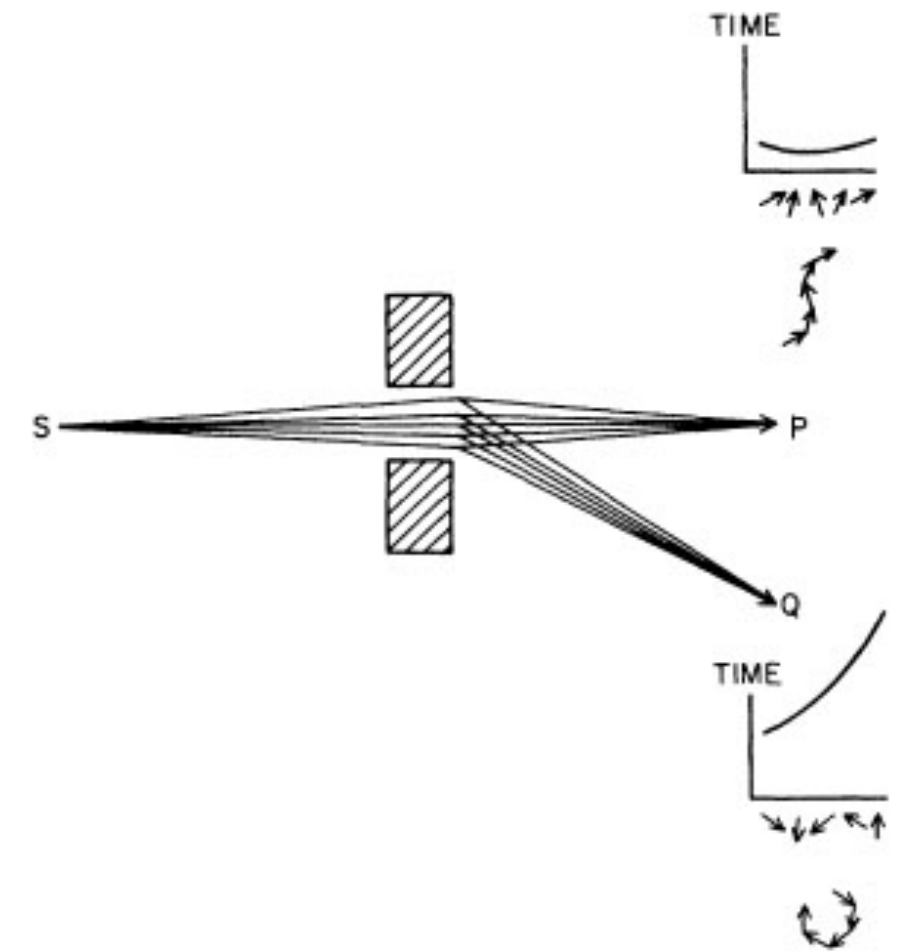
a mirror has to have enough size (be big enough) to reflect normally:

if mirror too small for the core (nearby) of neighboring paths,

light scatters in many directions, no matter where you put mirror.

Let us see how....

Let us investigate this central core of light more closely
by putting source at S, photomultiplier at P,
and pair of blocks between them
to keep paths of light from wandering too far away
(Figure right).



Now, put 2nd photo-multiplier at Q, below P,
and assume again, for simplicity,
that light can get from S to Q
only by paths of double straight lines.

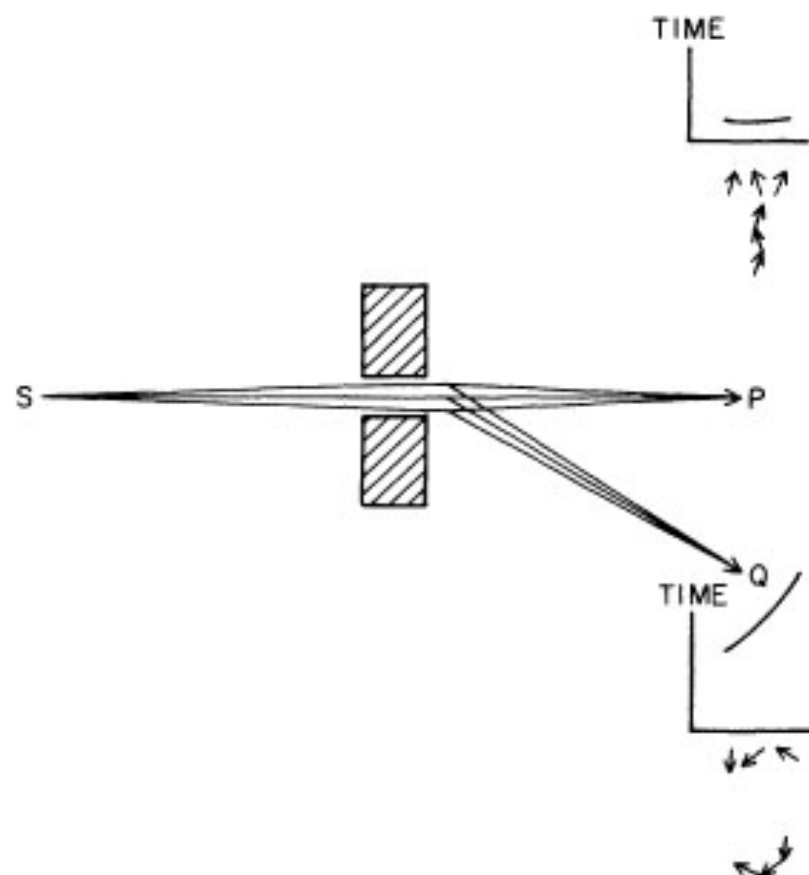
Now, what happens?

When gap between blocks wide enough
to allow many neighboring paths to P and to Q,
arrows for paths to P add up (because all paths to P take nearly same time),
while paths to Q cancel out (because those paths have sizable difference in time).

Thus photomultiplier at Q doesn't click.

But as push blocks closer together,
at certain point, detector at Q starts clicking!

When gap nearly closed and there are only few neighboring paths, arrows to Q also add up,
because there is hardly any difference in time between them, either.



Of course, both final arrows are small,
 so not much light either way through such small hole,
 but detector at Q clicks almost as much as one at P!

So when try to squeeze light too much
 to make sure it's going in only straight line,
 it refuses to cooperate and begins to spread out.

So idea that light goes in straight line is a convenient approximation
 to describe what happens in the world familiar to us;

Similar to crude approximation that says when light reflects off mirror,
 angle of incidence = angle of reflection.

Just as we were able to do a clever trick to make light reflect off mirror at many angles,
 can do a similar trick to get light to go from one point to another in many ways.

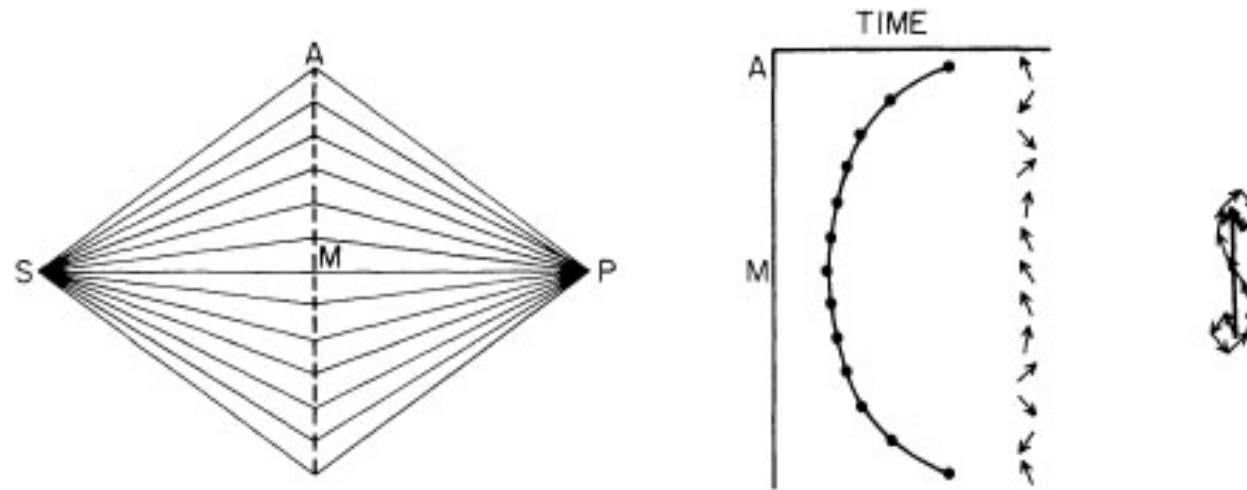
1st, simplify situation -

draw vertical dashed line (Figure below)

between light source and detector

(line means nothing; just artificial line)

and say that only paths we are going to look at are double straight lines.



Graph that shows

time for each path looks same as in case of mirror (drawn sideways, this time):

curve starts at A, at top, and then comes in,

because paths in middle are shorter and take less time.

Finally, curve goes back out again.

Now have some fun.

Let's "fool the light", so that all paths take exactly same amount of time.

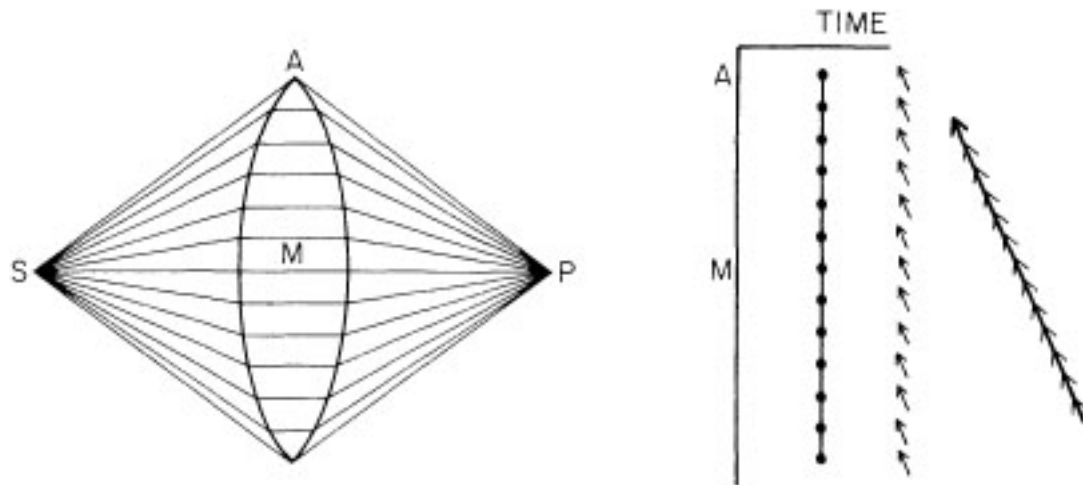
How can we do this?

How can we make shortest path, through M,
take exactly same time as longest path through A?

Well, light goes slower in water than does in air;
also goes slower in glass (much easier to handle!).

So, if put in just right thickness of glass on shortest path, though M,
can make time for that path exactly same as for path through A.

Paths next to M, which are just little longer, won't need quite as much glass (Figure below).



Nearer get to A,

less glass have to put in to slow up light.

By carefully calculating and putting

in just right thickness of glass to compensate
for time along each path,
can make all times same.

When draw arrows for each way light could go,

find that have succeed in straightening them all out -

and there are, in reality, millions of tiny arrows -

so net result is sensationally large, unexpectedly enormous final arrow!

Of course, you know what I am describing;

it's a focusing lens.

By arranging things so that all times equal, can focus light -

can make probability very high that light will arrive at particular point,
and very low that will arrive anywhere else.

I have used these examples to show how theory of quantum electrodynamics,
which looks at first like an absurd idea with no causality, no mechanism,
and nothing real to it,

produces effects that we are familiar with like:

light bouncing off mirror,

light bending when goes from air into water,

and light focused by lens.

It also produces other effects such as diffraction grating, etc.

In fact, theory continues to be successful at explaining **every** phenomenon of light.

Have shown with examples

how to calculate probability of event that can happen in alternative ways:

draw arrow for each way event can happen, and add arrows.

“Adding arrows” means arrows are placed head to tail and “final arrow” is drawn.

Square of resulting final arrow represents probability of event.

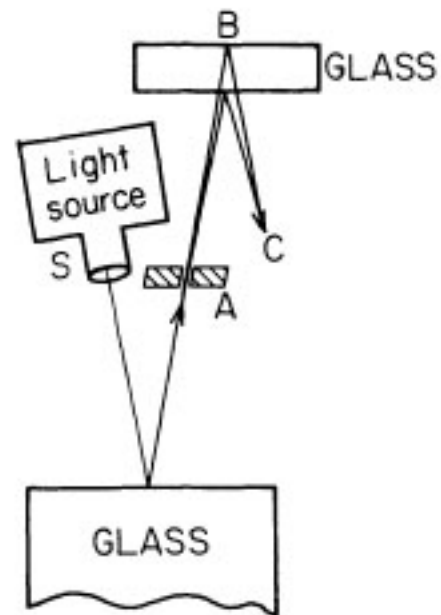
In order to give full picture of this quantum theory,

I am now going to show you how physicists calculate probability of compound events;
i.e., events that can be broken down into series of steps,
or events that consist of number of things happening **independently**.

Example of such a compound event can be demonstrated by modifying 1st experiment

Aimed some red photons at single surface of glass to measure partial reflection.

Instead of putting photomultiplier at A (Figure below),



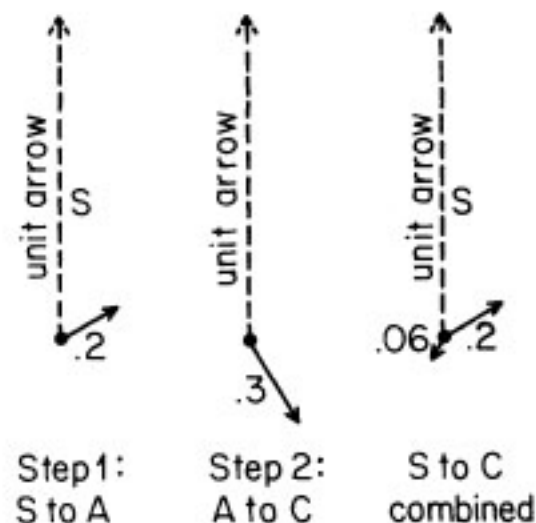
put in screen with hole in it to let photons that reach point A go through.

Then put sheet of glass at B,

and place photomultiplier at C.

How do we figure out probability

that photon will get from source S to detector C?



We can think of an event as **sequence of 2 steps**.

Step 1: photon goes from source to point A, reflecting off single surface of glass.

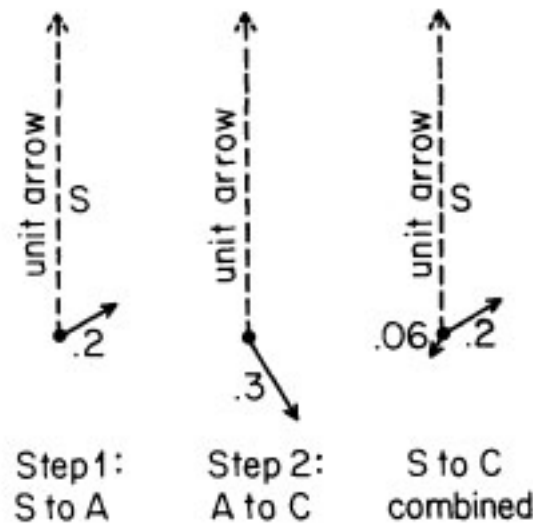
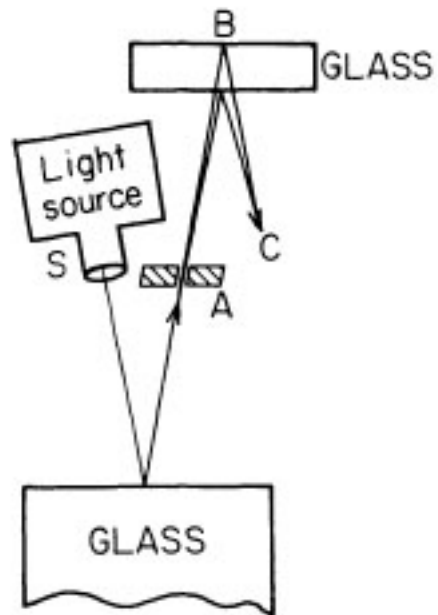
Step 2: photon goes from point A to photomultiplier at C, reflecting off sheet of glass at B.

Each step has final arrow - an “amplitude” (now I use words interchangeably) -

that can be calculated according to rules know so far.

Amplitude for 1st step has length of 0.2

(square is 0.04, probability of reflection by single surface of glass),
and turned at some angle - say, 2 o'clock (Figure left).



To calculate amplitude for 2nd step,
temporarily put light source at A
and aim photons at layer of glass above.

Draw arrows for front and back surface reflections
and add them -
say end up with final arrow with length of 0.3,
and turned toward 5 o'clock.

Now,

how do we combine 2 arrows to draw amplitude for entire event?

Look at each arrow in new way:

as instructions for shrink and turn.

In this example,

1st amplitude has length of 0.2 and turned toward 2 o'clock.

If begin with “unit arrow” - arrow of **length 1** pointed straight up

can shrink unit arrow from 1 down to 0.2, and turn from 12 o'clock to 2 o'clock.

For those of you that had the QM class:

$$A_1 = \text{amplitude for process 1} = a_1 e^{i\theta_1}$$

$$A_2 = \text{amplitude for process 2} = a_2 e^{i\theta_2}$$

a_i corresponds to the arrow length for process i

θ_i corresponds to the stopwatch time for process i

If two processes occur independently of each other

$$A = \text{amplitude for combined process} = A_1 A_2 = a_1 a_2 e^{i(\theta_1 + \theta_2)}$$

i.e., you multiply the lengths and add the times!

Amplitude for 2nd step can be thought of

as shrinking unit arrow from 1 to 0.3 and turning from 12 o'clock to 5 o'clock.

Now, to combine amplitudes for both steps,

shrink and turn in succession.

First, shrink unit arrow from 1 to 0.2 and turn from 12 o'clock to 2 o'clock;

then shrink arrow further, from 0.2 down to $\frac{3}{10}$ of that,

and turn by amount from 12 to 5 - that is, turn from 2 o'clock to 7 o'clock.

Resulting arrow has length of 0.06 and pointed toward 7 o'clock.

It represents probability of 0.06 squared, or $0.0036 = 0.36\%$.

Observing arrows carefully,

see that result of shrinking and turning 2 arrows in succession

is same as adding their angles (2 o'clock + 5 o'clock) and multiplying their lengths ($0.2 * 0.3$).

To understand why must add angles is easy:

angle of arrow determined by amount of turning by imaginary stopwatch hand.

So total amount of turning for 2 steps in succession

is simply sum of turning for 1st step plus additional turning for 2nd step.

Why we called process “multiplying arrows”

this takes bit more explanation, but it’s interesting.

Look at multiplication,

for moment, from point of view of Greeks (a digression).

Greeks wanted to use numbers that were not necessarily integers,

so represented numbers with lines.

Any number can be expressed as transformation of unit line

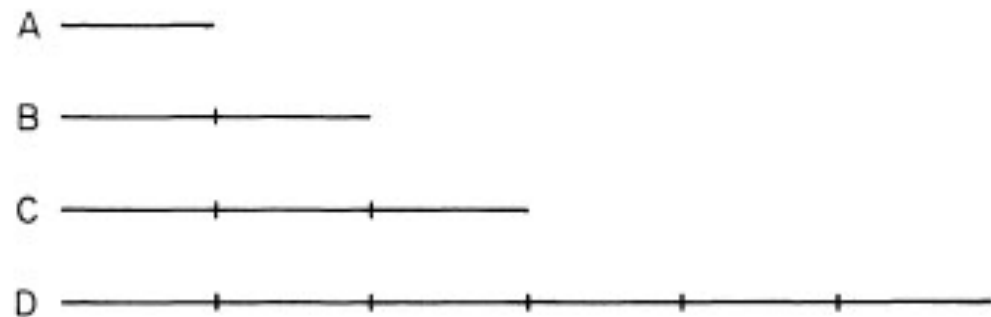
by expanding or shrinking it.

For example,

if line A is unit line (Figure below),

line B represents line 2

and line C represents 3.



Now how multiply 3 times 2?

Apply transformations in succession:

starting with line A as unit line,

expand it 2 times and then 3 times

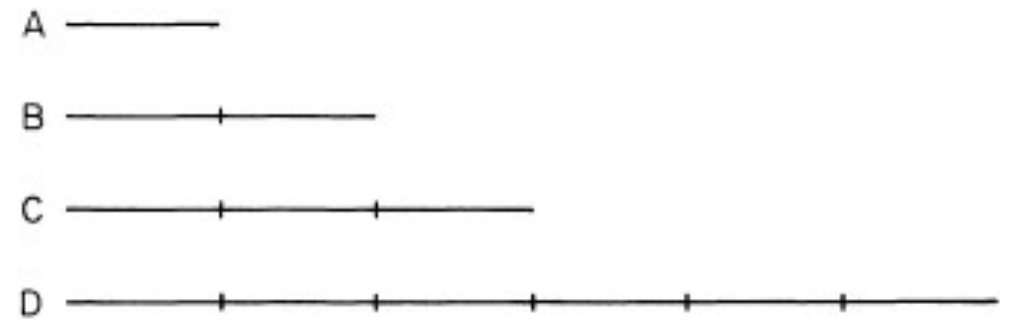
(or 3 times and then 2 times - order doesn’t make any difference).

Result is line D,

whose length represents 6.

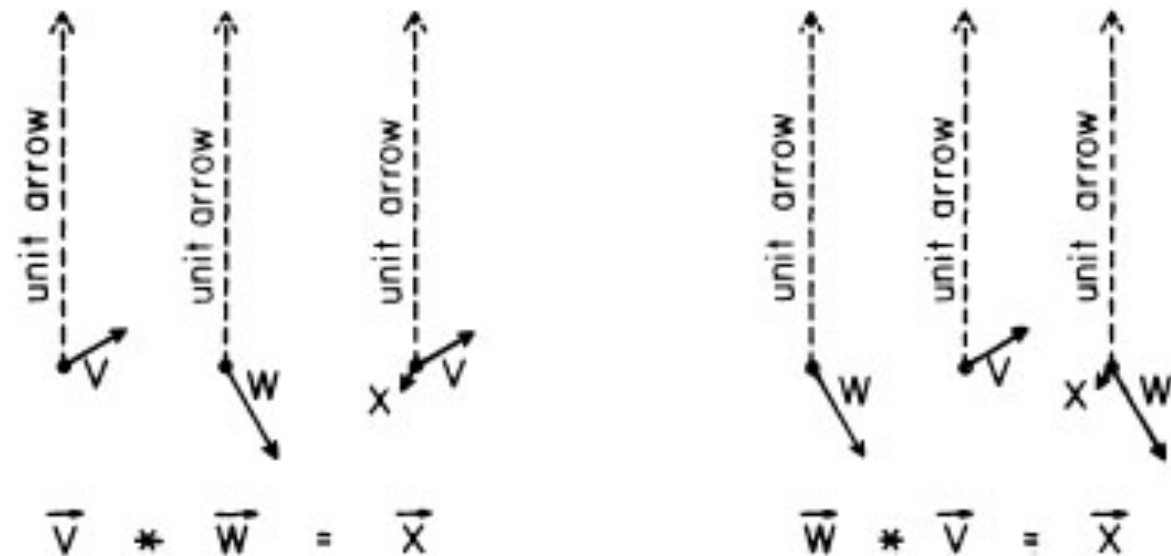
What about multiplying $1/3$ times $1/2$?

Taking line D to be unit line,
now, shrink it to $1/2$ (line C)
and then to $1/3$ of that.



Result is line A, which represents $1/6$. Multiplying arrows works same way (Figure below).

Apply transformations to unit arrow in succession
just happens that transformation of arrow
involves 2 operations, shrink and turn.



To multiply arrow V times arrow W,
shrink and turn unit arrow
by prescribed amounts for V,
and then shrink it and turn it

by amounts prescribed for W - again, order doesn't make any difference.

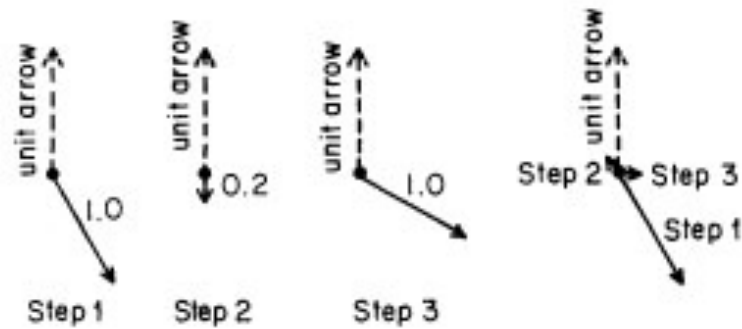
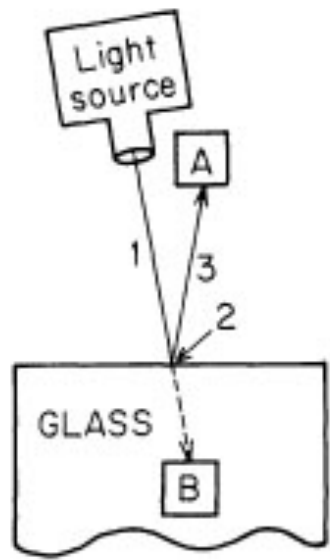
So multiplying arrows

follows same rule of successive transformation that work for regular numbers.

Go back to 1st experiment from earlier lecture

partial reflection by single surface -

with this idea of successive steps in mind (Figure below).



Can divide path of reflection into 3 steps:

- 1) light goes from source down to glass,
- 2) is reflected by glass, and
- 3) goes from glass up to detector.

Each step can be considered as certain amount of shrinking and turning of the unit arrow.

You'll remember that earlier, did not consider all ways light could reflect off glass,
which requires drawing and adding lots of little tiny arrows.

In order to avoid all detail,

gave impression that light goes down to particular point on surface of glass
and that it doesn't spread out.

When light goes from one point to another, it does, in reality, spread out (unless fooled by a lense),
and there is some shrinkage of unit arrow associated with that.

For moment, however,

I would like to stick to simplified view that light does not spread out,
and so it is appropriate to **disregard shrinkage**.

It is also appropriate to assume that since light doesn't spread out,
every photon that leaves source ends up at either A or B.

So: in 1st step no shrinking, but is turning

corresponds to amount of turning by imaginary stopwatch hand
as times photon going from source to front surface of glass.

In this example, arrow for 1st step ends up with length of 1 at some angle - say 5 o'clock.

2nd step is reflection of photon by glass.

Here, is sizable shrink - from 1 to 0.2 - and half turn.

(numbers seem arbitrary now: depend upon whether light reflected by glass
or some other material. Later, explain them).

Thus 2nd step represented by amplitude of length 0.2 and direction of 6 o'clock (half a turn).

Last step is photon going from glass up to detector.

Here, as in 1st step, no shrinking,

but is turning - say distance slightly shorter than step 1, and arrow points toward 4 o'clock.

Now “multiply” arrows 1,2, and 3 in succession (add three angles, and multiply lengths).

Net effect of 3 steps - 1) turning, 2) shrink and half turn, and 3) turning -

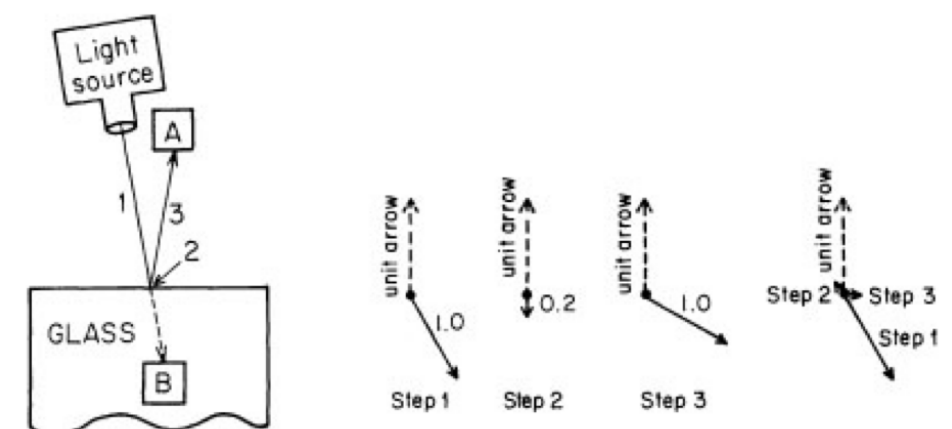
is same as earlier:

turning from steps 1 and 3 - (5 o'clock plus 4 o'clock)

is same amount of turning got

when let stopwatch run for whole distance (9 o'clock);

extra half turn from step 2 makes arrow point in direction opposite stopwatch hand, as did earlier,



and shrinking to 0.2 in 2nd step

leaves arrow whose square represents 4% partial reflection observed for single surface.

In this experiment,

there is a question that we didn't look at earlier:

What about photons that go to B

the ones that are transmitted by surface of glass?

Amplitude for photon to arrive at B must have length near 0.98,

since $0.98 * 0.98 = 0.9604$, which is close enough to 96%.

This amplitude can also be analyzed by breaking it down into steps (Figure right).

1st step

is same as for path to A

photon goes from light source down to glass

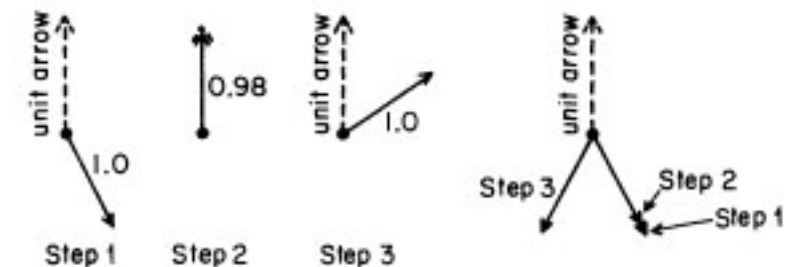
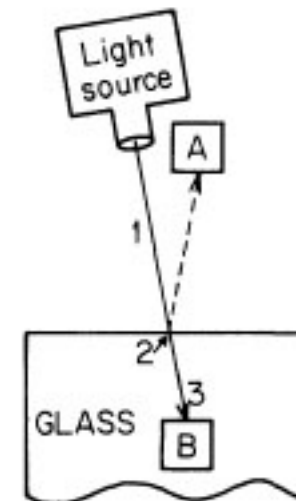
unit arrow is turned toward 5 o'clock.

2nd step

is photon passing through surface of glass:

there is no turning associated with transmission,

just bit of shrinking - to 0.98.



3rd step -

photon going through interior of glass

involves no additional turning and no shrinking.

Net result is arrow of length 0.98 turned in some direction

whose square represents probability that photon will arrive at B - 96%.

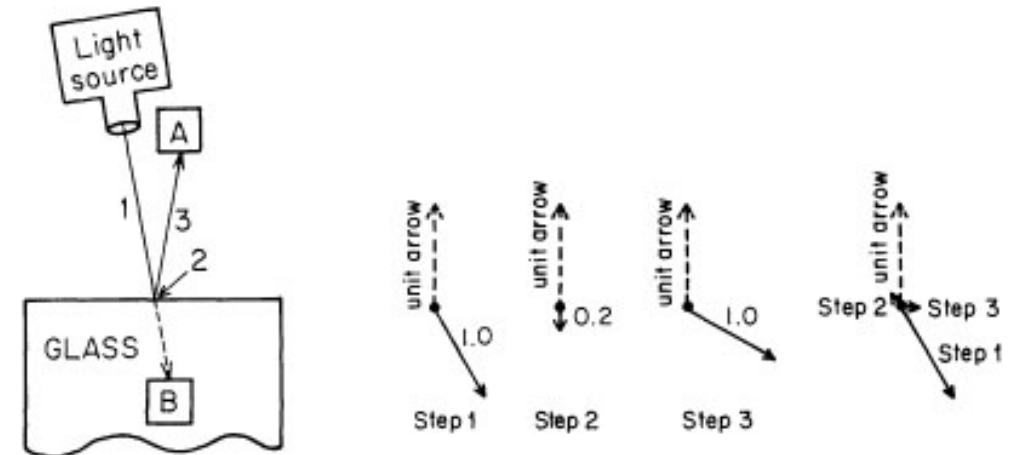
Now look at partial reflection by 2 surfaces again.

Reflection from front surface

is same as for single surface,

so 3 steps for front surface reflection

are same as saw moment ago.



Reflection from back surface can be broken down into 7 steps (Figure right).

Involves

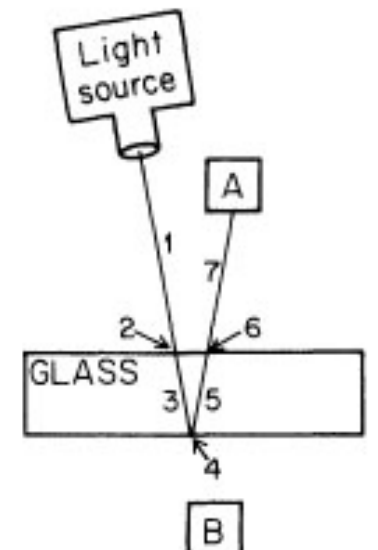
turning equal to total amount of turning of stopwatch hand timing photon over entire distance (steps 1, 3, 5, and 7)

shrinking to 0.2 (step 4) and two shrinks to 0.98 (steps 2 and 6).

Resulting arrow ends up in same direction as before,

but length is about 0.192 ($0.98 * 0.2 * 0.98$),

which we approximated as 0.2 earlier.



In summary,

here are rules for reflection and transmission of light by glass :

- 1) reflection from air back to air (off front surface) involves shrink to 0.2 and half turn;
- 2) reflection from glass back to glass (off back surface) also involves shrink to 0.2, but no turning;
- and 3) transmission from air to glass or from glass to air involves shrink to 0.98 and no turning in either case.

Perhaps too much of good thing,

but cannot resist showing you cute further example of how things work
and are analyzed by these rules of successive steps.

Move detector to location below glass,

and consider something didn't talk about earlier -

probability of **transmission** by two surfaces of glass (Figure right).

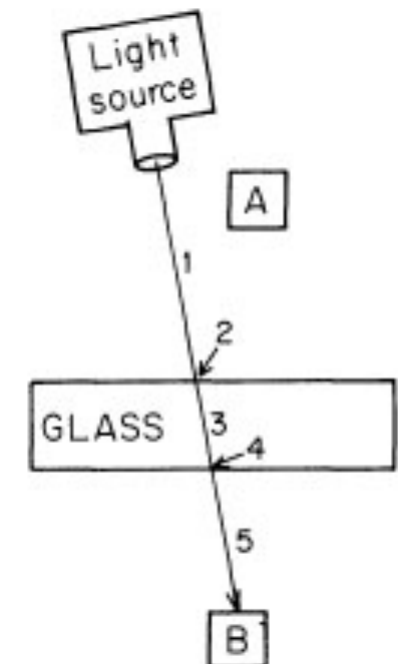
Of course you know answer:

probability of photon to arrive at B

is simply 100% minus probability to arrive at A,
which worked out beforehand.

Thus, if found chance to arrive at A is 7%,

chance to arrive at B must be 93%.



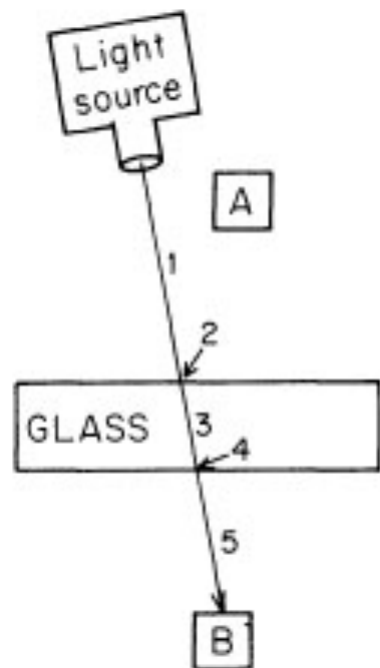
And as chance for A varies from zero through 8% to 16%
(due to different thicknesses of glass),
chance for B changes from 100% though 92% to 84%.

That is right answer,

but we were expecting to calculate all probabilities by squaring a final arrow.

How do we calculate amplitude arrow for transmission by layer of glass,
and how does it manage to vary in length so appropriately
as to fit with length for A in each case,

so probability for A and probability for B always add up to exactly 100%?



Look into **details**.

For photon to go from source to detector below glass, at B,

5 steps involved.

Let's shrink and turn unit arrows as go along.

1st 3 steps are same as in previous example:

photon goes from source to glass (turning, no shrinking);

photon transmitted by front surface (no turning, shrinking to 0.98);

photon goes through glass (turning, no shrinking).

4th step -

photon passes through back surface of glass

is same as 2nd step,

as far as shrinks and turns go:

no turns, but shrinkage to 0.98 of 0.98,

so arrow now has length of 0.96.

Finally, photon goes through air again, down to detector

means more turning, but no further shrinking.

Result is arrow of length 0.96,

pointing in some direction determined by successive turnings of stopwatch hand.

An arrow whose length is 0.96 represents probability of 92% (0.96 squared),

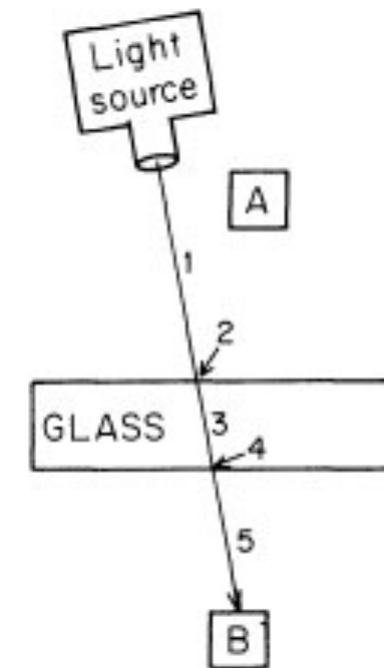
which means average of 92 photons reach B out of every 100 that leave source.

Also means that 8% of photons reflected by 2 surfaces and reach A.

But found out earlier that 8% reflection by 2 surfaces only right sometimes (“twice a day”)-

that in reality,

reflection by 2 surfaces fluctuates in cycle from 0 to 16% as thickness of layer steadily increases.



What happens when glass is just right thickness to make partial reflection of 16%?

For every 100 photons that leave source, 16 arrive at A

and 92 arrive at B,

which means 108% of light has been accounted for - horrifying!

Something is wrong.

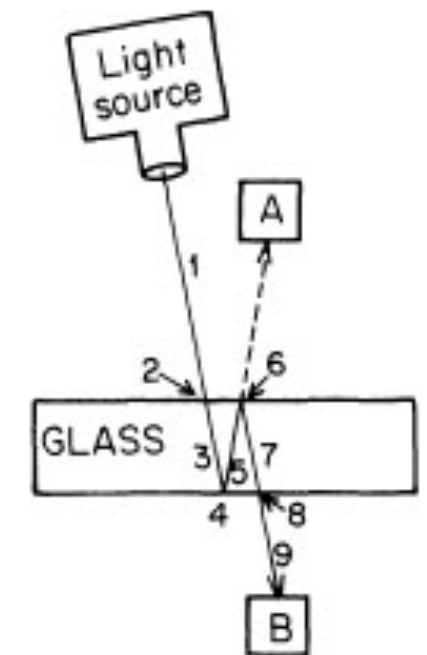
We simply neglected to consider **all** ways light could get to B!

For instance,

could bounce off back surface,

go up through glass as if were going to A,

but then reflect off front surface, back down toward B (Figure right).



Path takes 9 steps.

Let's see what happens successively to unit arrow

as light goes through each step.

1st step - photon goes through air - turning; no shrinking.

2nd step - photon passes through glass - no turning, but shrinking to 0.98.

3rd step - photon goes through glass - turning, no shrinking.

4th step - reflection off back surface - no turning, but shrinking to 0.2 of 0.98, or 0.196.

5th step - photon goes back up through glass - turning; no shrinking.

6th step - photon bounces off front surface

(really a “back” surface, because photon stays inside glass)

no turning, but shrinking to 0.2 of 0.196, or 0.0392.

7th step - photon goes back though glass - more turning; no shrinking.

8th step - photon passes through back surface

no turning, but shrinking to 0.98 of 0.0393, or 0.0384.

Finally, 9th step - photon goes through air to detector - turning; no shrinking.

Result of all this shrinking and turning is amplitude of length 0.0384

call it 0.04, for practical purposes -

and turned at angle that corresponds to total amount of turning by stopwatch

as times photon going through longer path.

Arrow represents second way that light can get from source to B.

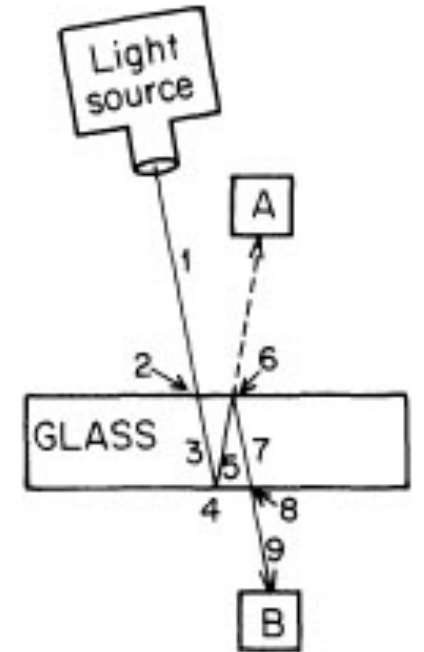
Now have two alternatives,

so must add 2 arrows - arrow for more direct path, whose length is 0.96,

and arrow for longer way, whose length is 0.04 - to make final arrow.

2 arrows are not usually in same direction,

because changing thickness of glass changes relative direction of 0.04 arrow to 0.96 arrow.



But look how nicely things work out:

extra turns made by stopwatch timing photon during steps 3 and 5 (on way to A)

are exactly equal to extra turns makes timing photon during steps 5 and 7 (on way to B).

That means when 2 reflection arrows are cancelling each other

to make final arrow representing zero reflection,

arrows for transmission reinforcing each other to make arrow of length $0.96 + 0.04$, or 1

when probability of reflection is zero,

probability of transmission is 100% (Figure below).

Reflection

$0.2 - 0.2$
0

0.4
16 %
 $0.2 + 0.2$

Transmission

$0.96 + 0.04$
100 %
1

$0.96 - 0.04$
84 %
0.92

When arrows for reflection are reinforcing each other to make amplitude of 0.4, arrows for transmission are going against each other, making amplitude of length $0.96 - 0.04$, or 0.92 when reflection is calculated to be 16%, transmission is calculated to be 84% (0.92 squared). See how clever Nature is with Her rules to make sure that always come out with 100% of photons accounted for!

Finally,

let me tell you that there is extension to the rule
that tells us when to multiply arrows:

arrows are to be multiplied

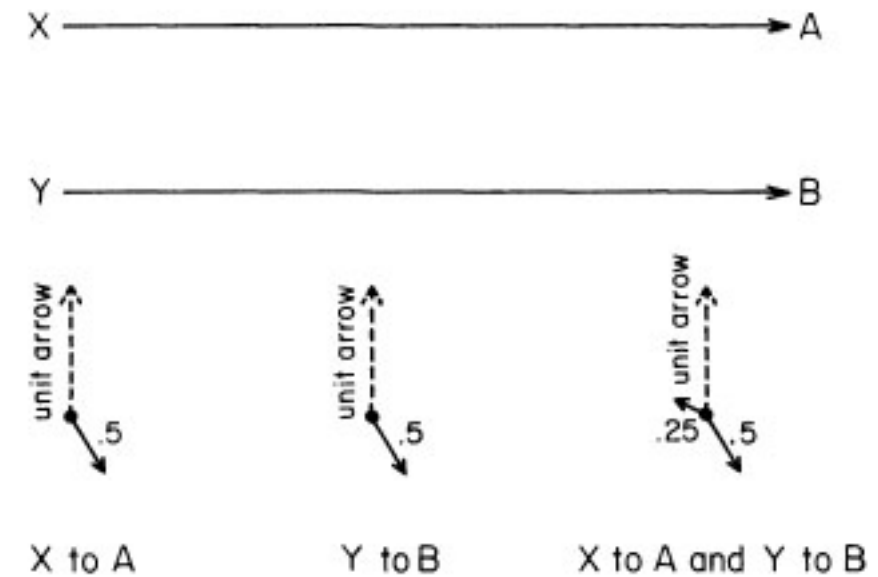
not only for an event that consists of a number of things
happening independently and possibly simultaneously.

For example, suppose have 2 sources, X and Y,

and 2 detectors, A and B (Figure right),

and want to calculate probability for following event:

after X and Y each lose a photon, A and B gain a photon.



In this example, photons travel through space to get to detectors

they are neither reflected or transmitted

so **now is good time** to stop disregarding fact that light spreads out as it goes along.

Now present complete rule for monochromatic light

travelling from one point to another point through space

nothing approximate here, and no simplification.

This is all there is to know

about monochromatic light going through space (disregarding polarization):

angle of arrow depends on imaginary stopwatch hand,

which rotates certain number of times per inch (depending on color of photon);

length of arrows inversely proportional to distance light goes

in other words arrow shrinks as light goes along.

This rule checks out with what is taught in a physics class

amount of light transmitted over distance varies inversely as square of distance

because arrow that shrinks to half its original size has square $1/4$ as big.

Suppose arrow for X to A is 0.5 in length and is pointing toward 5 o'clock, as is arrow for Y to B (Figure below).

Multiplying one arrow by other,

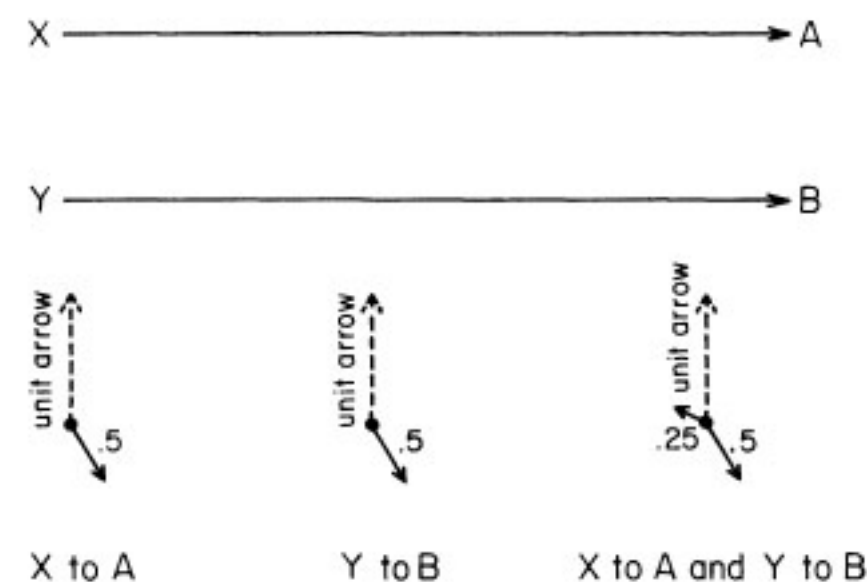
get final arrow of length 0.25, pointed at 10 o'clock.

But wait!

There is another way this event could happen:

photon from X could go to B,

and photon from Y could go to A.



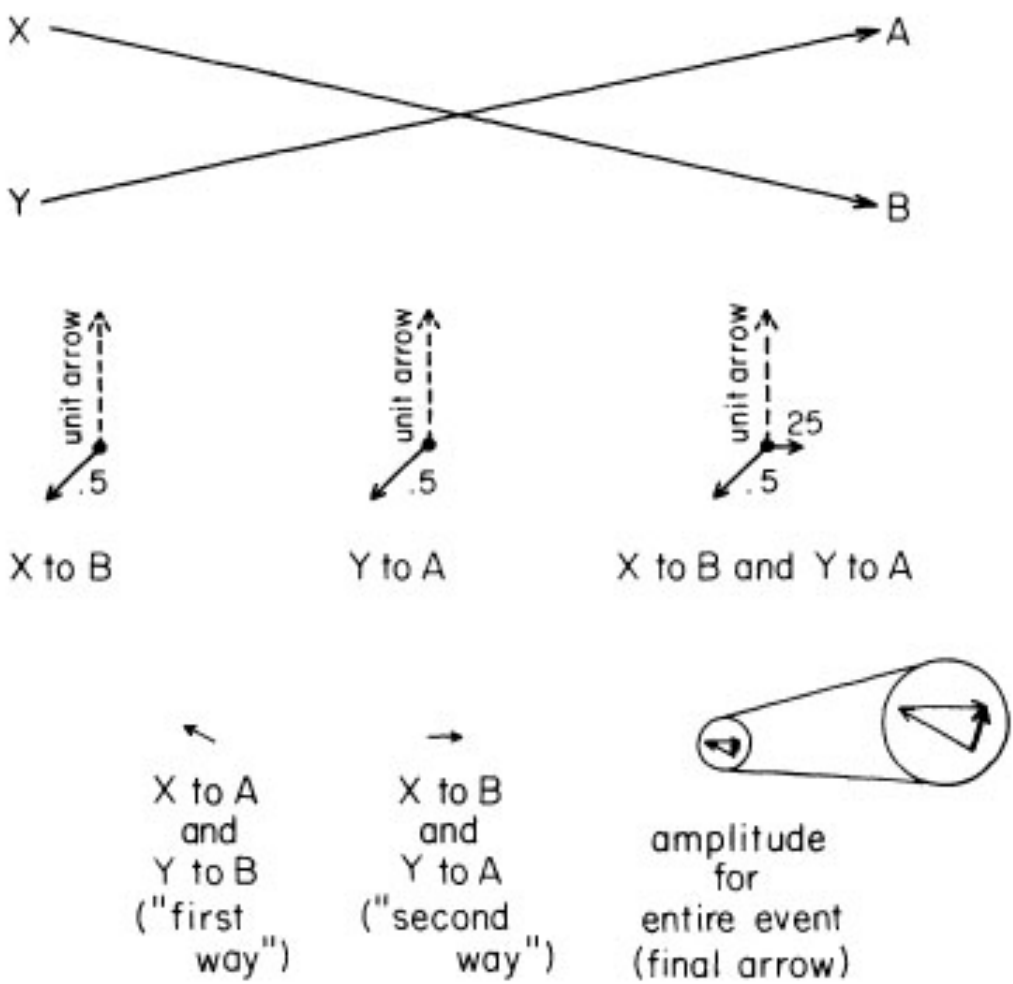
Each of these sub-events has amplitude,
and these arrows must be drawn
and multiplied to produce an amplitude
for this particular way event could happen (Figure right).

Since amount of shrinkage over distance
is very small compared to amount of turning,
arrows from X to B and Y to A
have essentially same length
as other arrows, 0.5,
but their turning is quite different:

stopwatch hand rotates 36,000 times per inch for red light,
so even tiny difference in distance
results in substantial difference in timing.

Amplitudes for each way event could happen
are added to produce final arrow.

Since their lengths are essentially same,
is possible for arrows to cancel each other out
if their directions are opposed to each other.



Relative directions of 2 arrows

can be changed by changing distance between sources or detectors:

simply moving detectors apart or together a little bit

can make probability of event amplify or completely cancel out,

just as in case of partial reflection by 2 surfaces.

This phenomenon, called Hanbury-Brown-Twiss effect,

has been used to distinguish between single source and double source of radio waves in deep space,

even when 2 sources are extremely close together.

In this example,

arrows were multiplied

and then added to produce final arrow (amplitude for event),

whose square is probability of event.

It is emphasized

that no matter how many arrows draw, add, or multiply,

objective is to calculate single final arrow for event.

Mistakes often made by physics students at first

because do not keep this important point in mind.

They work for so long analyzing events involving single photon

that begin to think that arrow somehow associated with photon.

But these arrows are probability amplitudes,

that give, when squared, probability of complete event.

Later we will begin process of simplifying and explaining properties of matter

to explain where shrinking to 0.2 comes from,

why light appears to go slower through glass or water than through air, and so on

because have been cheating so far:

photons don't really bounce off surface of glass.

Will show how photons do nothing but go from one electron to another,

and how reflection and transmission are really result of electron picking up photon,

“scratching its head”, so to speak, and emitting new photon.

This simplification of everything we have talked about so far is very pretty and helpful.

Electrons and Their Interactions - Next lecture starts here

Continuing our discussion on theory of quantum electrodynamics - interaction of light and electrons.

Most of phenomena familiar with involve interaction of light and electrons - all of chemistry and biology, for example.

Only phenomena not covered by this theory are phenomena of gravitation and nuclear phenomena; everything else contained in this theory

Found out earlier that have no satisfactory mechanism to describe even simplest of phenomena, such as partial reflection of light by glass.

Have no way to predict whether give photon will be reflected or transmitted by glass.

All can do is calculate probability that particular event will happen -

whether light will be reflected, in this case.

(This about 4%, when light shines straight down on single surface of glass;

probability of reflection increases as light hits glass at more of a slant.

When deal with probabilities under ordinary circumstances, we are following specific

“rules of composition”:

1) if something can happen in alternative ways, add probabilities for each of different ways

2) if event occurs as succession of steps - or depends on number of things happening

independently - then multiply probabilities of each of steps (or things).

In the wild and wonderful world of quantum physics,

probabilities calculated as square of length of arrow:

where would have expected to add probabilities under ordinary circumstances,

find ourselves “adding” arrows;

where normally would have multiplied probabilities, “multiply” arrows.

Peculiar answers that we get from calculating probabilities in this manner

match **perfectly** with the results of experiment!!

I rather delighted that we must resort to such peculiar rules
and strange reasoning in order to understand Nature,
and enjoy telling people about it.

There are no “wheels and gears” beneath this analysis of Nature;
if we want to understand Her, this is what you have to do.

Before I proceed with discussion,

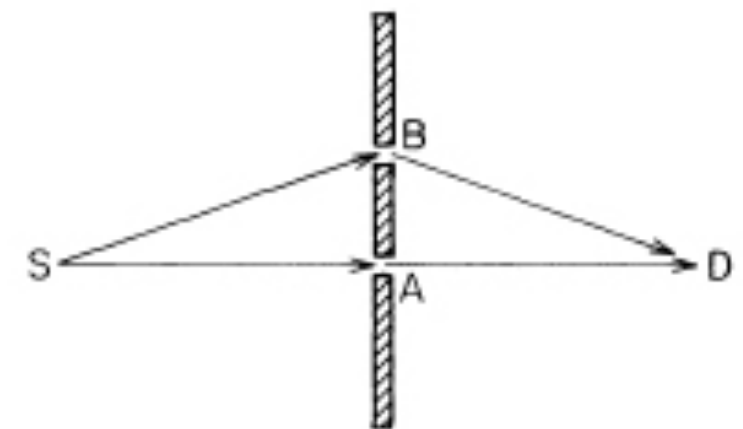
Let me show another example of how light behaves when using the rules.

Let us talk about very weak light of one color - one photon at a time -
going from source, at S, to detector , at D (Figure below).

Put screen in between source and detector
and make two very tiny holes
few millimeters apart from each other,
at A and B.

(If source and detector are 100 centimeters apart,
holes themselves have to be smaller than 1/10 of millimeter).

Put A in line with S and D,
and put B somewhere to side of A,
not in line with S and D.



When close hole at B,

get certain number of clicks at D

which represents photons that came through A

(say detector clicks average of one time for every 100 photons that leave S, or 1%).

When close hole at A and open hole at B,

know from earlier that get nearly same number of clicks, on average, because holes are small.

(When “squeeze” light too much, rules of ordinary world - such as light goes in straight lines - fall apart.)

When open both holes get complicated answer, because phenomenon of **interference** is present:

If holes are certain distance apart, get more clicks than expected 2% (maximum about 4%);

if two holes are slightly different distance apart, might get no clicks at all.

One would normally think that opening 2nd hole

would always increase amount of light reaching detector,

but that's not what actually happens.

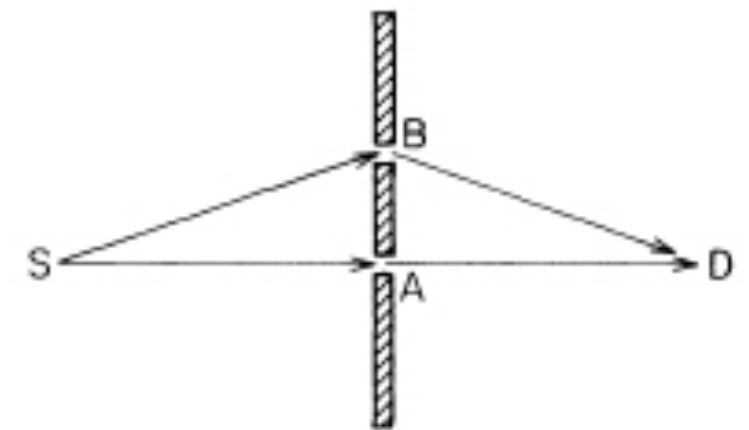
And so saying that light goes “either one way or the other” is false.

Still catch myself saying,

“Well, it either goes this way or that way”,

but when I say that, have to keep in mind that means **only in sense of adding amplitudes:**

photon has an amplitude to go one way, and amplitude to go other way



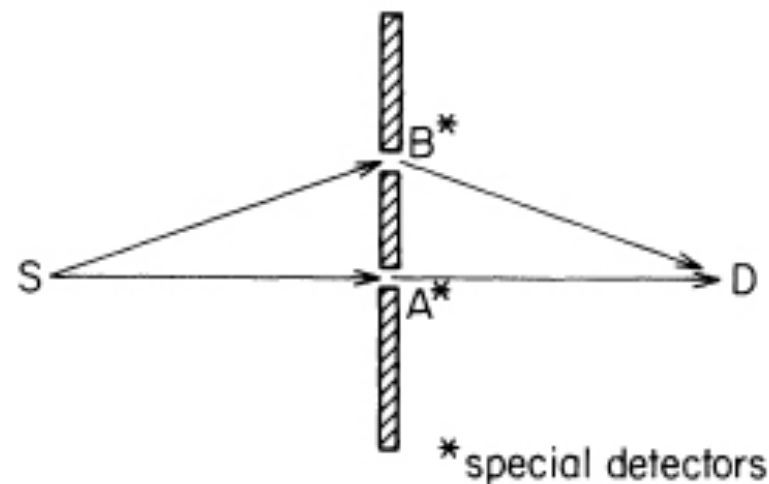
If amplitudes oppose each other, light won't get there
even though both holes are open.

Now, here's extra twist to strangeness of Nature to tell you about.

Suppose put in some special detectors - one at A and one at B

(possible to design detector that can tell whether photon went through it)

so can tell through which hole(s) photon goes through when both holes open (Figure below).



Since probability that single photon will get from S to D

affected only by distance between holes,

must be some sneaky way that photon

divides in two and then comes back together again, **right?**

According to hypothesis,

detectors at A and B should always go off together

(at half strength, perhaps?),

while detector at D should go off with probability of from zero to 4%,

depending on distance between A and B.

Here's what actually happens:

detectors at A and B never go off together

either A or B goes off.

Photon does not divide in two; it goes one way or other.

Furthermore,

under such conditions detector at D goes off 2% of time

simple sum of probabilities for A and B (1% + 1%).

2% not affected by spacing between A and B;

so interference disappears when detectors are put in at A and B!

Nature has got it cooked up so

we'll will never be able to figure out how She does it:

if we put instruments in to find out which way light goes,

we can find out, all right, but wonderful interference effects disappear.

But if we don't have instruments that can tell which way light goes,

then interference effects come back! Very strange, indeed!

To understand this paradox,

I remind you of most important principle: in order to correctly calculate probability of event,

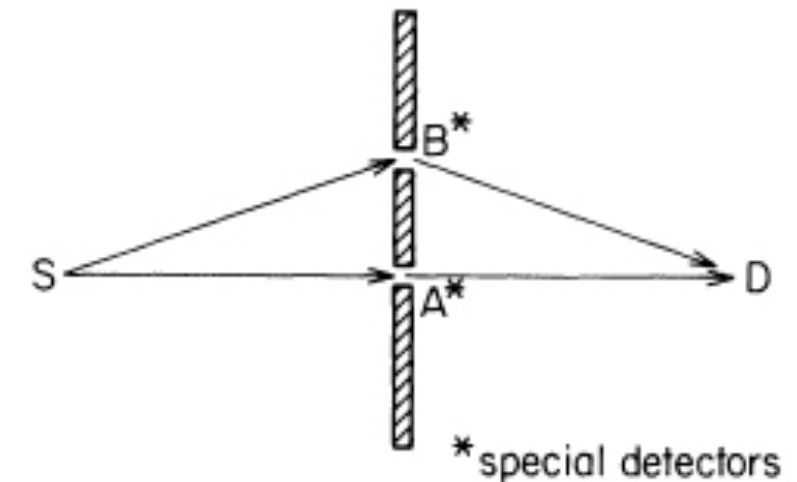
one must be very careful to define complete event clearly

in particular, what initial conditions and final conditions of experiment are.

Look at equipment before and after experiment, and look for changes.

When were calculating probability that photon gets from S to D

with no detectors at A or B, event was, simply, detector at D makes click.



When click at D was only change in conditions,

there was no way to tell which way photon went, so there was interference.

When put in detectors at A and B, **changed** problem.

Now, turns out, are 2 complete events — 2 sets of final conditions - that are distinguishable:

1) detectors at A and D go off, or

2) detectors at B and D go off.

When there are number of possible final conditions in experiment,
must calculate probability of each as separate, complete event.

To calculate amplitude that detectors at A and D go off,

multiply arrows that represent following steps:

photon goes from S to A, photon goes from A to D, and detector at D goes off.

Square of final arrow is probability of this event - 1%

same as when hole at B was closed, because both cases have exactly same steps.

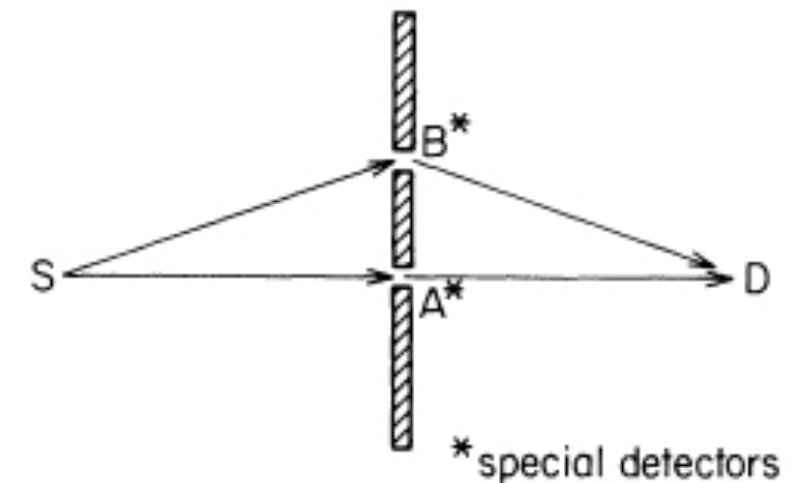
Other complete event is detectors at B and D goes off.

Probability of this event calculated in similar way, and is also same as before - about 1%.

If want to know how often detector at D goes off

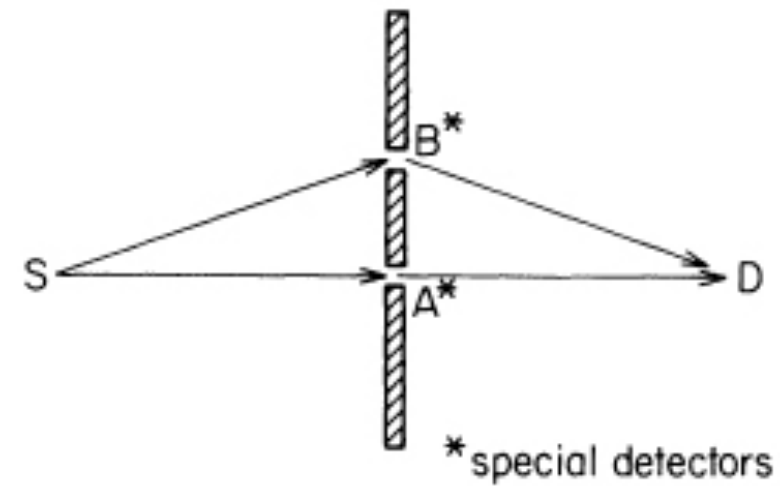
and don't care whether was A or B that went off in process,

probability is simple sum of two events - 2%.



In principle,

if there is something left in system
that one could have observed
to tell which way photon went,
have different “final states”
(distinguishable final conditions),
and add probabilities - not amplitudes
for each final state.



Complete story on this situation very interesting:

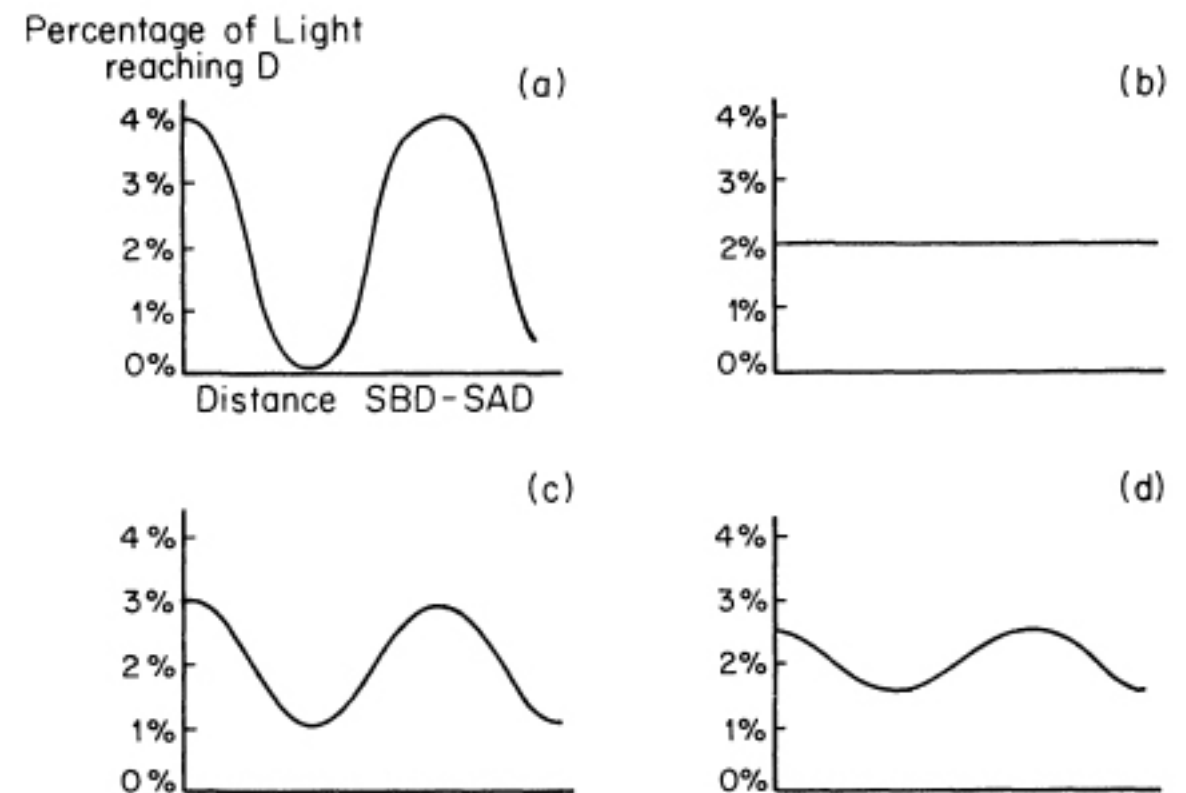
if detectors at A and B not perfect,

and detect photons only some of time,

then there are 3 distinguishable final conditions:

- 1) detectors at A and D go off;
- 2) detectors at B and D go off, and
- 3) detector at D goes off alone,

with A and B unchanged (left in their initial state).



Probabilities for 1st two events

calculated in way explained above

(except will be extra step

shrink for probability that detector at A [or B] goes off, since detectors not perfect).

When D goes off alone, can't separate 2 cases, and Nature plays with us by bringing in interference

same peculiar answer would have had if there were no detectors

(except that final arrow is shrunk by amplitude that detectors do not go off).

Final result is mixture,

simple sum of all 3 cases (Figure above), as reliability of detectors increases, get less interference.

I have pointed out these things because the more you see how strangely Nature behaves, and better understand why it is hard is to make model that explains how even the simplest phenomena work.

So theoretical physics has given up on that.

We saw earlier how an event can be divided into alternative ways

and how arrow for each way can be “added”.

Also saw how each way can be divided into successive steps,

how arrow for each step can be regarded as transformation of unit arrow,

and how arrows for each step can be “multiplied” by successive shrinks and turns.

Thus we are familiar with all necessary rules for drawing and combining arrows (that represent bits and pieces of events) to obtain final arrow, whose square is probability of an observed physical event.

It is **natural to wonder**

how far can we **push** this process of splitting events into simpler and simpler sub-events.

What are the **smallest** possible bits and pieces of events?

Is there **limited** number of bits and pieces that can be compounded
to form all phenomena that involve light and electrons?

Is there **limited** numbers of “letters” in this language of quantum electrodynamics
that can be combined to form “words” and “phrases”
that describe nearly every phenomenon of Nature?

Answer is yes; that number is 3.

There are **only three basic actions**

needed to produce **all** of phenomena associated with light and electrons.

Before I tell you what 3 basic actions are,

I should properly introduce you to **actors**.

Actors are photons and electrons.

Photons, particles of light, discussed at length earlier.

Electrons discovered in 1895 as particles:

could count them;

could put one of them on oil drop and measure electric charge and mass.

Gradually became apparent that motion of these particles accounted for electricity in wires.

Shortly after electrons discovered

it was thought that atoms were like little solar systems,
made up of central, heavy part (called nucleus) and electrons,
which went around in “orbits”, much like planets do when they go around sun.

If think that’s way atoms are, then you are back in 1910.

In 1924 Louis De Broglie found that there was wavelike character associated with electrons,
and soon afterwards, C.J. Davisson and L.H. Germer of Bell Laboratories
bombarded nickel crystal (the slits) with electrons and showed that they, too,
bounced off at crazy angles (just like X-rays do), and that these angles could be calculated
from De Broglie’s formula for wavelength of electron.

When we look at photons on a large scale,
much larger than distance required for one stopwatch turn
phenomena that we see are very well approximated by rules such as “light travels in straight lines”,
because there are enough paths around path of minimum time to reinforce each other,
and enough other paths to cancel non minimum time paths out.

But when space through which photon moves becomes too small (such as tiny holes in screen),
these rules fail —-> discover that light doesn’t have to go in straight lines,
there are interferences created by two holes, and so on as we saw.

Same situation exists with electrons:

when seen on large scale, travel like particles, on definite paths —> **know where electron is.**

But on small scale, such as inside an atom,

space so small that there is no main path, no “orbit”;

there are all sorts of ways electron could go, each with an amplitude.

Phenomenon of interference becomes very important,

and have to sum arrows to **predict where electron is likely to be.**

Rather interesting to note that electrons looked like particles at first,

and wave-like character later discovered.

On other hand, apart from Newton making mistake and thinking that light was “corpuscular”,

light looked like waves at first, and characteristics as particle discovered later.

In fact, both objects behave somewhat like waves, and somewhat like particles.

When you study QM, find that these classical words make no sense!! Wave-particle duality is a scam!

In order to save ourselves from inventing new words such as “wavicles”,

we have chosen to call objects “particles”,

but we know they obey these rules for drawing and combining arrows we have been explaining.

Appears that **all** “particles” in Nature

quarks, gluons, neutrinos, and so forth (discussed later)

behave in this same quantum mechanical way.

So now, present to 3 basic actions, from which all phenomenon of light and electrons arise.

Action #1: A photon goes from place to place.

Action #2: An electron goes from place to place.

Action #3: An electron emits or absorbs a photon.

Each of these actions has an amplitude

an arrow

that can be calculated according to certain rules.

In a moment, I will tell you those rules, or laws,

out of which we can make whole world (aside from nuclei, and gravitation, as always!).

Now, the stage on which these actions take place is not just space,

it is space and time.

Next two slides were for your review - hope you read - will not cover in lecture

Until now, have disregarded problems concerning time,

such as exactly when photon leaves source and exactly when arrives at detector.

Although space 3-dimensional,

I am going to reduce it to 1 dimension on graphs I am going to draw:

They show a particular object's location in space on horizontal axis, and time on vertical axis.

1st event I am going to draw in space and time - or space-time, as might call it -
is a baseball standing still (Figure right).

On Thursday morning, which I label as T_0 ,

the baseball occupies certain space, which I label as X_0 .

Few moments later, at T_1 , occupies same space, because it is standing still.

Few moments later, at T_2 , baseball still at X_0 .

So diagram of baseball standing still is vertical band,

going straight up, with baseball all over it inside the band.

What happens if have baseball drifting in weightlessness of outer space,

going straight towards a wall?

Well on Thursday morning (T_0) starts at X_0 (Figure right),

but little bit later, not in same place - has drifted over little bit, to X_1 .

As baseball continues to drift,

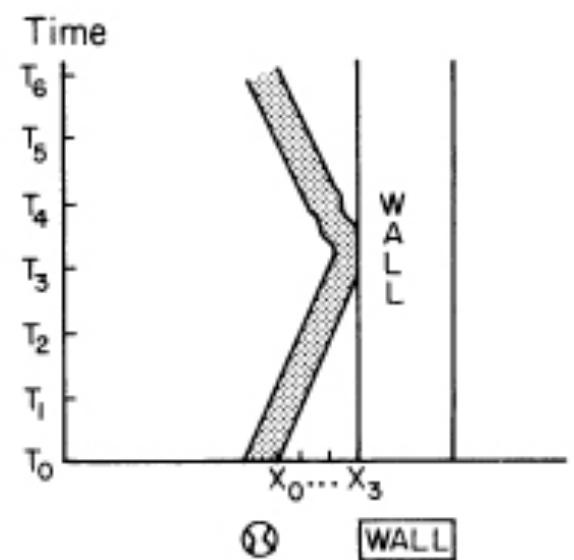
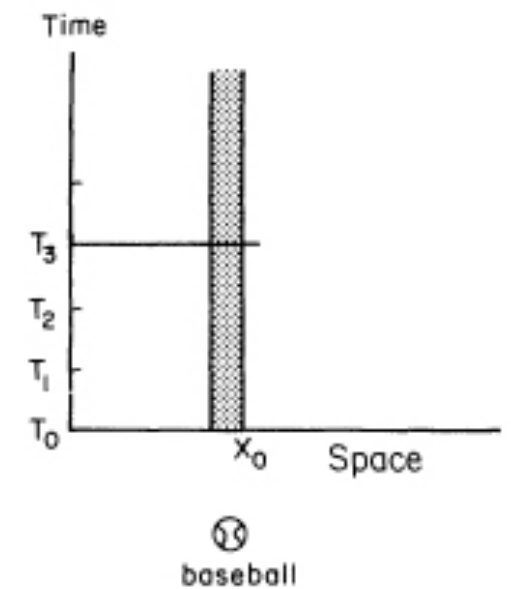
creates slanted “band of baseball” on diagram in space-time.

When baseball hits wall (standing still $\rightarrow$ vertical band),

goes back other way,

to exactly where it came from in space (X_0),

but to different point in time (T_6).



As for time scale,

most convenient to represent time not in seconds,
but in much smaller units.

Since will be dealing with photons and electrons,
which move very rapidly,

also going to have 45° angle represent something going speed of light.

For example, for particle moving at speed of light from (X_1, T_1) to (X_2, T_2) ,
horizontal distance between X_1 and X_2 is same as vertical distance
between T_1 and T_2 (Figure right).

Factor by which time stretched out

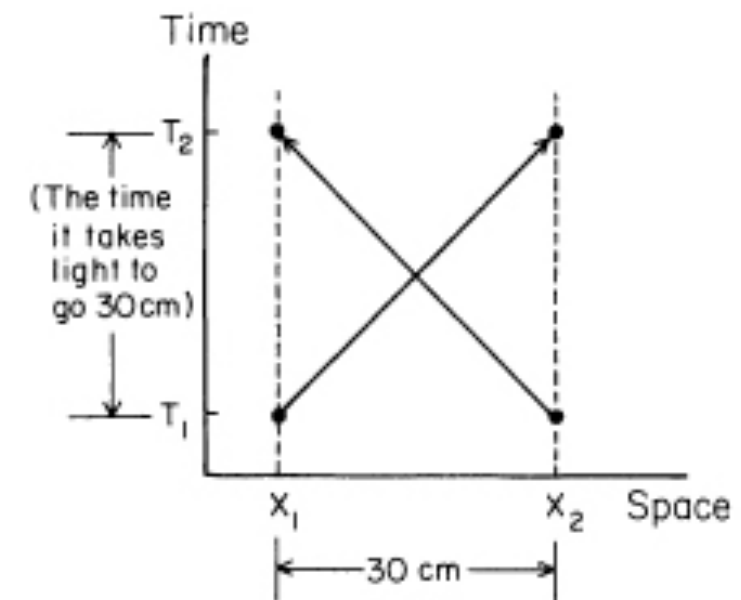
(to make 45° angle represent particle going speed of light) called c ,
and find c 's flying around everywhere in Einstein's formulas
result of unfortunate choice of second as unit of time,
rather than time takes light to go 1 meter.

Thus time scale we will use in graphs

will show particles going at speed of light to be travelling at 45° angle through space-time.

Amount of time takes light to go 30 centimeters - from X_1 to X_2 or from X_2 to X_1 -
is about one-billionth of a second.

just using " ct " on vertical axis



Now on with our discussion:

Now, look at 1st basic action in detail

photon goes from place to place.

Draw action as wiggly line from A to B

for no good reason.

Should be more careful:

should say,

photon that is known to be at given place

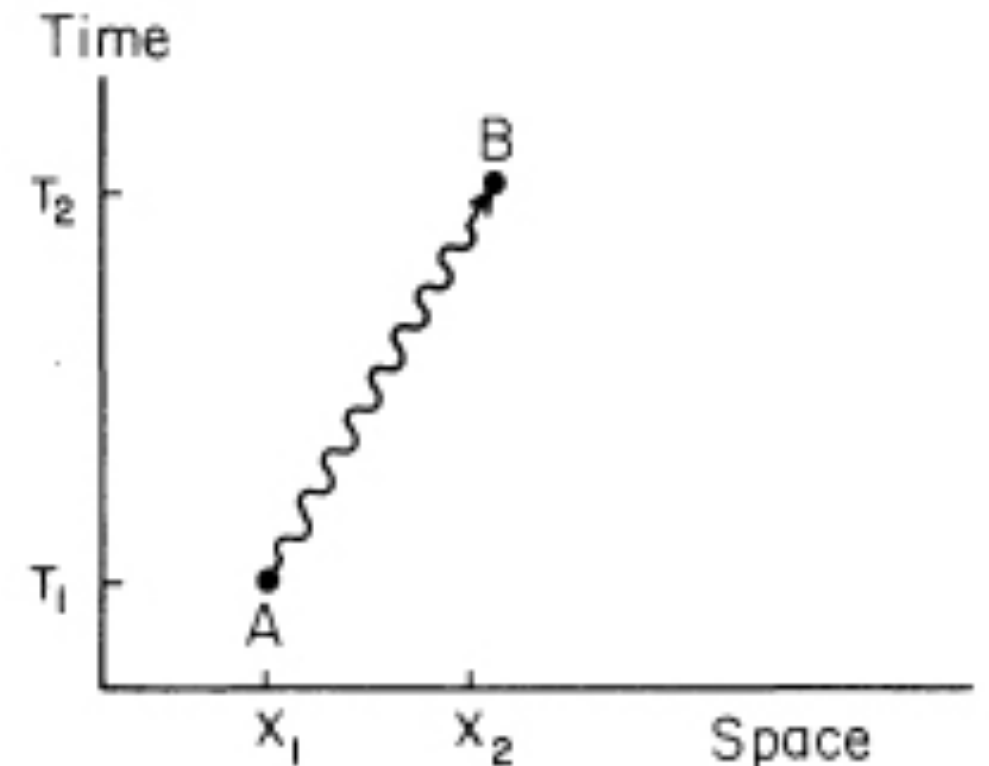
**at given time has certain amplitude to get to another
place at another time.**

On space-time graph (Figure right),

photon at point A - at X_1 and T_1 -

has amplitude to appear at point B - at X_2 and T_2 .

Size of this amplitude we call **$P(A \text{ to } B)$** .



There is formula for size of this arrow, **P(A to B)**.

It is one of great laws of Nature, and and is very simple.

It depends on difference in distance and difference in time between 2 points.

These differences are expressed mathematically as $(X_2 - X_1)$ and $(T_2 - T_1)$.

In these lectures, we are plotting point's location space in 1 dimension, along x-axis.

To locate point in 3-dimensional space, “room” has to be set up,

and distance of point from floor and from each of 2 adjacent walls

(all at right angles to each other) has to be measured.

3 measurements labeled X_1 , Y_1 and Z_1 .

Actual distance from this point to 2nd point with measurements X_2 , Y_2 , Z_2

calculated using “three-dimensional Pythagorean Theorem”:

square of actual distance is

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2$$

Excess of this over time difference, squared -

$$(X_2 - X_1)^2 + (Y_2 - Y_1)^2 + (Z_2 - Z_1)^2 - (T_2 - T_1)^2$$

assuming $c=1$ for the moment

- is sometimes called “**the Interval**”, or I ,

remember earlier discussion

- and it is the combination that Einstein's theory of relativity **says** that **P(A to B)** must depend on.

Most of contribution to final arrow for $P(A \text{ to } B)$ is just where would expect it -

where difference in distance is equal to difference in time (that is, when interval I is **zero**).

But in addition; there is contribution when I not zero,

that is **inversely proportional** to the value of I :

it points in direction of 3 o'clock when I is more than zero (when light going faster than c),

and points toward 9 o'clock when I is less than zero (Figure below).

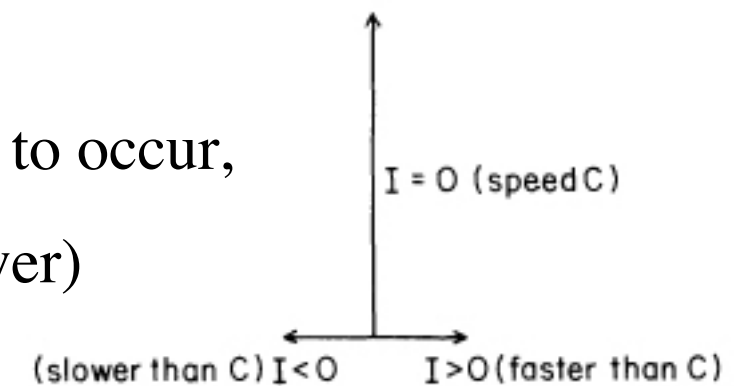
These later contributions cancel out in many macroscopic circumstances.

Major contribution to $P(A \text{ to } B)$ occurs at conventional speed of light -

when $(X_2 - X_1)$ is equal to $(T_2 - T_1)$ - where one would expect it all to occur,

but in QED world there is also amplitude for light to go faster (or slower)

then conventional speed of light.



Found out earlier that light doesn't go only in straight lines;

now, find out that it doesn't go only at speed of light!

May surprise you that there is amplitude for photon to go at speeds faster or slower than conventional speed, c .

Amplitudes for these possibilities very small compared to contribution from speed c ;

in fact, they cancel out when light travels over long distances.

(in actual macroscopic experiments!!)

However, when distances short - as in many of the diagrams we will be drawing - other possibilities become vitally important and must be considered.

So that's 1st basic action,

1st basic law of physics

photon goes from point to point.

That explains **all** about optics; that's the **entire** theory of light! (all we used earlier)

Well, not quite: left out **polarization** (as always),

and **interaction of light with matter**, which brings 2nd action.

2nd action fundamental to quantum electrodynamics is:

An electron goes from point A to point B in space-time.

Imagine electron as simplified, fake electron, with no polarization - what physicists call a “**spin-zero**” electron.

In reality, electrons have type of polarization, which **doesn't** add anything to main ideas we need; only **complicates** formulas somewhat.

Formula for amplitude for this action, which call **E(A to B)**

also depends on $(X_2 - X_1)$ and $(T_2 - T_1)$ (in same combination as described before)

as well as on a **number** we will call “**n**”,

A number that, once determined, enables all calculations to **agree** with experiment.

(will see later how we determine **n**'s value.)

—> a rather complicated formula, and I don't know how to explain it in simple terms.

However, can tell you that formula for $P(A \text{ to } B)$

a photon going from place to place in space-time

is **same** as that for $E(A \text{ to } B)$ - an electron going from place to place

if we set $\mathbf{n}$ to zero.

As I said the formula for $E(A \text{ to } B)$ is complicated,

but there is interesting way to explain what it amounts to.

$E(A \text{ to } B)$ can be represented as

the **giant sum of lots of different ways** an electron could go from point A to point B in space-time

(Figures at right):

electron could take “one-hop flight”,

going directly from A to B;

could take “two-hop flight”,

stopping at intermediate point C;

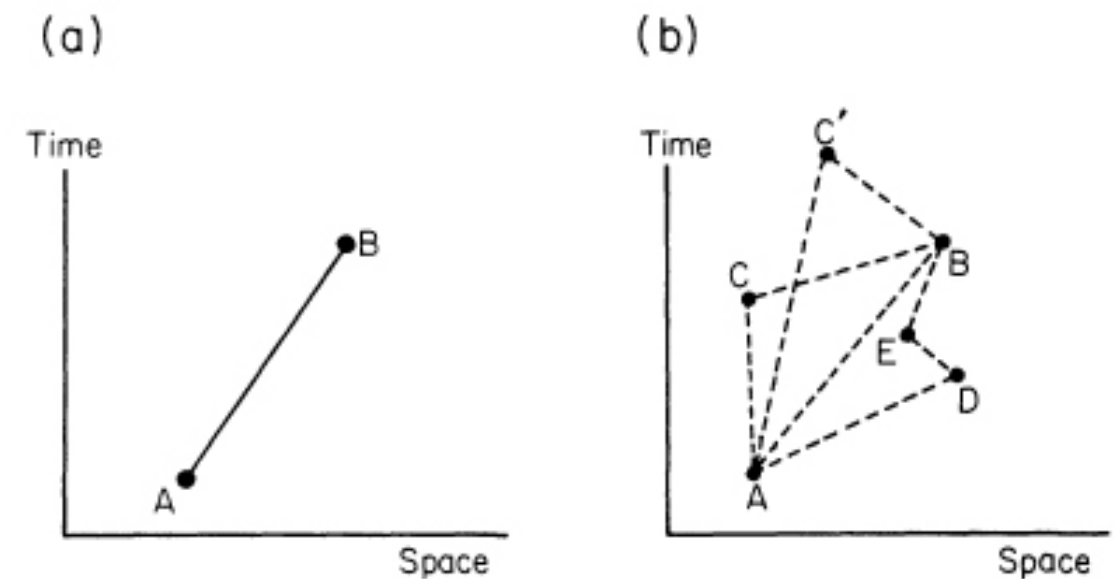
could take “three-hop flight”,

stopping at points D and E, and so on.

In such an analysis, amplitude for each “hop”

from one point B to another point E - is $P(B \text{ to } E)$,

i.e., **same** as amplitude for photon to go from point B to point E.



Amplitude for each “stop” represented by n^2 ,

n being same number mentioned before which we use to make the calculations come out right.

Formula for $E(A \text{ to } B)$ is thus series of terms:

$$P(A \text{ to } B) \text{ [“one-hop” flight]} + P(A \text{ to } C) * n^2 * P(C \text{ to } B) \text{ [“two-hop” flights, stopping at C]} \\ + P(A \text{ to } D) * n^2 * P(D \text{ to } E) * n^2 * P(E \text{ to } B) \text{ [“three-hop” flights, stopping at D and E]} + \dots$$

for all possible intermediate points C, D, E, and so on.

Note that as n gets larger,

the nondirect paths make a greater contribution to final arrow.

When n is zero (as for photon),

all terms with an n in them drop out (because = zero), leaving only first term.

which is $P(A \text{ to } B)$. **Thus $E(A \text{ to } B)$ and $P(A \text{ to } B)$ are closely related.**

3rd basic action is : electron emits or absorbs photon

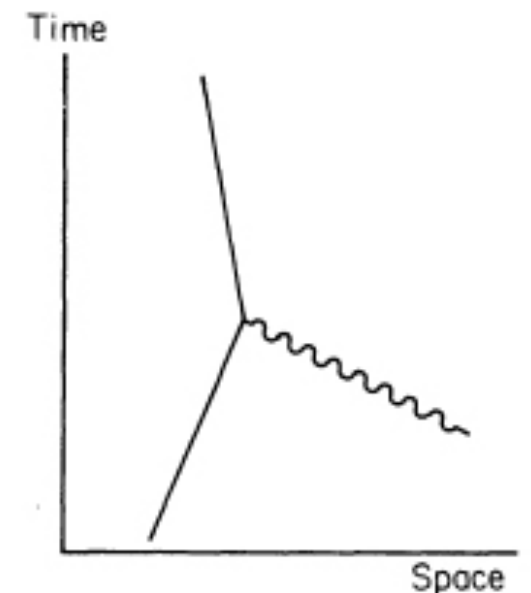
doesn't make any difference which.

Call this action a “junction”, or “coupling” or “vertex”.

To distinguish electrons from photons in my diagrams,

will draw each electron going through space-time as straight line.

Every coupling, therefore, is **junction** between two straight lines and wavy line (Figure above).



two “fermions” and one “boson” - more about that later

There is no complicated formula for amplitude of an electron to emit or absorb photon;
doesn't depend on anything - it is just a **number**!

This junction number = **j**

value is about **- 0.1**:

i.e., **shrink to about 1/10, and half turn.**

This number = amplitude to emit or absorb photon $\Leftrightarrow$ “**charge**” of particle.

Well, that's all there is to these basic actions

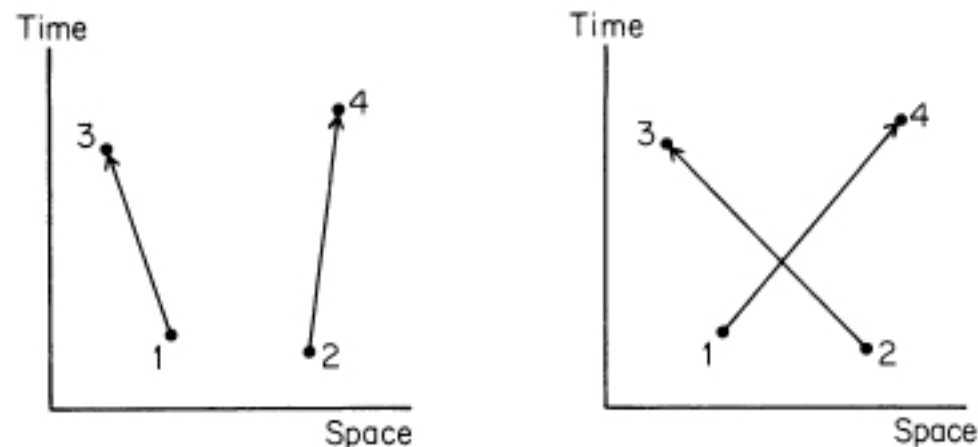
except for some slight complications due to polarization that we are always leaving out.

Next job is to put these 3 actions together

to represent circumstances that are somewhat more complicated.

For 1st example,

calculate probability that 2 electrons, at points 1 and 2 in space-time,
end up at points 3 and 4 (Figure below).



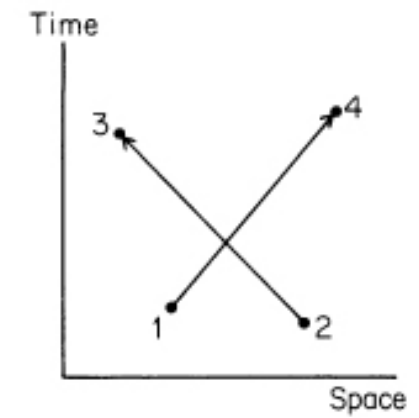
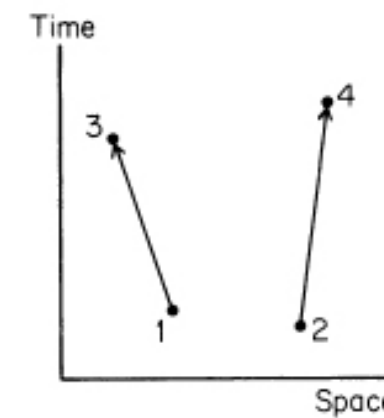
This event can happen in two main ways.

1st way is that electron at 1 goes to 3

computed by putting 1 and 3 into formula $E(A \text{ to } B)$,

which write $E(1 \text{ to } 3)$

and electron at 2 goes to 4 - computed by $E(2 \text{ to } 4)$.



These are two “sub-events” happening independently and simultaneously,

so two arrows are **multiplied** to produce arrow for this 1st way event could happen.

Therefore write formula for “first way arrow” as $E(1 \text{ to } 3) * E(2 \text{ to } 4)$.

Another way event could happen is that electron at 1 goes to 4

and electron at 2 goes to 3 - again, two independent/simultaneous sub-events.

“second way arrow” is $E(1 \text{ to } 4) * E(2 \text{ to } 3)$,

and add it to “first-way” arrow.

Had I included effects of polarization of electron,

“second-way” arrow would have been “subtracted” - turned 180° and added (more later).

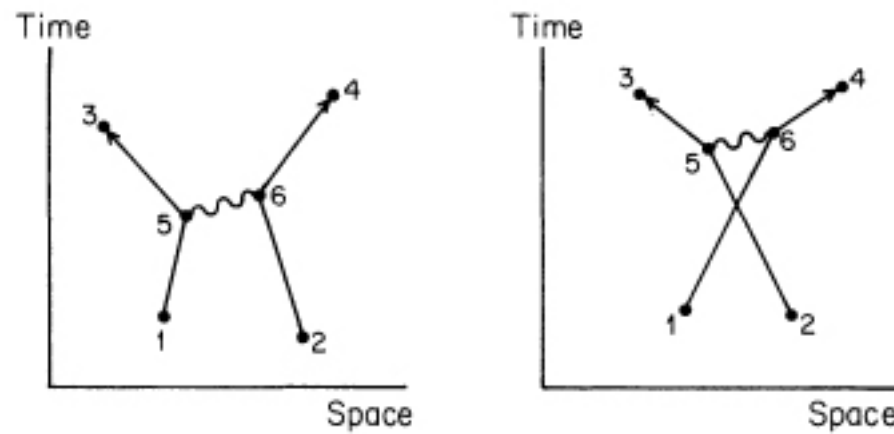
This is good approximation for amplitude of this event.

To make more exact calculation that will agree more closely with results of experiment,

must consider other ways this event could happen.

For instance, for each of two main ways the event can happen,

one electron could go charging off to some new and wonderful place and emit photon (Figure below).



Meanwhile, other electron could go to some other place and absorbs photon.

Calculating amplitude for 1st of these new ways involves multiplying amplitudes for :

electron goes from 1 to new and wonderful place,

5 (where emits photon),

and then goes from 5 to 2;

other electron goes from 2 to other place, 6

(where absorbs photon), and then goes from 6 to 4.

Must remember to include amplitude that photon goes from 5 to 6.

Going to write amplitude for this way event could happen in high-class mathematical fashion,

and you can follow along:

$$E(1 \text{ to } 5) * j * E(5 \text{ to } 3) * E(2 \text{ to } 6) * j * E(6 \text{ to } 4) * P(5 \text{ to } 6)$$

a lot of shrinking and turning. (you can easily figure out notation for other case, where electron at 1 ends up at 4, and electron at 2 ends up at 3.)

Initial/Final conditions of experiment for these more complicated ways are same as for simpler ways

- electrons start at points 1 and 2 and end up at points 3 and 4 -

so cannot distinguish between these alternatives and first two - **only measure initial/final states!!!!**

Therefore we must add arrows for these two ways to two ways just previously considered.

But wait: positions 5 and 6 could be anywhere in space and time

yes, anywhere(do not measure them)

and arrows for all of those positions have to be calculated and added together.

See it's getting to be lot of work.

Not that rules are so difficult - like playing checkers:

rules simple, but use them over and over.

So difficulty in calculating comes from having to pile so many arrows together.

That's why takes 4 years of graduate work for students to learn how to do it efficiently

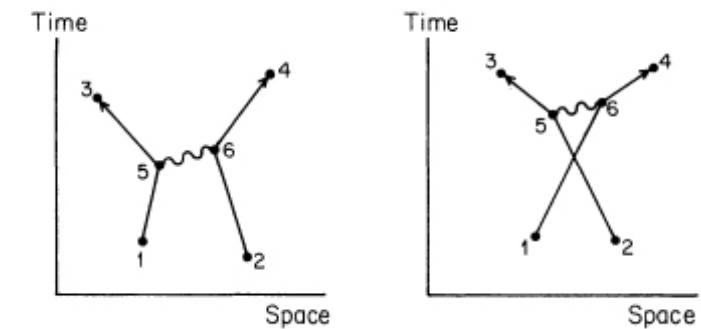
and here we are looking at easy problem!

(when problems get too difficult, now we just put them on a computer!)

I would like to point out something about photons being emitted and absorbed:

if point 6 is later than point 5,

might say that photon emitted at 5 and absorbed at 6 (Figure below).



If 6 earlier than 5,

might prefer to say photon emitted at 6 and absorbed at 5,

but could just as well say that photon going backwards in time!

However, **don't have to worry**

about which way in space-time photon went;

all included in formula for $P(5 \text{ to } 6)$,

and say photon was **“exchanged”**.

Isn't it beautiful how simple Nature is!

Now, in addition to photon that is exchanged between 5 and 6,

another photon could be exchanged - between two points, 7 and 8

(Figure right).

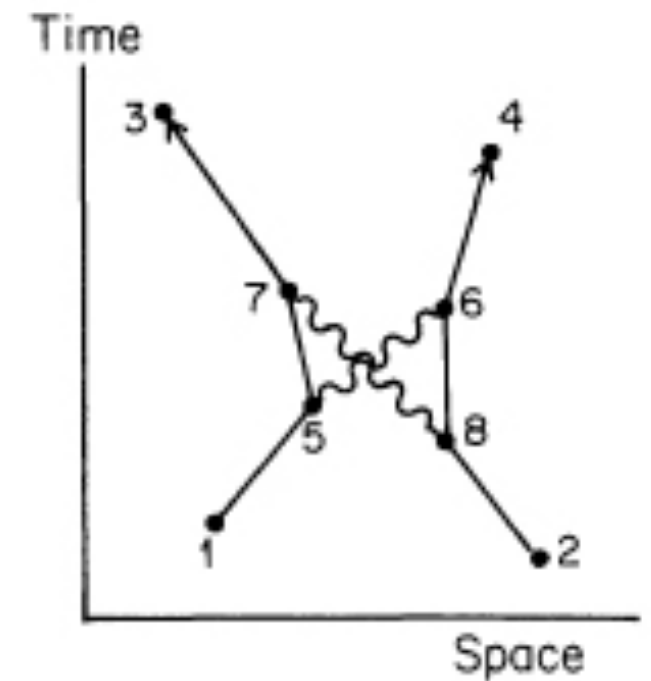
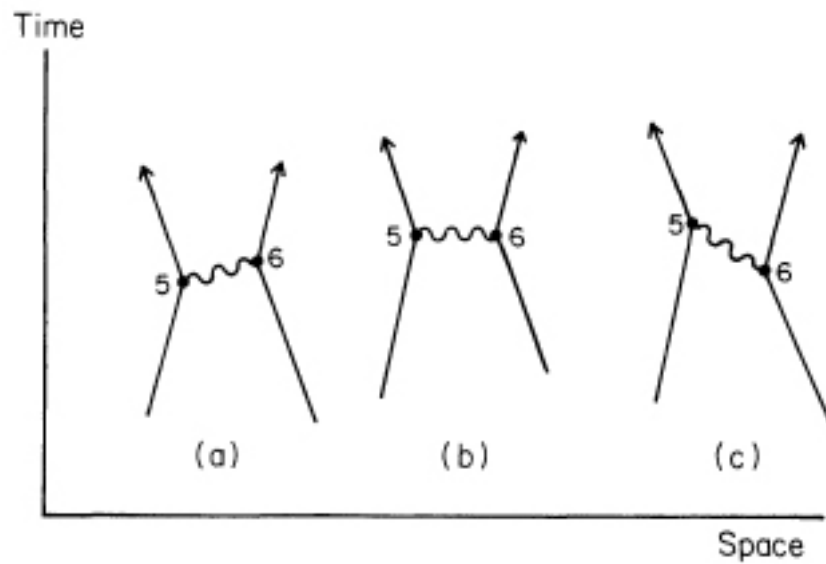
Too messy to write down all basic actions whose arrows have to be multiplied,

but - as may have noticed

every straight line gets an $E(A \text{ to } B)$, every wavy line gets a $P(A \text{ to } B)$, and every coupling gets a j .

Thus, there are six $E(A \text{ to } B)$'s two $P(A \text{ to } B)$'s, and four j 's - for every possible 5,6,7, and 8!

That makes billions of tiny arrows that have to be multiplied and then added together!



Appears that calculating amplitude for simple case is hopeless business,
but when you're graduate student you've got to get your degree, so you keep on going.

But there is hope for success.

It is **found** in that magic number, j ;

1st two ways event could happen have no j 's in calculation;

next way has $j*j$,

and last way we looked at had $j*j*j*j$.

Since $j*j$ is less than 0.01,

it means length of arrow for this way generally less than 1% of arrow for first two ways;

and arrow with $j*j*j*j$ in it less than 1% of 1% - one part in 10,000

compared to arrows that have no j .

If got enough time on computer, can work out possibilities that involve j^6

one part in a million - and match accuracy of experiments.

That's how calculations of simple events are made.

That's way it works;

That's all there is to it! All graduate students laugh at me!!!!!!

Let us look more deeply. In particular, let us look at another event now.

Begin with photon and an electron,

and end with photon and an electron. **Only things we actually measure!!!!!!**

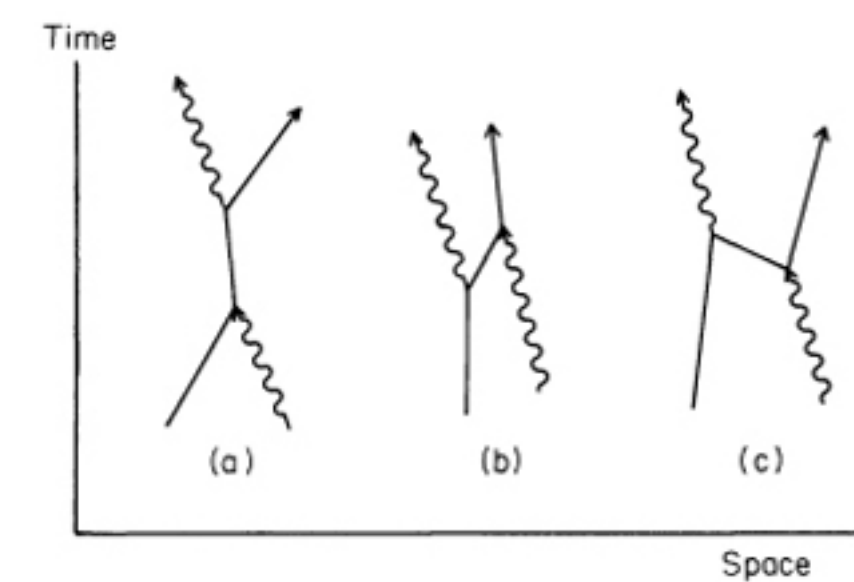
One way for this event to happen is:

photon absorbed by electron, electron continues on a bit, and new photon comes out.

This process is called scattering of light.

When we make diagrams and calculations for scattering,

must include some strange possibilities (Figures below).



For example,

electron could emit photon before absorbing one (b).

Even more strange is possibility (c)

that electron emits photon

then **travels backwards in time** to absorb photon,

and then proceeds forward in time again.

Path of such “backwards-moving” electron can be so long

that it **appears as a real particle** in an actual experiment in laboratory.

Such behavior must be included in these diagrams and the equations for E(A to B).

The backwards-moving electron

when viewed with time moving forward appears to have the same properties as ordinary electron, except it's attracted to normal electrons —> says it has “positive charge”.

(if included effects of polarization, would be apparent why sign of $\mathbf{j}$ for backwards-moving electron appears reversed, making charge appear positive,)

For this reason it is called a “**positron**”.

Positron is sister particle to electron, and is example of an “**anti-particle**”.

This phenomenon is general.

It happens with all particles.

Thus, it seems that every particle in Nature has amplitude to move backwards in time and therefore has an anti-particle.

When particle and its anti-particle collide, they annihilate each other and form other particles.

(For positrons and electrons annihilating, usually a photon or two is produced.)

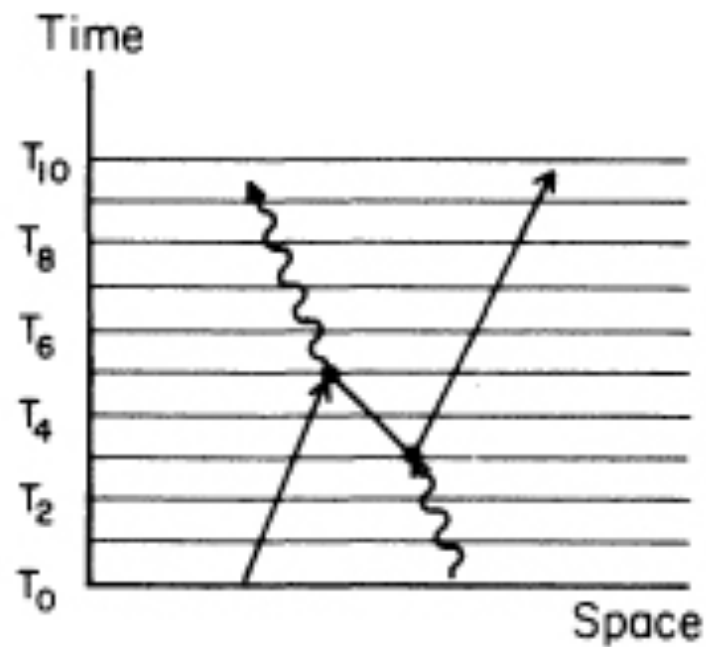
What about photons?

Photons look exactly same in all respects when travel backwards in time - as saw earlier - so they are own anti-particles.

See how clever Nature is at making an exception part of the rule!

Now let me show what a backwards-moving electron looks like to us, as it move forwards in time.

With sequence of parallel lines to aid eye, I am going to divide diagram into blocks of time, T_0 to T_{10} (Figure below).



Start at T_0 with electron moving toward photon,
which is moving in opposite direction.

All of a sudden - at T_3 - photon turns into 2 particles, positron and an electron.

Positron doesn't last very long: soon runs into an electron - at T_5 ,
where annihilates and produces new photon.

Meanwhile, electron created earlier by original photon continues on through space-time.

That is the way to think about it!

Next thing I would like to talk about is an **electron in an atom**.

In order to understand behavior of electrons in atoms,

have to **add** one other feature to our discussion, namely, the nucleus

- heavy part at center of an atom that contains at least one proton
- **(the proton is a “Pandora’s Box” that we will open later).**

I will not give you the correct laws for behavior of nucleus now - they are very **complicated**.

But in this case, where nucleus is quiet, we can **approximate its behavior**

as that of particle with an amplitude to go from one place to another in space-time

according to formula for $E(A \text{ to } B)$, but with a **much higher value** for the number **n** .

Since the nucleus is so heavy compared to electron, we can deal with it **approximately** by saying that stays in essentially one place as it moves through time.

Simplest atom, called **hydrogen**,
is proton and electron.

By exchanging photons,
proton keeps electron nearby, dancing around it (Figure right).

Amplitude for photon exchange is

$$(-j) * P(A-B) * j$$

two couplings

+ amplitude for photon to go from place to place.

Amplitude for proton to have coupling with photon is $-j$ (- because a positive charge).

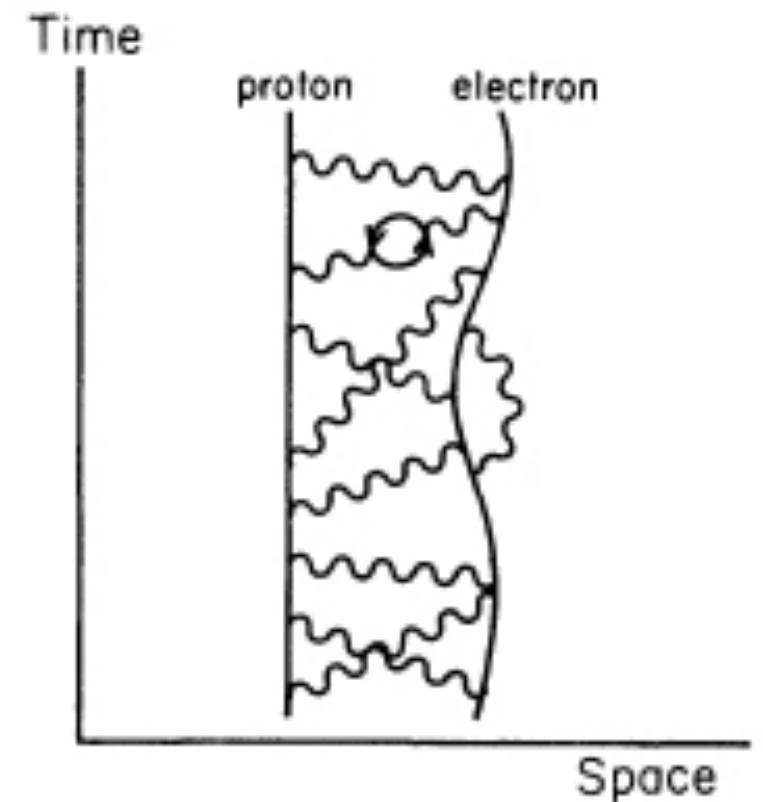
Atoms that contain more than one proton
and corresponding number of electrons
also scatter light

(atoms in air scatter light from sun and make sky blue),

but diagrams for these atoms

would involve so many straight and wiggly lines

-> complete mess!



Now, show you diagram of an electron in hydrogen atom scattering light (Figure right).

As electron and nucleus exchanging photons,

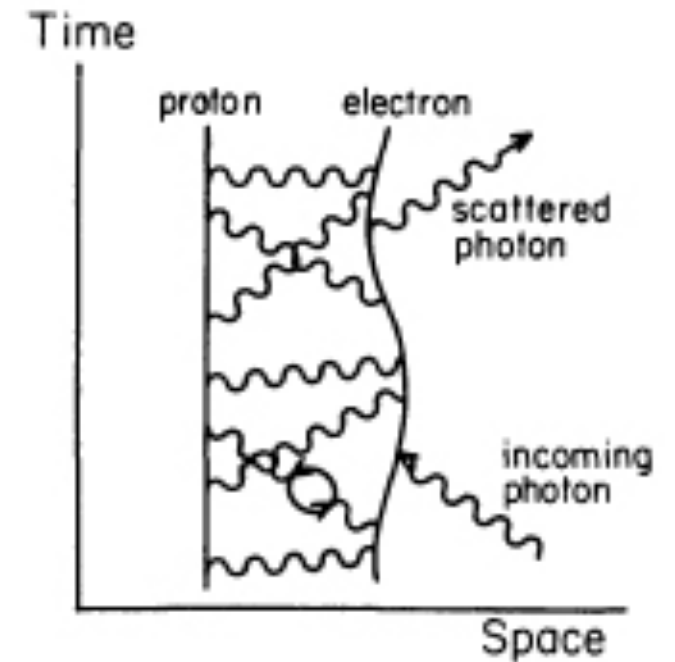
photon comes from outside atom,

hits electron and absorbed;

then new photon emitted.

(As usual, are other possibilities to be considered,

such as new photon emitted before old photon is absorbed.)



Total amplitude for all ways electron can scatter photon can be summed up as single arrow corresponding to a certain amount of shrink and turn.

(Later, we will call this arrow “S”.)

Amount of shrink and turn depends on nucleus and arrangement of electrons in atoms,
and is different for different materials.

Now, look again at partial reflection of light by layer of glass.

How does it actually work?

Talked about light being reflected from front surface and back surface.

This idea of surfaces was **simplification** made in order to keep things easy at the beginning.

Light is **not** really affected by surfaces.

An incoming photon is **scattered** by electrons in atoms inside the glass;

it is absorbed and a **new** photon comes back up to detector.

Interesting that instead of adding up all billions of tiny arrows

that represent amplitude for all electrons inside glass to scatter incoming photon,

we can add just two arrows for “front surface” and “back surface” reflections

and come out with same answer.

Let's see **why**.

To discuss reflection by a layer from new point of view must **take into account** the dimension of time.

Previously, when talked about light from monochromatic source,

used imaginary stopwatch that times photon as it moves -

hand of stopwatch determined angle of amplitude for given path.

In formula for $P(A \text{ to } B)$

(amplitude for photon to go from point to point) there is **no mention** of any timing.

What happened to stopwatch?

What happened to the turning?

Earlier, simply said light source was **monochromatic** without giving any details.

To correctly analyze partial reflection by layer, **need to know more** about a monochromatic light source.

The amplitude for photon to be emitted by source **varies**,
in general, **with time**:

As time goes on, the **angle** of amplitude for photon to be emitted by source changes.

Source of **white** light - **many** colors mixed together

emits photons in **chaotic** manner:

angle of amplitude changes

abruptly and irregularly in fits and starts.

But **when construct** monochromatic source,

are making device that has been **carefully** arranged

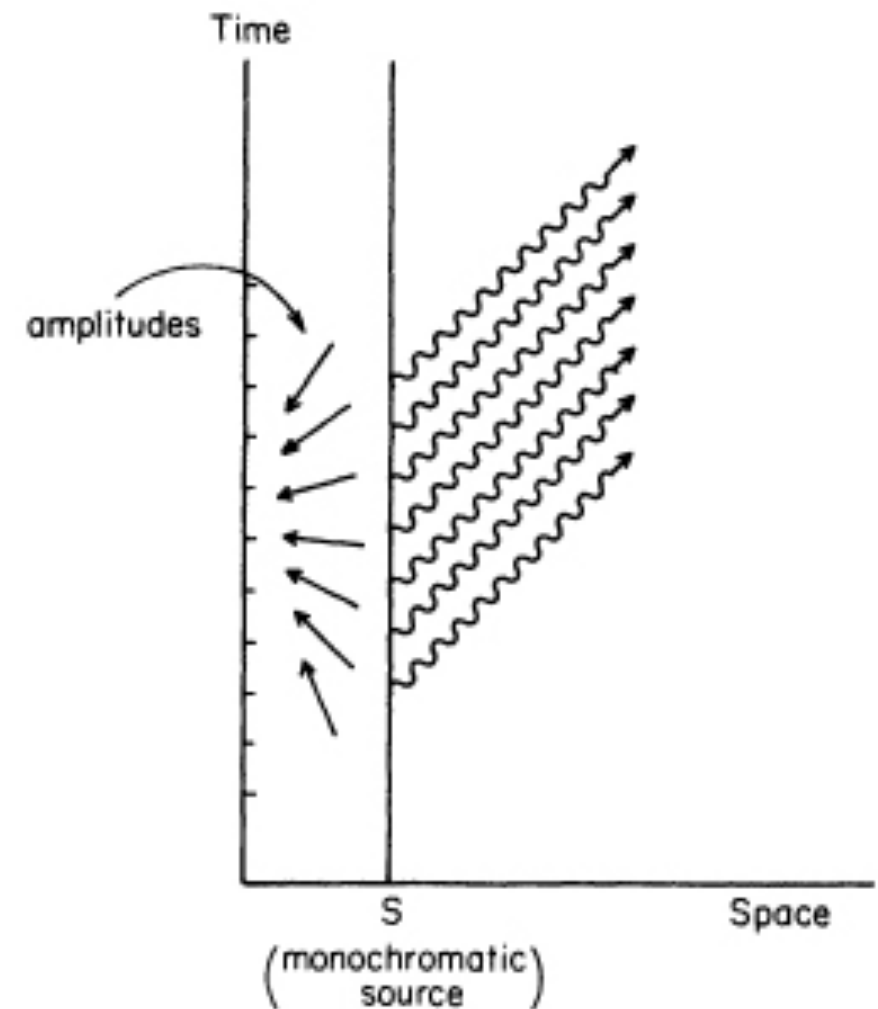
so that **amplitude** for photon

to be emitted at certain time is easily calculated:

it changes its angle at **constant** speed, like stopwatch hand.

(**Actually**, arrow turns at same speed as imaginary stopwatc

used before, but in opposite direction - Figure right).



So, all photons being emitted by a monochromatic source have the same length arrow

and their direction depends on their time of emission from the source;

The direction changes at constant rate

Rate of turning **depends** on color of light:

amplitude for blue source turns nearly **twice** as fast as that for red source, just as before.

So timer used for “imaginary stopwatch” **was** monochromatic source: - in reality,

—> angle of amplitude for given path **depends** on what time photon **emitted** from source,

Once photon has been emitted,

no further turning of arrow as photon goes from one point to another in space-time.

Remember emission time is dependent on where we see photon - farther away takes longer to get there
-> earlier emission -> different angle —> same result as saying stopwatch turns during motion
argument!

Although formula $P(A \text{ to } B)$

says there is amplitude for light to go from one place to another at speeds other than c ,

distance from source to detector in our experiment is relatively **large** (compared to an atom),

so **only** surviving contribution to $P(A \text{ to } B)$'s length(amplitude)

that counts comes from speed c .

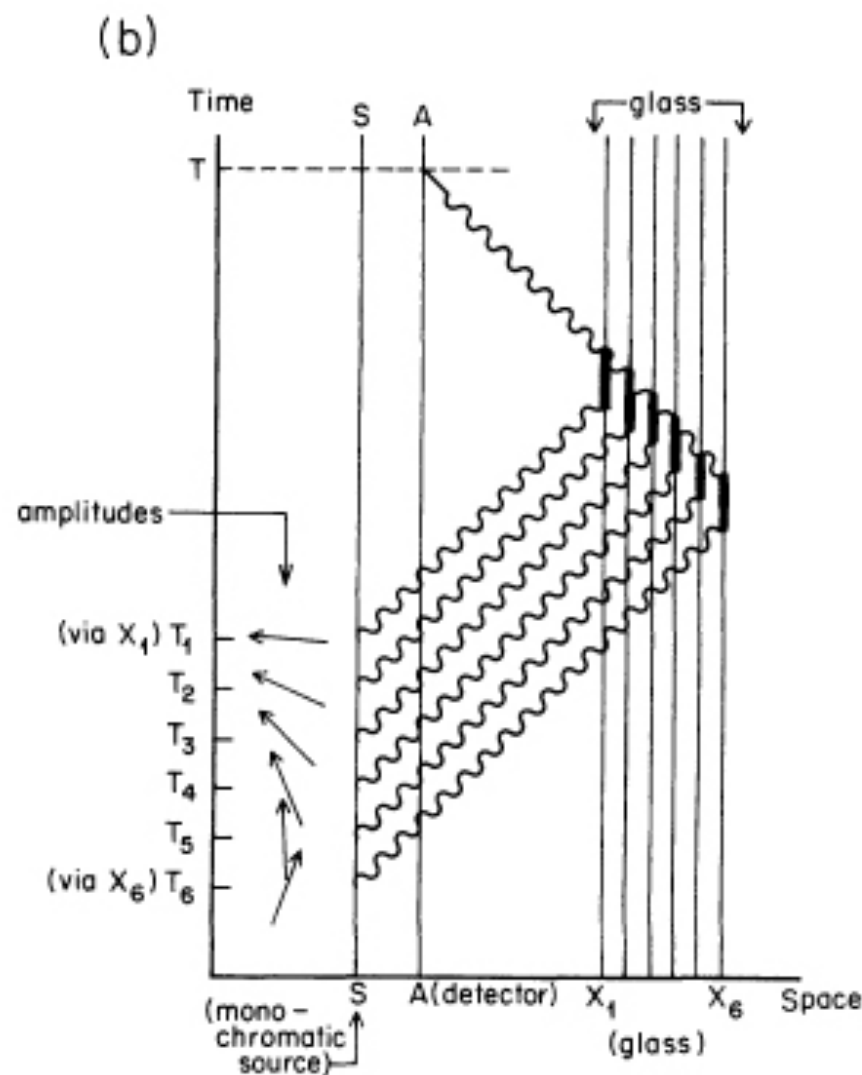
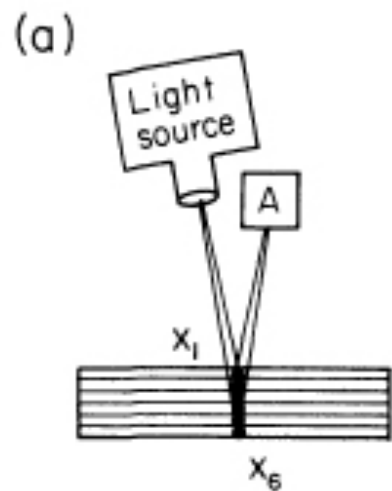
To **begin** new calculation of partial reflection,

start by **defining event completely**:

detector at A makes click at certain time, T.

Then, **divide** layer of glass into number of very thin sections

say, six (Figure below).



Note: directions at emission are different so that all photons can arrive at detector at same time!!!!

From analysis we did earlier

in which **found** that nearly all light was reflected was from middle of mirror,

know that although

each electron is scattering light

in **all** directions,

when all arrows for each section are **added**,

only place where arrows **don't cancel out**

is **where** light goes straight down to

middle section and then

scatters in one of two directions

straight back up to detector

or straight down through glass.

Final arrow for event will thus be **determined**

by adding **six arrows** representing

scattering of light from six middle points

X_1 to X_6

arranged **vertically** throughout glass.

All right,

let's calculate arrows for each of these ways light could go
via six points, X_1 to X_6 .

There are four steps involved in each way

(which means four arrows will be multiplied):

STEP #1: A photon is emitted by the source at a certain time.

STEP #2: The photon goes from the source to one of the points in the glass.

STEP #3: The photon is scattered by an electron at that point.

STEP #4: A new photon makes its way up to the detector.

Say amplitudes for steps 2 and 4 (photon goes to or from point in glass)

involve no shrinking or turning,

because assume that none of light gets lost or spread out between source and glass
or between glass and detector.

For step 3 (electron scatters photon) amplitude for scattering is constant

a shrink and turn by certain amount, S

and is same everywhere in glass.

(amount is, as mentioned before, different for different materials).

For glass, turn of S is 90° .

Therefore, of 4 arrows to be multiplied,

only arrow of step 1

amplitude for photon to be emitted from source at certain time is different from one alternative to next.

Time at which photon would have been emitted to reach detector

A at time T (Figure right(b)) is not same for 6 different paths.

Photon scattered by X_2 would have to have been emitted

slightly earlier than photon scattered by X_1 because that path longer.

Thus arrow at T_2 is turned slightly more than arrow at T_1

because amplitude for monochromatic source

to emit photon at certain time rotates counterclockwise as time goes on.

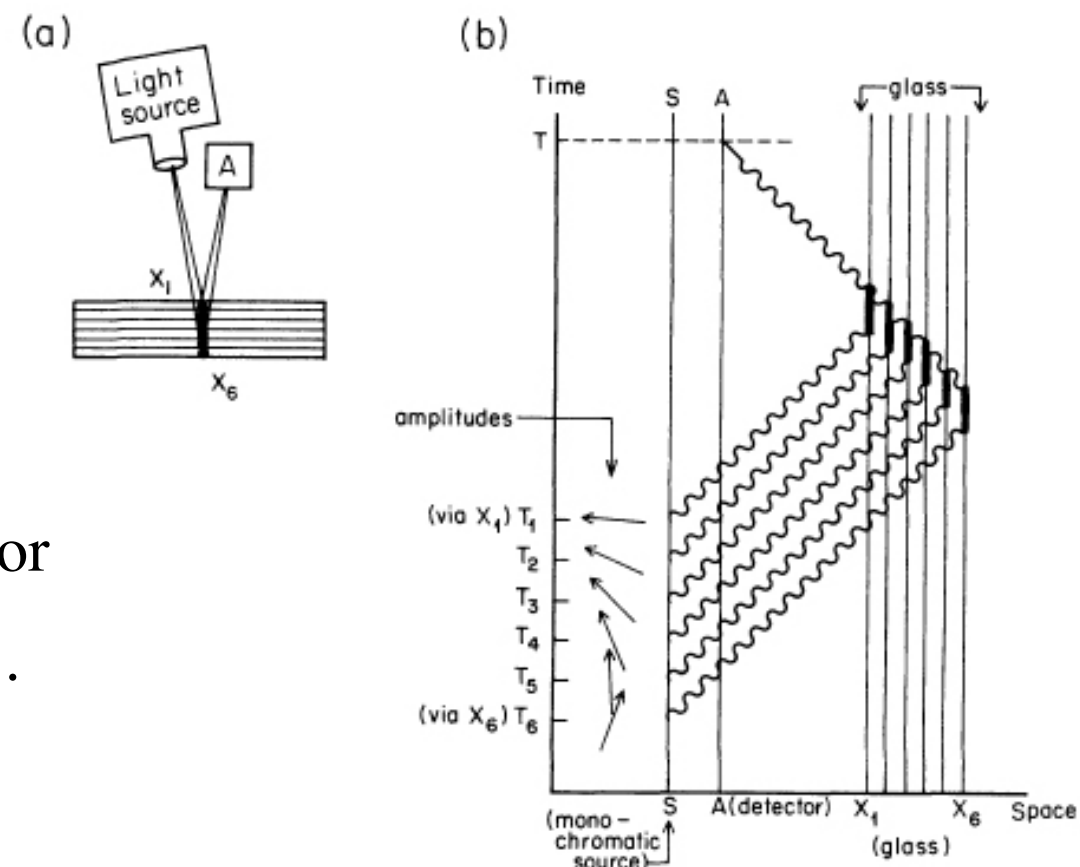
Same goes for each arrow down to T_6 :

all 6 arrows have the same length,

but are turned at different angles

that is, pointing in different directions

because they represent photon emitted by source at **different times**.



After shrinking arrow at T_1 by amounts prescribed in steps 2,3 and 4

turning it 90° prescribed in step 3

end up with arrow 1 (Figure right(c)).

Same goes for arrows 2 through 6.

Thus arrows 1 through 6 are all same (shortened) length, and turned relative to each other in exactly same amount as arrows at T_1 through T_6 .

Next add arrows 1 to 6.

Connecting arrows in order from 1 to 6, get something like an arc, or part of circle.

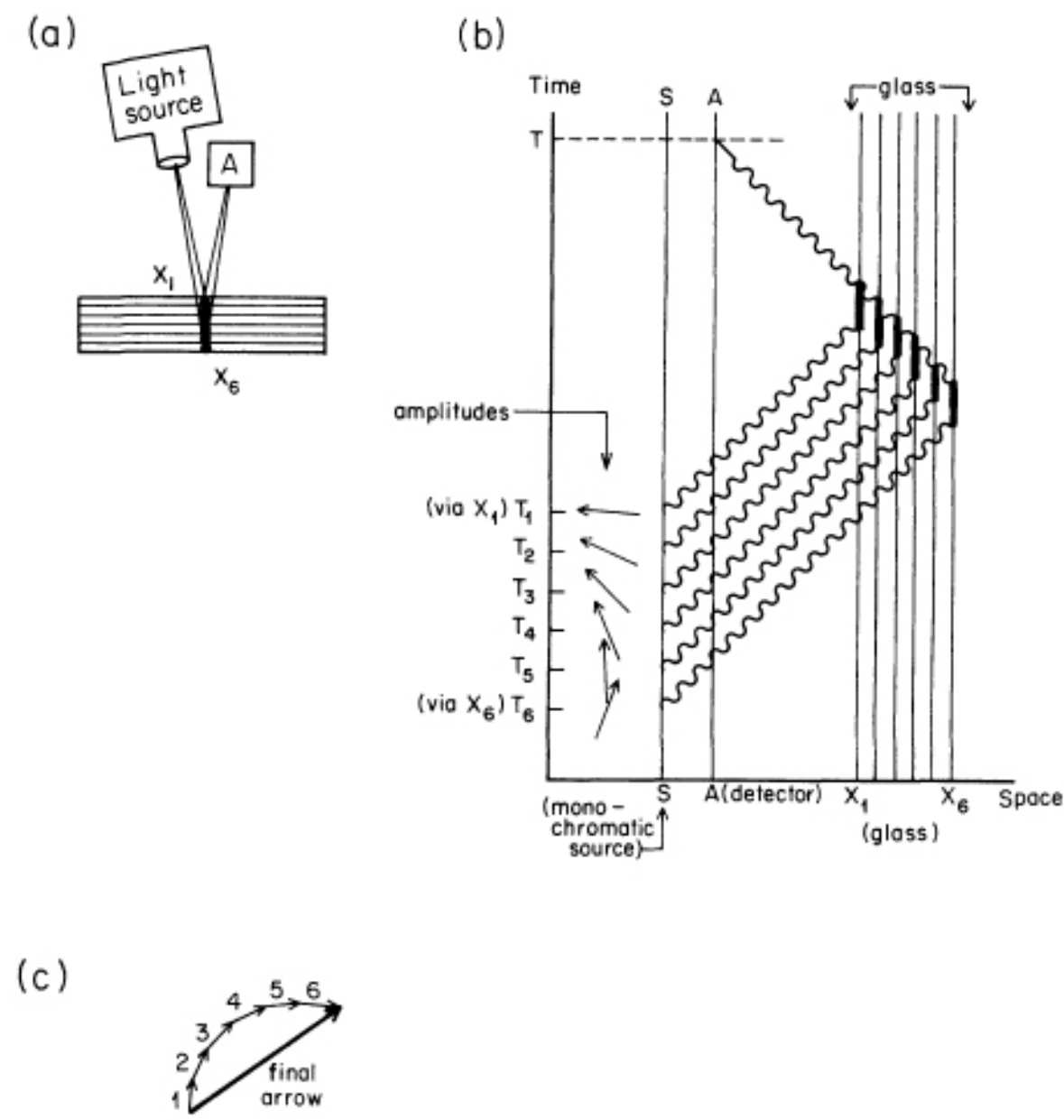
Final arrow forms chord of arc.

Length of final arrow increases with thickness of glass -

thicker glass means more sections, more arrows, and therefore **more** of circle - **until** half circle is reached (and final arrow is diameter).

Then length of final arrow decreases as thickness of glass continues to increase, and circle becomes complete to begin new cycle.

Square of this length is probability of event, and varies in cycle of 0 to 16%.



There is **mathematical trick** can use to get same answer (Figure right(d)):

If draw arrows from center of “circle” to tail of arrow 1 and to head of arrow 6, get two radii.

If radius arrow from center to arrow 1 turned 180° (“subtracted”), then can be combined with other radius arrow to give same final arrow!

That’s what were doing earlier: these two radii are two arrows said represented “front surface” and “back surface” reflections.

Each have famous length of 0.2.

Thus can get correct answer for probability of partial reflection by **imagining** (falsely) that all reflection comes from only front and back surfaces.

In intuitively easy analysis,

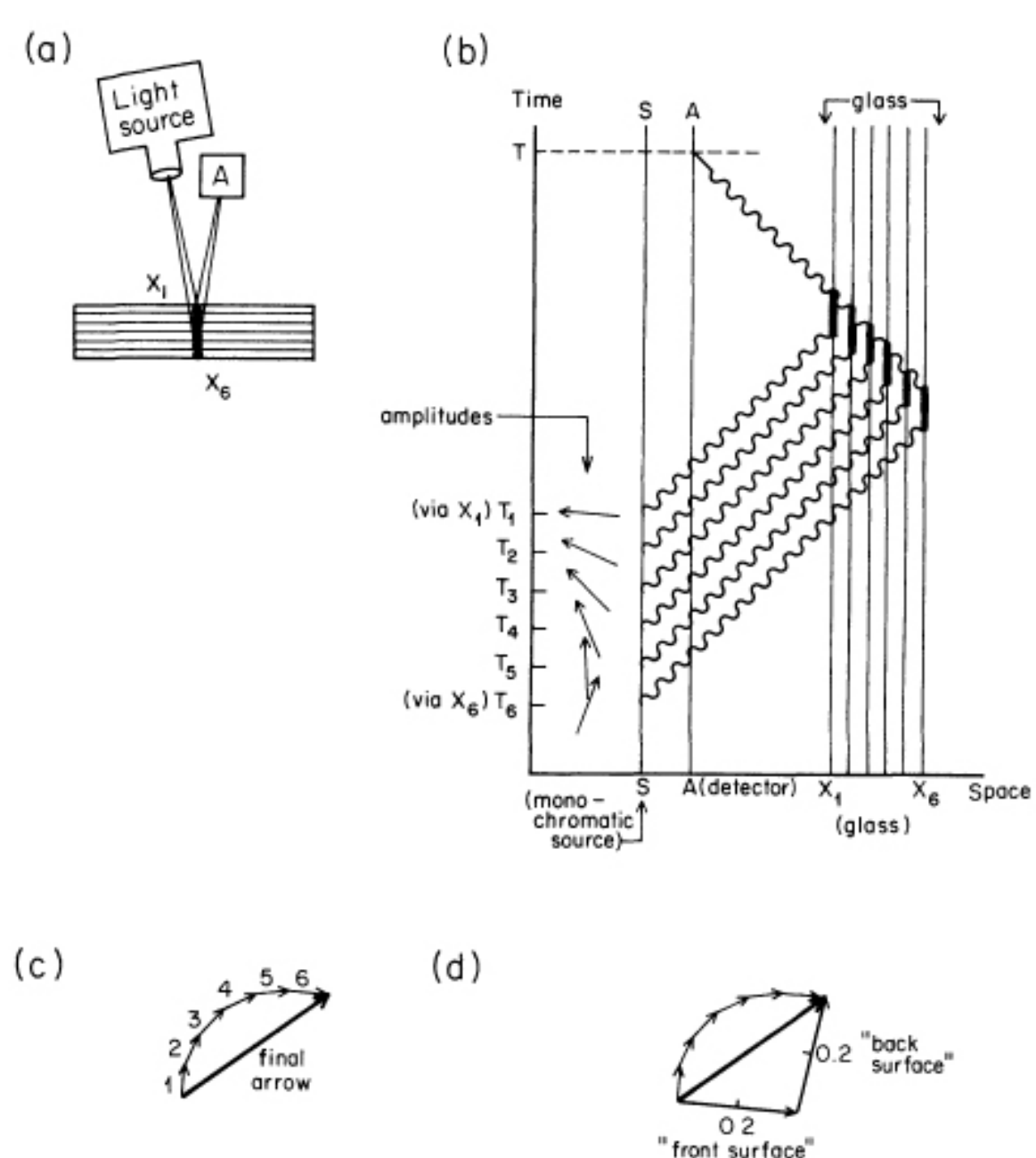
“front surface” and “back surface” arrows are mathematical **constructions** that give right answer, whereas in the analysis we just did -

with space-time drawing and arrows forming part of circle

is more accurate representation of what is really going on:

partial reflection is scattering of light by electrons inside glass.

Now, what about light that goes through layer of glass? Similar discussion - **you read!** Skip, in lecture.



First, is amplitude that photon goes straight through glass without hitting any electrons (Figure right(a)).

This is **most important arrow** in terms of length.

But are 6 other ways photon could reach detector below glass:

photon could hit X_1 and scatter new photon down to B;

photon could hit X_2 and scatter new photon down to B, and so on.

These 6 arrows all have same length as arrows that form “circle” in previous example:

length based on same amplitude of electron in glass to scatter photon, S.

But this time,

all 6 arrows point in **same direction**,

because **length** of all 6 paths that involve one scattering is same.

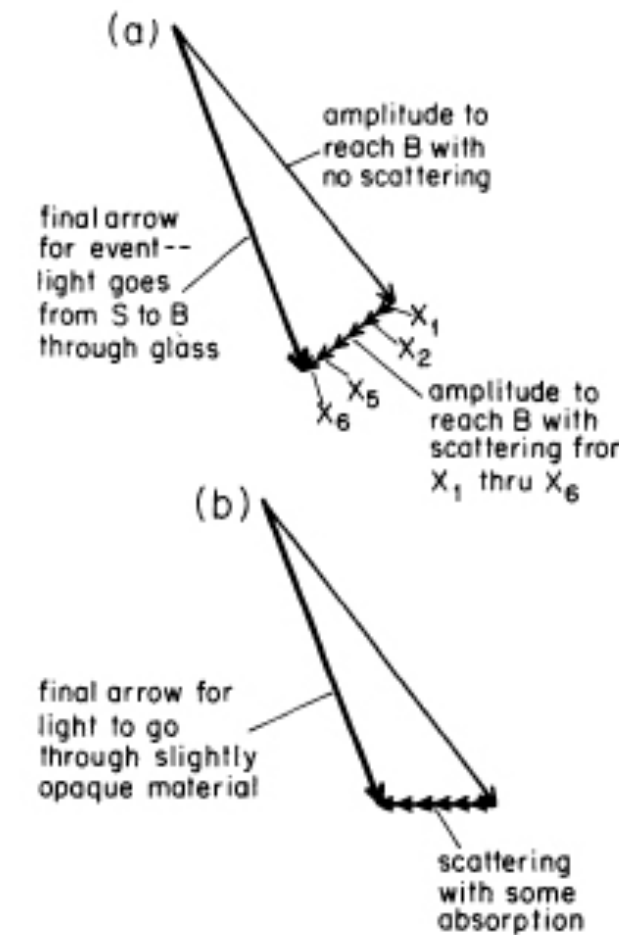
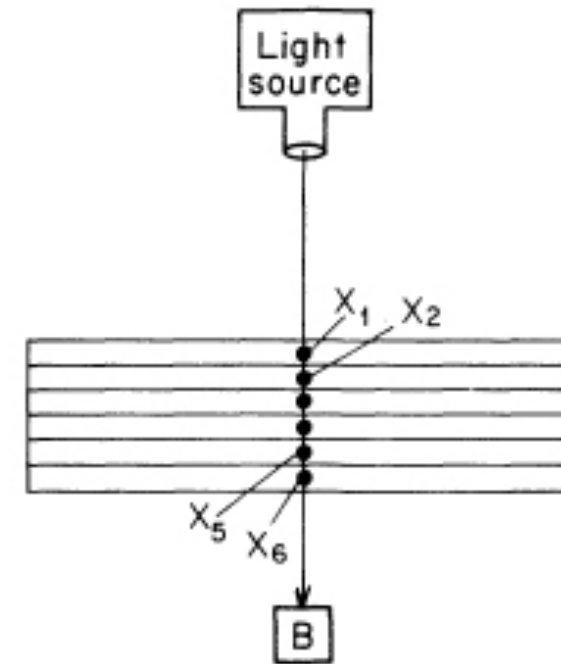
Direction of these minor arrows is

at right angles to main arrow **for** transparent substances such as glass.

When minor arrows are added to main arrow,

result in final arrow that has same length as main arrow,

but is turned in slightly different direction.



SKIP

Thicker the glass,
more minor arrows there are,
and more final arrow turned.

That's how focusing lens **really works**:

final arrows for all paths can be made to point in **same** direction
by inserting extra thicknesses of glass into shorter paths.

Same effect would appear

if photons went slower through glass than through air:
there would extra turning of final arrow.

SKIP

That's why we said earlier

that light appears to go slower in glass (or water) than through air.

In reality, “slowing” of light

is extra turning caused by atoms in glass (or water) scattering light. **Can even stop light!!!**

Degree to which is extra turning of final arrow

as light goes through given material called its “index of refraction”.

For substances that absorb light, minor arrows are at less than right angles to main arrow (Figure above(b)).

This causes final arrow to be shorter than main arrow, indicating that probability of photon going through partially opaque glass smaller than through transparent glass.

Thus **all** phenomena **and** arbitrary numbers mentioned earlier

such as partial reflection with amplitude of 0.2,

“slowing” of light in water and glass, and so on

explained in more detail by just **3** basic actions

three actions that do, in fact, explain nearly everything else, too.

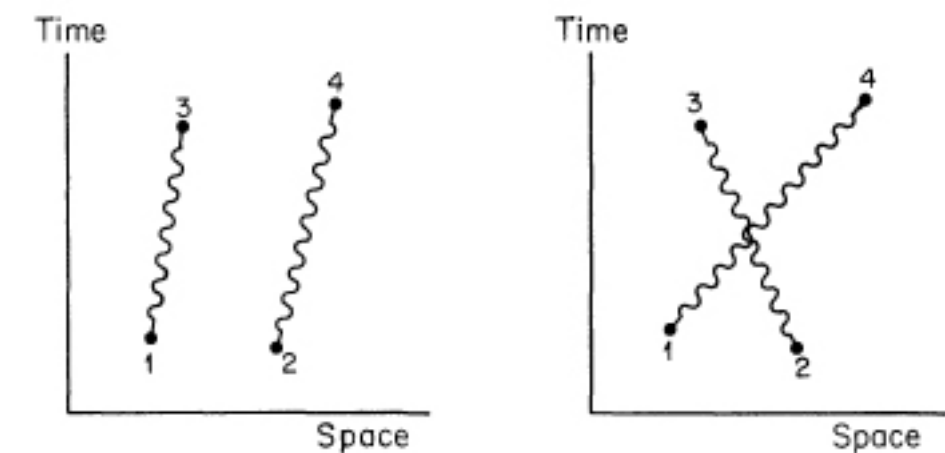
Hard to believe that nearly all vast apparent variety in Nature

results from monotony of repeatedly combining just 3 basic actions.

But it does.

Now outline bit of how some of variety arises.

Start with photons (Figure below).



What is probability that 2 photons,

at points 1 and 2 in space-time,

go to 2 detectors, at points 3 and 4?

Are 2 main ways event could happen

and each depends on 2 things happening concomitantly:

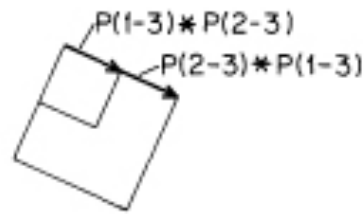
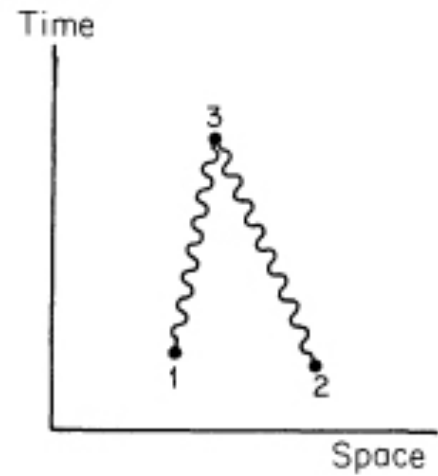
photons could go directly - $P(1 \text{ to } 3) \cdot P(2 \text{ to } 4)$

or could “cross over” - $P(1 \text{ to } 4) \cdot P(2 \text{ to } 3)$.

Resulting amplitudes for 2 possibilities are added, and get interference (as saw earlier),

making final arrow vary in length, depending on relative location of points in space-time.

What if make 3 and 4 same point in space-time (Figure below)?



Let's say both photons end up at point 3,
and see how this affects probability of event.

Now have $P(1 \text{ to } 3) * P(2 \text{ to } 3)$ and $P(2 \text{ to } 3) * P(1 \text{ to } 3)$,
which results in 2 identical arrows.

When added, sum is twice length of either one,

and produces final arrow whose square is 4 times square of either arrow alone.

Because 2 arrows are identical, are always “lined up”.

In other words,

interference doesn't fluctuate according to relative separation between point 1 and 2;
it is always positive.

If didn't think about always positive interference of 2 photons,

would have thought that get twice probability, on average.

Instead, get 4 times probability **all the time**.

When **many** photons are involved,

this more-than-expected probability increases **even further**.

This results in number of practical effects.

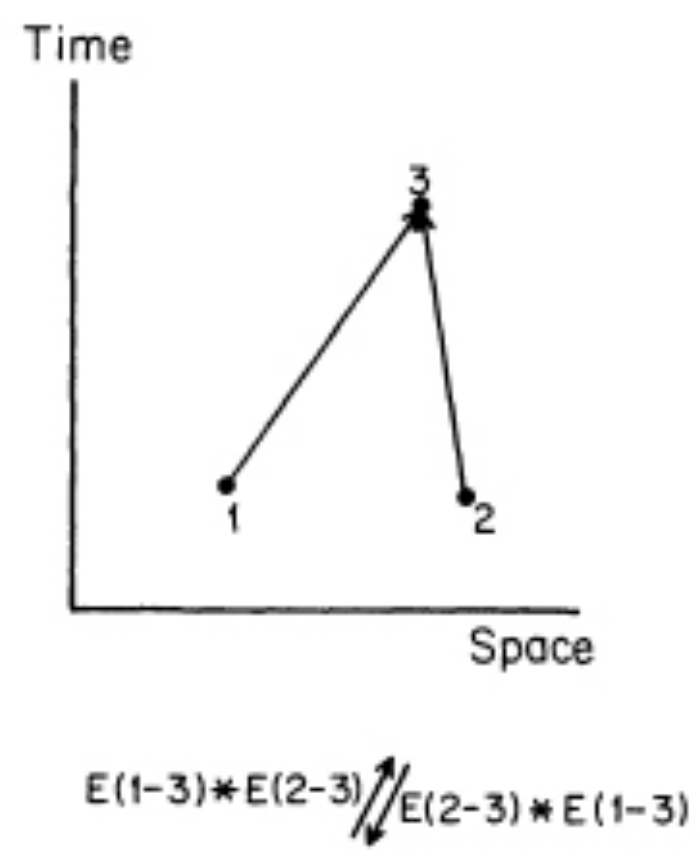
Can say that photons tend to get in same condition,
or “state” (way amplitude to find one varies in space).
Chance that atom emits photon is enhanced if some photons
(in state that atom can emit into) are already present.
This phenomenon of “stimulated emission” discovered by Einstein
when he launched quantum theory proposing photon model of light.
Lasers work on basis of this phenomenon.

If made **same** comparison
with fake, spin-zero electrons,
same thing would happen.

But in real world,
where electrons polarized,
something very different happens:

two arrows $E(1 \text{ to } 3) * E(2 \text{ to } 4)$ and $E(1 \text{ to } 4) * E(2 \text{ to } 3)$,
are subtracted - one of them is turned 180° before are added.

When points 3 and 4 are same,
two arrows have same length and direction
and thus cancel out when are subtracted (Figure right).



That means electrons, unlike photons,

do not like to go to same place; **avoid** each other like plague

no two electrons with same polarization can be at same point in space-time

this is called the “**exclusion principle**”.

Exclusion principle turns out to be origin of great variety of chemical properties of atoms.

One proton exchanging photons with one electron dancing around called a hydrogen atom.

Two protons in same nucleus exchanging photons with two electrons

(polarized in opposite directions) called a helium atom.

Skip next couple of slides in lecture; you read about applications.....

Chemists have complicated way of counting:

instead of saying “one, two, three, four, five protons”,

they say, “hydrogen, helium, lithium, beryllium, boron”.

Are only two states of polarization available to electrons,

so in atom with three protons in nucleus exchanging photons with three electrons -

condition called lithium atom -

third electron is farther away from nucleus than other two

(which have used up nearest available space), and exchanges fewer photons.

Causes electron to easily break away from own nucleus under influence of photons from other atoms.

Large number of such atoms close together

easily lose third electrons to form sea of electrons swimming around from atom to atom.

This sea of electrons reacts to any small electrical force (photons),

generating current of electrons - I am describing lithium metal conducting electricity.

Hydrogen and helium atoms do not lose their electrons to other atoms. -> they are “insulators”.

All atoms - more than one hundred different kinds -

SKIP

made up of certain number of protons exchanging photons with same number of electrons.

Patterns in which they gather complicated and offer enormous variety of properties:

some are metals, some are insulators, some are gases, others are crystals;

there are soft things, hard things, colored things, and transparent things -

terrific cornucopia of variety and excitement that comes from exclusion principle and

repetition again and again and again of three very simple actions $P(A \text{ to } B)$, $E(A \text{ to } B)$, and j .

(If electrons in world were unpolarized,

all atoms would have similar properties:

electrons would all cluster together, close to nucleus of own atom, and

would not be easily attracted to other atoms to make chemical reactions.

You might wonder how such simple actions could produce such a complex world.

It is because phenomena we see in world are result of enormous intertwining of tremendous numbers of photon exchanges and interferences.

Knowing 3 fundamental actions is very small beginning

toward analyzing any real situation,

where there is such a multitude of photon exchanges going on that is impossible to calculate

Experience has to be gained as to which possibilities are more important.

Thus invent such ideas as “index of refraction” or “compressibility” or “valence” to help us calculate

in approximate way when there’s an enormous amount of detail going on underneath.

Analogous to knowing rules of chess - which are fundamental and simple

compared to being able to play chess well,

SKIP

which involves understanding character of each position and nature of various situations

which is much more advanced and difficult.

Branches of physics that deal with questions such as why iron (with 26 protons) magnetic,

while copper (with 29) not, or why one gas transparent and another one not,

are called “solid-state physics”, or “liquid-state physics”, or “honest physics”.

Branch of physics that found 3 simple little actions (easiest part) called “fundamental physics”

Name chosen in order to make other physicists feel uncomfortable!

Most interesting problems - and certainly most practical problems - obviously in solid-state physics.

But as someone said there is nothing so practical as a good theory,

and the theory of quantum electrodynamics is definitely a good theory!

Now back to discussion:

Finally, return to number 1.0011596521865,

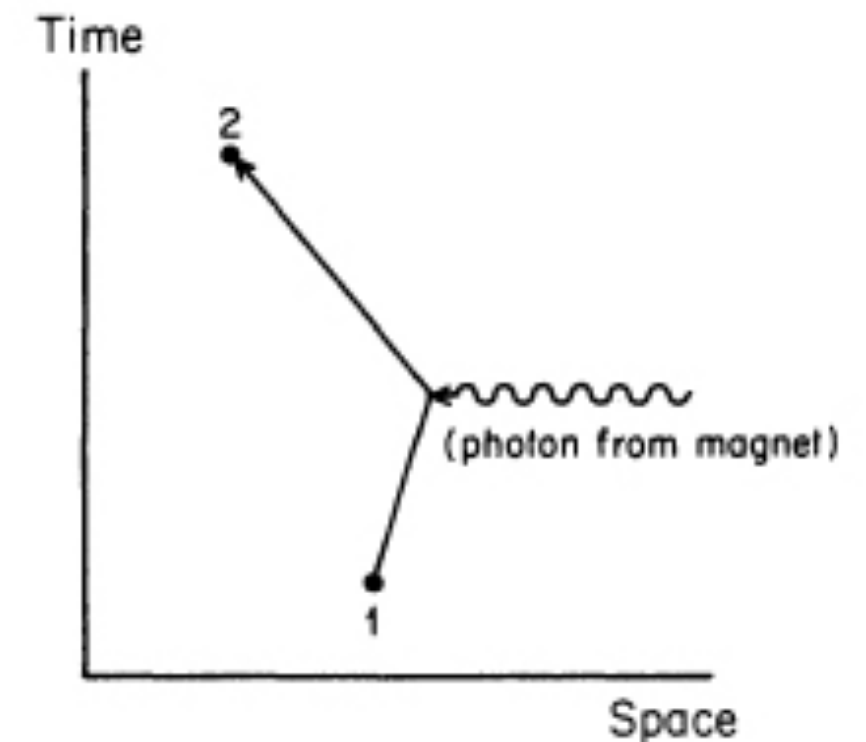
—> number from earlier that has been measured
and calculated so carefully.

Number represents response of electron to external magnetic field
something called “magnetic moment”.

When Dirac first worked out rules to calculate number,
used formula for $E(A \text{ to } B)$ and got very simple answer,
which we will consider in these units = 1.

Diagram for first approximation

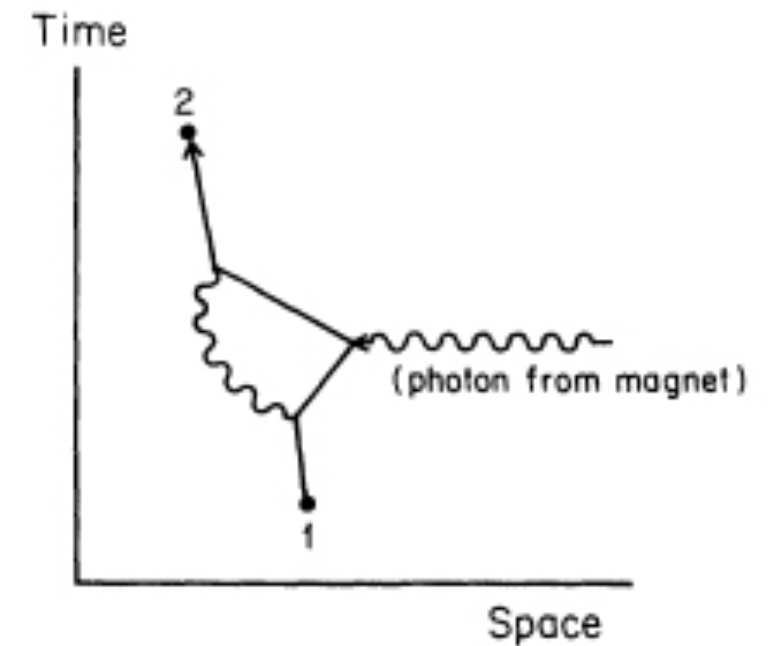
of magnetic moment of electron very simple
an electron goes from place to place in space-time
and couples with photon from magnet (Figure right).



After some years was discovered that this value not exactly 1,
but slightly more - something like 1.00116.

This correction worked out for first time in 1948 by Schwinger
as $j * j$ divided by 2π ,

and due to alternative way electron can go from place to place:
instead of going directly from one point to another,
electron goes along for while and suddenly emits photon;
then later absorbs (horrors) own photon (Figure right).



Perhaps something “immoral” about that, but electron does it!

To calculate arrow for this alternative,

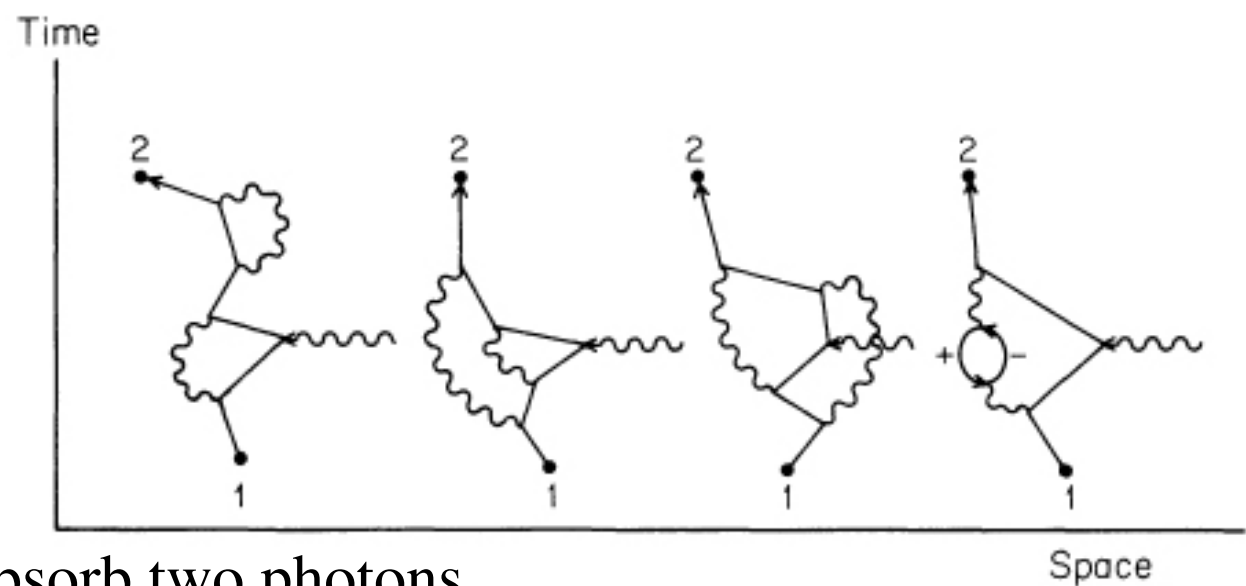
have to make arrow for every place in space-time that photon can be emitted
and every place can be absorbed

-> will be two extra $E(A \text{ to } B)$'s, a $P(A \text{ to } B)$, and two extra j 's, all multiplied together.

Students learn how to do this simple calculation in elementary quantum electrodynamics course,
in second year of graduate school

But wait: experiments have measured behavior of an electron so accurately
that have to consider still other possibilities in our calculations

all ways electron can go from place to place with four extra couplings (Figure below).



There are three ways electron can emit and absorb two photons.

There's also new, interesting possibility (at right Figure to left):

one photon emitted; makes positron-electron pair,
and - again, if hold "moral" objections electron and positron annihilate,
creating new photon that is ultimately absorbed by electron.

That possibility also has to be figured in!

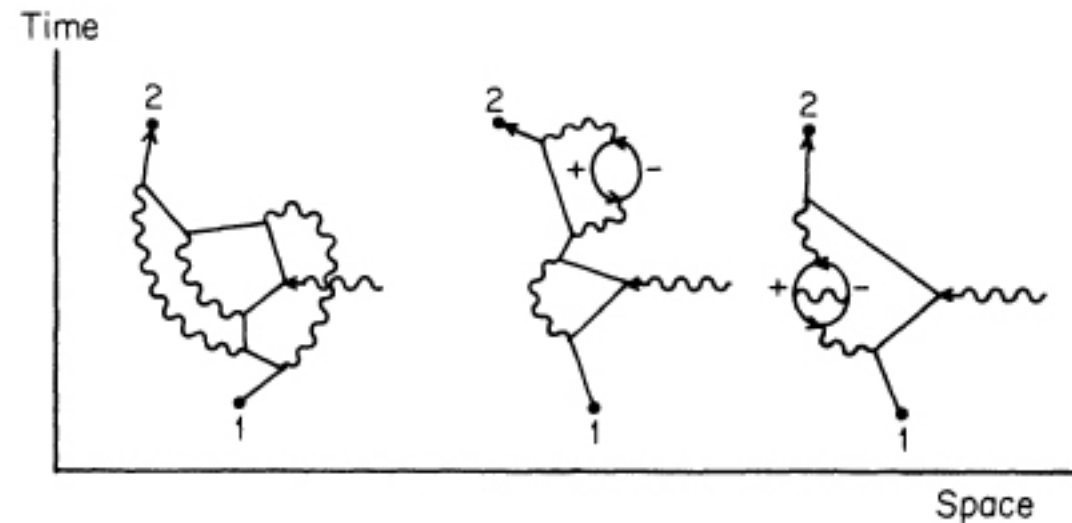
Took two "independent groups of physicist two years to calculate this next term,
and then another year to find out was a mistake
experimenters had measured value to be slightly different,
and it looked for awhile that theory didn't agree with experiment for first time,
but no: was mistake in arithmetic.

How could two groups make same mistake"?

Turns out that near end of calculation two groups compared notes
and ironed out differences between their calculations, so were not really independent.

Term with six extra j's

involves even more possible ways event could happen
draw few of them for you now (Figure below).



Took twenty years

to get extra accuracy figured into theoretical value of magnetic moment of an electron.

Meanwhile experimenters made even more detailed experiments

and added few more digits onto number

and theory still agreed with it.

So,

to make our calculations, make these diagrams,

write down what they correspond to mathematically,

and add amplitudes

—-> straightforward, “cook- book” process.

Therefore, can be done by machines.

Now that have super-duper computers, have begun to compute term with eight extra j's.

At present time

experiments says $1.0011596521873 \pm 0.00000000000028$

theory says $1.0011596521865 \pm 0.00000000000035$

Some of uncertainty in theoretical value is due to computer's rounding off numbers;

most of it due to the fact that value of j not exactly known.

Term for eight extra j's involves something like **900 diagrams**,

with hundred thousand terms each - fantastic calculation - and has been done.

Sure that in few more years,

theoretical and experimental numbers for magnetic moment of electron

will be worked out to still more places.

Of course, not sure whether two values will still agree.

That, one can never tell until one makes calculation and does experiments.

Have come full circle to number chose to “intimidate” you with at beginning.

Hope understand significance of number much better now:

represents extraordinary degree to which

have been constantly checking that strange theory of quantum electrodynamics is indeed correct.

Throughout have I delighted in showing that **price** of gaining such an accurate theory
has been erosion of common sense.

Must **accept** some very bizarre behavior:

amplification and suppression of probabilities,
light reflecting from all parts of a mirror,
light travelling in paths other than a straight line,
photons going faster or slower than conventional speed of light,
electrons going backwards in time,
photons suddenly disintegrating into positron-electron pair, and so on.

That we must do,

in order to appreciate what Nature really doing underneath nearly all phenomena we see in world.

With exception of technical details of polarization,

have described framework by which understand all these phenomena.

Draw amplitudes for every way an event can happen

and add them when would have expected to add probabilities under ordinary circumstances;
multiply amplitudes when would have expected to multiply probabilities.

Thinking of everything in terms of amplitudes

may cause difficulties at first because of their abstraction,
but after while, one gets used to strange language.

Underneath so many of phenomena see every day

are only 3 basic actions:

one is described by simple coupling number, j ;

other two functions - $P(A \text{ to } B)$ and $E(A \text{ to } B)$

both of which closely related.

That's all there is to it, and from it all rest of laws of physics come.

However, before finish, let me make a few **additional** remarks.

Can understand spirit and character of quantum electrodynamics

without including technical detail of polarization.

But feel uncomfortable unless I say something about what I have been leaving out.

Photons, turns out, come in **4** different varieties,

called **polarizations**,

that are **geometrically related** to directions of space and time.

Thus photons are polarized in X, Y, Z, and T directions.

(Perhaps **heard** somewhere that light comes in only two states of polarization

for example, photon going in Z direction can be polarized at right angles,

either in X and Y direction.

“polaroid glasses”

Well, you guessed it:

in situations where photon goes a long distance
and appears to go at speed of light,
amplitudes for Z and T terms exactly **cancel** out.

But for virtual photons going between proton and electron in atom,
it is T component that is most important.

In **similar** manner,

electron can be in one of 4 conditions that also related to geometry,
but in somewhat more **subtle** manner.

Call these conditions 1, 2, 3, and 4.

Calculating amplitude for an electron going from point A to point B in space-time
becomes somewhat **more complicated**,
because can **now ask** questions such as,

**“What is amplitude that electron liberated in condition 2
at point A arrives in condition 3 at the point B?”**

The **16** possible combinations

coming from 4 different conditions an electron can start at A
and 4 different conditions can end up in at B

are **related** in simple mathematical way to formula for that $E(A \text{ to } B)$ told you about.

For photon, **no** such modification necessary.

Thus photon polarized in X direction at A will **still** be polarized in X direction at B,
arriving with amplitude $P(A \text{ to } B)$.

Polarization **produces** large number of different possible couplings.

Could ask, for example,

**“What is amplitude that an electron in condition 2 absorbs photon polarized in X direction
and thereby turns into electron in condition 3?”**

All possible combinations of polarized electrons and photons **do not** couple,

but those that do, do so with same amplitude j ,

but **sometimes** with an additional turn of arrow by some multiple of 90° .

These possibilities for different kinds of polarization and nature of couplings

can all be deduced in very elegant and beautiful manner

from principles of quantum electrodynamics and two further assumptions:

1) results of an experiment are not affected if apparatus

with which are making experiments is turned in some other direction,

and **2)** also doesn't make any difference

if apparatus is in spaceship moving at some arbitrary speed. (-> principle of relativity.)

—-> Rotational Invariance and Lorentz Invariance

This elegant and general analysis

shows that every particle must be in one or another class of possible polarizations, which we call spin 0, spin $1/2$, spin 1, spin $3/2$, spin 2 and so on.

Different classes behave in different ways.

Spin 0 particle is simplest - has just one component - not effectively polarized at all.

(fake electrons and photons that have been considering are spin 0 particles.

So far, no fundamental spin 0 particles have been found.)

Real electron is example of spin $1/2$ particle,

and real photon is example of a spin 1 particle.

Both spin $1/2$ and spin 1 particles have four components.

Other types would have more components, such as spin 2 particles, with ten components.

Connection between relativity and polarization is simple and elegant,

but **can't** explain simply and elegantly! (would take lot of time to do it - another entire class.)

Although details of polarization are not essential

to understanding spirit and character of quantum electrodynamics,

they are, of course, essential to correct calculation of any real process,

and often have profound effects.

Have been **concentrating** on relatively simple interactions

between electrons and photons at very small distances, in which **only** few particles are involved.

Now make one or **two remarks**

about how these interactions **appear in larger world**,

where **very, very** large numbers of photons are being exchanged.

On such large scale, calculation of arrows gets very **complicated**.

There are, however, some situations **not so difficult** to analyze.

Are circumstances, **for example**, where amplitude to emit photon by source

is independent of whether another photon has been emitted.

Can happen when source very heavy (nucleus of atom),

or when very large number of electrons all moving same way, such as up and down

in antenna of broadcasting station or going around in coils of an electromagnet.

Under such circumstances large number of photons are emitted, all **exactly** same kind.

Amplitude of an electron to absorb photon in such an environment

independent of whether it or any other electron has absorbed other photons before.

Therefore entire behavior can be given by just this amplitude for an electron to absorb a photon,

which depends only on electron's position in space and time.

Physicists **use** ordinary words to describe this circumstance.

They say, in this case, the electron is moving in an external field.

Physicists use word “**field**”

to describe any quantity that **depends** on position in space and time.

Temperatures in air provide good example:

they vary according to where and when you make your measurements.

When take polarization into account, are more components to the field.

(are 4 components - corresponding to amplitude to absorb

each of different kinds of polarization (X,Y,Z,T) photon might be in.

Technically called vector and scalar electromagnetic potentials.

From combinations of these, classical physics derives more convenient components called **electric and magnetic fields**.)

In situation where electric and magnetic fields varying slowly **enough**,

amplitude for electron to travel over very long distance **depends** on path it takes.

As saw earlier in case of light,

most important paths are ones where angles of amplitudes from nearby paths are nearly same.

Result is that particle doesn't necessarily go in straight line.

This brings us all the way **back** to classical physics,

which **supposes** that there are fields

and that electrons **move** through them in such a way as to make a certain quantity least.

Physicists call this quantity “**action**” and formulate this rule as “**principle of least action**”.

One example of how rules of quantum electrodynamics produce phenomena on large scale.

Just wanted to **remind you** that effects see on large scale

and strange phenomena see on small scale

are both produced by interaction of electrons and photons,

and **are all described, ultimately,**

by theory of quantum electrodynamics.

Loose Ends

First,

going to talk about **problems** associated with theory of quantum electrodynamics itself,

supposing that all there is in world is electrons and photons.

Then will talk about **relation** of quantum electrodynamics to rest of physics.

Most **shocking characteristic** of theory of quantum electrodynamics

is **crazy framework** of amplitudes,

which **might think** indicates problems of some sort!

However, physicists have been fiddling around with amplitudes

for more than fifty years now,

and have **gotten very used to it.**

Furthermore, **all new particles and new phenomena** that are able to observe
fit perfectly with everything that can be deduced from such a **framework of amplitudes**,
in which probability of an event is square of final arrow
whose length determined by combining arrows in funny ways (with interferences, and so on).

So framework of amplitudes has **no experimental doubt** about it:

can have all philosophical worries you want as to what amplitudes mean
(**if**, indeed, mean anything at all),
but because physics is **experimental** science and framework agrees with experiment,
—> **good enough for us so far.**

There is **set of problems** associated with theory of quantum electrodynamics
that has to do with **improving** method of calculating sum of all little arrows -
various techniques that are available in different circumstances -
that take graduate students three or four years to master.

Since they are technical problems, I am not going to discuss them in detail here,

It is **just a matter** of continuously improving techniques for analyzing
what theory really has to say in different circumstances.

But there is **one** additional problem that is characteristic of theory of quantum electrodynamics itself,
which **took 20 years to overcome. Let me mention it now and say much more later!**

Has to do with our model of ideal electrons and photons and numbers n and j .

If electrons **ideal**,

and went from point to point in space-time **only** by direct path

(left diagram in Figure right),

then would be **no problem**:

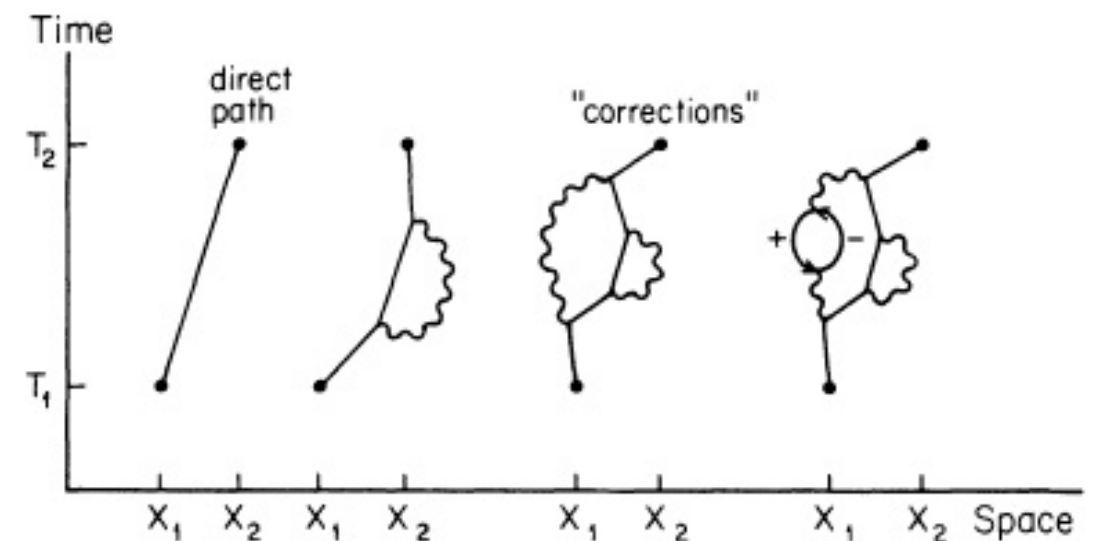
n would simply be mass of electron

(which we can **determine by observation**),

and **j** would simply be its “charge”

(amplitude for electron to couple with photon).

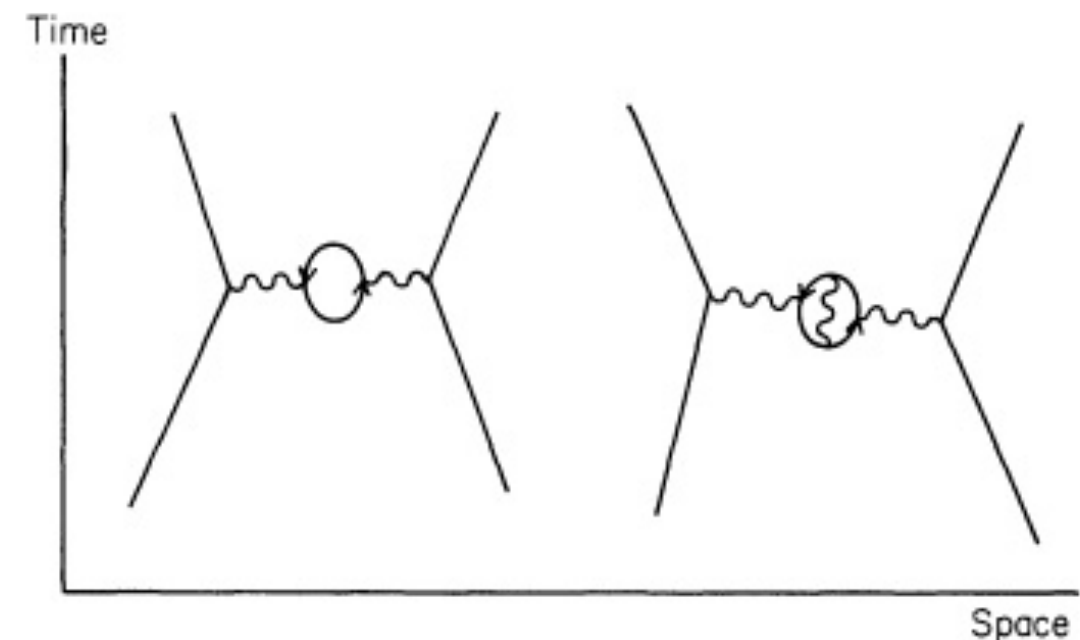
Also can also be determined by experiment.



But no such ideal electrons exist.

Mass observed in laboratory is that of **real** electron, which emits and absorbs **own** photons from time to time, and **therefore** depends on E(A to B),

which **involves** n (Figure right) - mass = f(n).



Experimentally measured amplitude for electron to couple with photon,
the mysterious number, e ,
is number determined by experiments that **include** all “corrections” **“include” —> real electrons**
for photon going from point to point in space-time,
of which two shown on last slide.

When calculating,
need number, j , that **does not** include these corrections,
but includes **only** photon going directly from point to point.

Difficulty exists with computing this j that is **similar** to difficulty in computing value of n .

Since the mass and charge of an electron
are **affected** by these and all other alternatives,
experimentally measured mass, m ,
and experimentally measured charge e , of electron are different
from numbers use in our calculations, n and j

If were definite mathematical connection between n and j on one hand,
and m and e on other, would still be **no** problem:
would **simply calculate** what values of n and j need to start with
in order to **end up** with observed values, m and e .

(if calculations didn't agree with m and e , would **jiggle original n and j around until they did.**)

Let's see how one actually calculates m.

Write a series of terms that is something like series we saw for magnetic moment of electron:

first term has no couplings - just $E(A \text{ to } B)$ -

and **represents ideal** electron going directly from point to point in space-times

Second term has two couplings and represents photon being emitted and absorbed.

Then come terms with four, six, and eight couplings, and so on

(some of these “corrections” were shown earlier).

When calculating terms with couplings,

must consider (as always) **all** possible points where couplings could occur,

right down to cases where two coupling points are on top of each other

with zero distance between them.

Problem is, when try to calculate all way down to **zero distance**,

equation blows up in our face and gives meaningless answers - things like **infinity**.

This **caused lot of trouble** when theory of quantum electrodynamics first came out.

People **getting** infinity for every problem they tried to calculate!

(One should be able to go down to zero distance in order to be mathematically consistent,

but that's where there is no n or j that makes any calculational sense;

that's where trouble is.)

Instead of including all possible coupling points down to distance of zero,
if one **stops** calculation when distance between coupling points very small
say, 10^{-30} centimeters,
billions and billions of times smaller than anything **observable** in experiment
(presently 10^{-18} centimeters) -
then there are definite values for n and j that we can use
so that calculated mass comes out to **match** the m observed in experiments,
and calculated charges **matches** the observed charge, e.

Now, **here's catch**: if somebody else comes along and stops their calculation at **different** distance
say, 10^{-40} centimeters - their **value** for n and j needed to get same m and e come out **different!**

20 years later,(1949), Hans Bethe and Victor Weisskopf **noticed** something:

if 2 people **who stopped** at different distances to determine n and j from same m and e
then calculated answer to another problem
each using appropriate but different values of n and j
when all the arrows from all terms were included,
their answers to this other problem came out nearly the **same!**

In fact, closer to zero distance that calculations for n and j were stopped,
better final answers for other problem would agree!

Schwinger, Tomonaga, and Feynman independently **invented** ways

to make definite calculations to confirm that it is true (got prizes for that).

People could finally calculate with theory of quantum electrodynamics!

So **appears** that only things that depend on small distances

between coupling points are values for n and j

= **theoretical numbers** that are not directly observable anyway

everything else, which can be observed, seems not to be affected.

Shell game that must play to find n and j is technically called “renormalization”.

But no matter how clever the word, is what I would call a **dippy** process!

Having to resort to such **hocus-pocus**

has **prevented** us from proving that theory of quantum electrodynamics

is mathematically self-consistent.

Surprising that theory still hasn't been proved self-consistent one way or other by now;

—> **suspect** renormalization not mathematically legitimate.

What is **certain** is that **do not have** good mathematical way

to describe theory of quantum electrodynamics:

such a bunch of words to describe the connection

between n and j and m and e **is not** good mathematics. (**at time of Feynman!**)

Another interesting way of describing this difficulty

is to say that perhaps **idea that two points can be infinitely close together is wrong**
or assumption that can use geometry down to last notch is false.

If **make minimum possible distance** between two points as small as 10^{-100} centimeters

(smallest distance involved in any experiment today around 10^{-16} centimeters),

infinities disappear, all right - but other inconsistencies arise,

such as **total probability** of an event adds up to slightly more or less than 100%,

or get **negative energies** in infinitesimal amounts.

Been **suggested** that these inconsistencies arise

because **haven't taken into account** effects of gravity

which are normally very, very weak,

but become important at distances of 10^{-33} cm.

Most profound and beautiful question

associated with observed coupling constant, e

—> **the amplitude for real electron to emit or absorb a real photon.**

Is **simple number** that has been experimentally determined to be close to - 0.08542455.

(physicist friends won't recognize number,

because like to remember it as inverse of its square:

about 137.03597 with uncertainty of about 2 in last decimal place.

Has been **mystery** ever since was discovered more than 50 years ago,

and all good theoretical physicists put this number up on their walls and worry about it.)

Immediately would **like to know** where number for coupling comes from:

is it related to π ,

or **perhaps** to base of natural logarithms?

Nobody knows.

It's one of **greatest damn mysteries of physics:**

magic number that comes to us with **no** understanding by humans.

Might say "hand of God" wrote number ,

and " we don't know how She pushed Her pencil".

Know what kind of a dance to do experimentally to measure number very accurately,

but don't know what kind of dance to do on computer to make this number come out

without putting it in secretly!

Good theory **would say** e is square root of 3 over 2π squared, or something.

Have been, from time to time, suggestions what e is, but none of them useful.

First, Arthur Eddington proved by pure logic that number physicists like had to be exactly 136,
experimental number at that time.

Then, as more accurate experiments showed number closer to 137,
Eddington discovered slight error in earlier argument,
and showed by pure logic again that number had to be integer 137!

Every once in a while, someone notices that certain combination of pi's and e's (base of natural logarithms), and 2's and 5's produces mysterious coupling constant,
but it is fact not fully appreciated by people who play with arithmetic that would be surprised how
many numbers can make out of pi's and e's and so on.

Therefore, throughout history of modern physics,
has been paper after paper by people who have produced an e to several decimal places,
only to have next round of improved experiments disagree with it.

Even though have to resort to dippy process to calculate j today,
possible that someday legitimate mathematical connection between j and e will be found.

Would mean that j is mysterious number, and from it comes e.

In such case **would doubtless** be another batch of papers that tell us
how to calculate j “with our bare hands”,
so to speak, proposing that j is 1 divided by 4π , or something.

That exposes all problems associated with quantum electrodynamics. Deal with later classes!

There is **no** theory that adequately explains these numbers.

Use numbers in all our theories, but **don't** understand them
what they are, or where they come from.

Believe that from fundamental point of view, this is very interesting and serious problem.

Lots of speculation about new particles is confusing,

but complete discussion of rest of physics to show how character of those laws
framework of amplitudes, diagrams that represent interactions to be calculated,
and so on

appears to be same as for theory of quantum electrodynamics
our best example of good theory.