

[SQUEAKING]

[RUSTLING]

[CLICKING]

JOSH Last time, we started talking about vision and the eye and the retina. So the idea is objects in the world reflect
MCDERMOTT: light. The photons are absorbed by the eye. And from the pattern of photons that are absorbed by the eye, your job is to figure out what's out there in the world.

So we talked about the eye and the retina and how the retina is wired up backwards and so how that has the consequence that you have a blind spot where the optic nerve exits the eye and that there are these other adaptations to deal with the fact that it's backwards, like the cell body is being pushed out of the way. We observed our blind spot. We talked about rods and cones, how they have different distributions over the eye. So the cones are very dense in the fovea, the center of gaze, whereas there are no rods, and then the cones drop off precipitously after that.

We talked about eccentricity, distance from the fovea. We talked about how the rods and the cones differ in their sensitivity to light. And so under high light conditions, photopic conditions, your vision is dominated by your cones. If you walk into a dark room or a movie theater, you gradually adapt to the low light levels. And what that means is there's an initial stage of adaptation within the cones, and then your vision transitions to being determined by the rods. And the rods eventually become extremely sensitive to light.

So we talked about differences between rod and cone vision, how rod vision is more sensitive than cone vision, how it has lower acuity, and how you're not able to see color when you are dependent on rods because there's only one type of rod. Yeah?

AUDIENCE: I might be misunderstanding, but can you go back? So in that one, why is the threshold of rods higher than the cones for the light, this kind of light?

JOSH Yeah, so under high light conditions, so photopic conditions, the rods are essentially not active. So that's why the
MCDERMOTT: sensitivity of the rods is very poor. And what happens when you enter low light levels is that there's a biochemical process inside the rod photoreceptor that restores that sensitivity. And it builds up over time. And so their sensitivity increases. And eventually, there's a transition where your threshold becomes determined by the rod rather than the cone. Does that make sense?

And so we talked about how one of the things to remember about the cones is that there are three different types of them that have different spectral sensitivities. And that's critical to color vision, which we'll get into a little bit later. And then we talked about the ganglion cells and how the ganglion cells come in a couple different flavors. So ganglion cells, remember, are like the output stage of the retina. So there's a few different stages of neurons, starting with the photoreceptors, culminating in the ganglion cells. And the ganglion cells then send their axons out through the optic nerve to the thalamus and then the cortex.

But there are a couple different types of ganglion cells, midget and parasol cells. They differ in a whole bunch of different ways, one of which is the receptive field size. And we talked about this other kind of critical property of receptive field size, which is that receptive fields tend to be small in the fovea and then get bigger as you move out to the periphery. So this is a major organizing principle within the visual system.

But you can see that if you look at the ganglion cells, there are these two classes of receptive field size. And those correspond to the midget and the parasol. So the parasol cells have larger receptive fields. But overall, there is an increase in the receptive field size as you move out to the periphery, and that means a decrease in your spatial resolution.

So spatial resolution is worse out in the periphery than it is in the fovea. And you're not normally very aware of this because typically anytime you want to see something, you look at it. And that means that you move your fovea to the thing that you want to see. But if you force yourself to fixate, you'll be able to detect this. And it's very measurable. So this is an eye chart that's intended to compensate for the eccentricity related changes in resolution.

And so then we ended last time posing this question of why it is that vision is organized in this way. So why is it that we've evolved eyes that have really good resolution at the center of gaze? And the intuition and the proposal that was tested in this particular paper that we started to talk about is that if there is a limit on the number of spatial samples that you can make-- and that limit might be due to the size of the optic nerve-- then you have some finite number of receptors that you could kind of place in the eye.

And there's different ways in which you could arrange them. But one way to do it is to put a very high density of receptors in one particular part of the eye and then enable the system to move the eye around so that it can do this dense sampling at a particular region of the image and then select where that is.

And so this was a paper that was attempting to investigate whether that's actually a good solution to the problem of vision. And so what they did in this paper is set up a machine system that had a receptor lattice the parameters of which could be optimized. So there's a lattice of receptors that are defined by the positions, μ , and the size of the receptive field of the receptor. That's σ . And this system was trained to perform a visual search task. So there were images like this. And the task was to find and identify the digit on a cluttered background.

And so what they found is that when they optimized the parameters of the lattice in order to enable the system to solve this problem, over the course of the optimization, you can see that the receptor positions and sizes kind of arrange themselves to resemble the organization that you see in the retina, which is to say there is a center of gaze, where the receptors are fairly densely spaced with small receptive fields, and then a periphery, where the spacing kind of gets bigger and the receptive fields get broader.

So that's kind of where we wrapped up. And so interestingly, in this kind of artificial setting, they could kind of simulate problem constraints that don't really exist in biology. And so one of those problem constraints was enabling the system to have a zoom function in addition to simply being able to translate the receptor lattice around.

And when the system was allowed to zoom, you get a different kind of solution emerging. So if it can zoom, then it turns out to be better to just have the receptors more or less uniformly spaced. So that's kind of like a CCD imaging device like is in your phone. But when you can translate only, then you end up with these foveal-like organizations.

And so what these graphs are plotting is the sampling interval or the spacing between the receptors as a function of the distance from what they define to be the center. And you can see that the spacing kind of increases. And these are two different data sets that they trained it on, whereas in the situation where the system can zoom, you see a much weaker dependence. And then this is the receptive field size as a function of eccentricity.

So the other thing that they show in this paper is how the models end up using their retinas when they're performing this task. So remember, this is like a visual search task, where you have to find and recognize the digit. And so these are different time intervals in a particular trial. So this is the stimulus. And you can see that the model is essentially moving its receptor lattice so that it's located over the digit. And in the case where it's allowed to translate, you get the fovea positioned over the digit, whereas when it's allowed to zoom, it hones in on the digit in a different way.

So the inference from this is that a fovea like what we have in our eyes is a useful way to attain high resolution vision given, one, a limit on the number of samples that can be transmitted to the brain, in this case from the optic nerve, and two, the ability to move the eyes to position the fovea at different locations. So the idea is that our visual system has kind of arrived at a good solution to the problem of getting high resolution without kind of having this enormous bandwidth that would be difficult to pass on to the brain. Any questions about that?

And so the bigger picture here, again, is the current era that we're in is giving us the opportunity to answer some of these questions about why our sensory systems are the way they are by taking in silico systems and being able to optimize them for problems and then looking at what kinds of solutions emerge as a consequence of that.

AUDIENCE: So basically, in the bottom left in examples for those two data sets, there's different placements of the fovea. And I'm curious how much we can generalize optimization machine learning experiments to the real world. It's so dependent on the type of data that you feed it.

JOSH MCDERMOTT: Yeah, I think it's a good question. I'm not actually sure that this is data set dependent. It could just be that you'd get similar variants with different random initializations. You'd have to run the experiments multiple times. But yeah, it is showing you that you can get-- there's some variability that is emerging. On the other hand, both of them kind of have a fovea. It's just like one of them's a little bit lopsided for whatever reason. So yeah, you'd have to do some more experiments, I think, to understand that. But I think the point is you could understand that by looking in more detail at the conditions that give rise to these things.

So these ganglion cells and other cells in the retina come in-- there's another set of flavors that they come in. And the one that we're going to talk about now, it has to do with center-surround receptive field organization and the two different types of that. And so this was a very early discovery about by the retina credited to Kuffler in the 1950s, who recorded from neurons in the retina, from ganglion cells, and was investigating the patterns of light that would cause the neurons to fire.

And what he discovered was evidence that you can think of the neuron as having this center-surround receptive field, where there is a center which is defined by the fact that if you stimulate that with light, that will increase the response of the cell. And a surround, which is defined by the fact that if you stimulate that with light, that will tend to decrease the response of the cell.

And so the consequence of that is that the best stimulus for the neuron in this case is going to be this one here, where you're just stimulating the center with light. And if you present an annulus, such that you're just hitting the surround, you get no response. And if you have a bigger disc, such that you're stimulating both the center and the surround, the response will be weaker than if you just stimulate the center. So the neuron is kind of computing a difference between the center and the surround.

And this is an example of how this used to work. So what this is, what I would normally show you, is a video of an experiment where they were recording from one of these ganglion cells. And then this is a projector screen, just like this one, on which light is being shined. But what you have to really be able to do in order to appreciate this is hear the sound because what happens is the electrode that's measuring the response of the cell is being played out as an audio signal. And so every time there's an action potential, you'll hear a click. So you have to be able to hear that.

But what you would see if we could play this, and what you would hear, is that there is this region of space that corresponds to the center of the receptive field. And when there's a spot of light that is in there, there's a brisk response. And if you make the disc larger, the response decreases. So we're going to have to play those next time.

So there's two types of these receptive fields, one that we describe as an on-center, off-surround receptive field. So the center produces an excitatory response and the surround an inhibitory response, and then another one that we describe as an off-center, on-surround cell, where it's kind of the opposite, so two different types. And so you can think of the on-center cells as responding to local luminance increments. And the off-center is responding to local luminance decrements.

So this was probably the first example of a discovery of the structure of receptive fields. And that kind of led to a pretty large and rich research program that we'll talk more about when we get into V1 and extended throughout, like a lot of the rest of the visual system. But in these early stages of the visual system, this notion of receptive field is especially powerful, in part because the neurons can be described in fairly simple terms mathematically. And so we have the ability to I'm going to write down pretty simple mathematical models that do a really good job of capturing the responses of the neurons.

So what we're going to do now is get into that just a little bit. So this is an image. It's an image of Mickey Mouse. But it's just an image. And so what an image is, and for our purposes, is it's a spatial array of intensities. So that's I of x, y . So I is the intensity. And there will be an intensity value at every x position and every y position.

Now, we're going to try to describe the neuron and its receptive field as a filter. So the filter will be G . And G will operate on the input I . And that will produce a response L . Now, because these are neurons-- and we talked a little bit about this a class or two ago in response to a question-- the output of the neuron is action potentials.

And so in addition to this filtering operation, there is an output non-linearity that will actually generate spikes that will turn the response into a firing rate. And then we'll sample spikes from that firing rate. So that's going to be the mathematical model of the neuron. And we will attempt to see whether this does a good job of accounting for those responses

So the assumption that we will make here is that the response of the neuron is given by a linear filter that acts on the stimulus. So what does that mean? It means that the response of the filter is given by this equation. So this is the image. So this is just an array of gray values. This is the filter, which is another array of numbers that are a function of spatial position.

And then we're taking the dot product of the filter in the image. What does that mean? It means that at every point in this spatial array, we take the product of the image intensity and the filter coefficient. And then we add up all those products. That's what this double integral is with respect to both x and y .

So here's an example filter G . That's a center-surround filter. So if we're looking at it in image coordinates, that's what it would look like. If you just take a slice through it, which is sometimes kind of easier to think about, you get this looking thing. That's just in the x direction. So we're going to look at this in one dimension. So now, we just have a single interval with respect to x .

So here's our filter. And this is one example image. This is a very special image because it has the same form as the filter. And the consequence of that is that when we evaluate the response of the filter, we take the product here of G of x and I of x . So at every different position here, we're taking the number here, and we're multiplying it by the number here. And because I of x is the same as G of x , where G of x is negative, I of x is negative, and you get a positive number. When G of x is positive, I of x is positive, and you also get a positive number. So you have all these positive numbers.

So now, when you add all these things up, you take the integral, you get something big, a big response. By comparison, here's the same filter but now applied to a fairly uniform image. So what's going to happen here is now the image is kind of positive everywhere. And that means that where the filter is negative, we get negative numbers as the product. And where the filter is positive, we get positive numbers. And so now when we add all this stuff up, the negative numbers kind of partially cancel out the positive numbers. And so we get a smaller response.

So that's the basic concept for how these filters work. So how can we measure receptive fields? Well, remember a couple of lectures ago, we talked about the idea of the spike-triggered average as one way to measure these filters. And so the way that works is you take a random flicker stimulus. So this is just a sequence of random images. It's a movie.

| present that to an organism whose retina you are recording from. This is the stimulus history. It's a sequence of frames. And then every time there's an action potential, you take the stimulus history that immediately preceded that and you add that to your buffer. And you take all of those little sets of frames, and you get an average sequence of frames. This is like the average movie that preceded a spike. So that's the spike-triggered average. Every time there's a spike, you trigger this averaging process.

So here are examples of these kinds of receptive fields that you can measure in the retina. So here's one. See that? I'll try it again. So there's a little blob that kind of shows up right. Here's another, different color blob.

So that's what you get. If we look at an example of this over time-- so the spike-triggered average is this movie right. That's a sequence of frames. So here's each one of those frames or some subsampling of them. And you can see that there's this little blob here that starts out dark. It kind of then becomes light. And this is the time preceding the spike.

Here is just a quantification of that. So there's this inner region that's been defined here by the dashed line, and then an outer region, and then this region beyond that. And so this is the plot that kind of averages all of the pixels inside the inner outline and then outside the outer outline. And there's three different plots here because there are three different color channels. So this is a movie, each image of which has R, G, and B channels. So there's actually three images in every frame. So we can compute the averages separately for the red, green, and blue channels.

And so you can see there are some differences between the channels. But essentially, in the center, so inside the inner outline, you get this particular shape where there's an increase and then a decrease. And then outside the outer outline, you get this thing that's kind of the opposite. So that's the center-surround receptive field organization. So what's happening in the center is kind of the opposite of what's happening in this surround.

But you can also see, there's these temporal dynamics too, so this sort of picture here with these spatial receptive fields is part of the story. But there's also these temporal dynamics. So this is how you would evaluate a filter at one particular point in space. So that would be like one receptive field at one particular location. Any questions about that?

So this was-- and I should have said, this is how you would evaluate the filter response. And then this is how you measure it. So another kind of important idea in this space is the idea of convolution. So convolution is an operation that produces the output of a linear filter at each location in an image or every time point in some time series. And convolution comes up a lot in different contexts in signal processing.

In how we think about vision, convolution has a very particular role because it's often used to simulate the effects of a large set of neurons with the same type of receptive field at different spatial positions. And so the idea is like we've got these centers-surround receptive fields. So in this particular example back here, this is a receptive field that was measured from one neuron, and it's at a particular location on the retina.

But there's a ton of ganglion cells. Most of them have these center-surround receptive field organizations. And they just differ in where those receptive fields are located. So the idea is that the image is being analyzed by this kind of huge set of neurons that all have receptive fields that look like this, that are just at different positions within the image. And so one way to capture that population, or to simulate it, is by taking a particular filter or receptive field type and convolving it with the image.

So this is what this looks like. So here's an example grayscale image. So it's just a black cross. These are the image intensities that would correspond to that image. So each one of these things is like a pixel. And these are two different filters. They're really simple filters. So they're two pixels wide. One of them is minus 1, plus 1 in this direction. And this one is plus 1, minus 1 this direction.

And so what is shown at the bottom here is the convolution of this image with each of these filters. And so the operation that's being performed here is the same one that we just saw. So it's a dot product. But it's being performed with the filter at every possible location within the image. So let's just do a couple of examples to make sure we get it.

So let's start out with this filter here. We've got minus 1, plus 1. So when that one is put in the top left corner of the image, we've got the pixel value 10 occurring here and 10 occurring here. So we multiply minus 1 times 10 and plus 1 times 10. That gives us minus 10 and plus 10. And we add those together. And that gives us 0. And that's the response that we get here.

So if we just shift it over 1 pixel, so now, we have an intensity value of 10 here and 2 here. So that gives us minus 1 times 10, which is minus 10, and plus 1 times 2, which is 2. Minus 10 plus 2 is minus 8. And then we move it over again. And now, we have 2 and 2. And so those cancel out, and you get 0. You move it over again. Now, you get 2 and 10. And that gives you plus 8 and so on and so forth.

So you evaluate the filter at every different position within the signal. So that's the convolution operation. That's just the mathematical operation. And that's the definition. And you can do the same thing with this particular edge operator. And it works exactly the same way. So you get a 0 here because you have 10 and 10. And then here you get plus 8 because you have the 10 corresponding to the plus 1 and the 2 corresponding to the minus 1, and so on and so forth.

So that's a mathematical operation. It comes up all the time in signal processing. In the context of thinking about the visual system, as I said, we use convolution to simulate a whole bunch of neurons with receptive fields that have the same form but that are just located at different positions, as you would have in the retina.

Here's a movie that will give you another little example of this. So what you're going to see here-- so this is a filter. This is a very simple filter, where the filter kind of resembles a Gaussian. So it's essentially taking a local average. Here's the signal. And this is the result of the convolution. And this is a movie that will go like this.

So there's one other little detail here. So this differs from the example on the previous slide in that the convolution is actually being evaluated for positions of the filter where it doesn't completely overlap with the signal. And so this is a choice that you make anytime you want to perform a convolution. And in different contexts, you might want to do different things.

So the thing that we were kind of seeing started out like here and then moved all the way to over here. And so what happens is that at every position, the result of applying the filter is this little moving average. And so down here, it's negative because most of the values in the signal are negative and so on and so forth, so that's the idea.

So what do we do with all of this? So one place where this was historically kind of influential in how people thought about things was in its application to certain illusions. So this is a very famous illusion, known as Mach bands, named after a person named Mach. So the image here is of this form. So this is a plot of the image intensity as a function of position. And so it starts out at some low intensity. And then there's a ramp up to some high intensity.

Now, the illusion is that people will typically report seeing bands, a dark band here and a bright band there. So for me, the bright band is a lot more salient than the dark band, but I can see them both. It's an illusion because there's no actual band in the image. So the intensity just kind of starts out low and then increases and then kind of levels off.

So that's just this old illusion. And people got interested in the fact that if you take one of these center-surround receptive fields and convolve it with the image, you get out a response profile that actually kind of resembles what you see. So here's the stimulus, the image. Here's a one-dimensional version of the center-surround receptive field.

And the idea here is that when you're evaluating the response of the filter at the image, when you get over here, what happens is that the inhibitory lobe of the receptive field kind of hits the ramping part of the stimulus. And so that causes the response to go down a little bit. And then you get to the ramp, and the thing gradually goes up. And then you get up here to the top, and it's the same thing, where you end up with the inhibitory portion of the receptive field kind of in the ramp. And so it's lower than this part. And that creates an increase in the response.

So people kind of observed this and thought, well, maybe this somehow explains the fact that we see these bands when we look at this image. Here's a similar thing. There's lots of illusions that are like this one, where-- and the illusion here is that you probably, when you look at this thing, you have the impression that there are these dark blobs at the intersections of the white lines. Everybody see those dark blobs? You probably don't see them right at the center of gaze but then kind of at all the other locations.

And so there's this analogous type of explanation here, which is that if you think about what the response of a center-surround filter would be if you convolved it with this image, the response of an on-center, off-surround receptive field would be reduced at one of these intersections, compared to if you're off of one of the intersections. And so the idea is that when this receptive field is centered here, you get more white stuff in the surround than when it's centered here. And because this is an inhibitory surround, when there is a high intensity inhibitory surround, that produces a lower response.

So there's less dark stuff in the surround when the receptive field is on an intersection. So there's a smaller response to the filter. We've observed that you see a dark spot. And the hypothesis is that the smaller response of the filter could explain why you see a dark spot. Yes?

AUDIENCE: Can you explain why you don't see dark spot on the one that you're focusing on?

JOSH MCDERMOTT: Does anybody have an idea? What do we know to be different about the fovea compared to the periphery? Yeah.

AUDIENCE: High density of receptors and small receptive fields.

JOSH MCDERMOTT: So one possibility is that at the fovea, the receptive field is so small that this stuff doesn't happen. That's at least possible. But that would be one kind of explanation. And can anybody explain why would this spot be dark as opposed to light? Yeah.

AUDIENCE: Because the surrounding squares are dark. If you reversed their colors, would you see spots or would you see even darker--

JOSH Yeah. So if you reverse the colors, it would be reversed. But I was actually asking-- I was looking for an answer in terms of the type of receptive field that we're trying to explain this with. And so specifically, what I was going for here is this is an on-center receptive field. And so normally, the response of that cell kind of signals increments in intensity. And so the idea is that if there's a smaller response from that cell, that would correspond to a decrease in intensity. So that was the underlying logic I was looking for.

But there's something-- so this is like the classic explanation of stuff. It's in all the textbooks. There's something deeply funny about these explanations. And that is that, essentially, they're implicitly assuming that-- it's like there's this notion that your visual system gets the output of these filters and then somehow doesn't know where they came from. So you get the outputs of these filters. And there's a lower response at one point compared to another.

But if in principle, whatever's happening downstream knows that it's dealing with the outputs of center-surround receptive fields, you might imagine that it would be able to decode that and not incorrectly think that there's a dark blob at the center of these intersections. So essentially, these explanations implicitly assume suboptimal decoding of the responses, which is kind of weird.

And so that's the classical explanation. And it's worth knowing about. I think we actually don't really understand why these illusions exist. My intuition is that it has to do with the fact that these are very weird images. They're very unnatural in the sense that they have these very high contrast edges arranged in particular ways. And in general, there's lots of funny illusions that result from images like this, where you have lots of high contrast edges and sharp corners and things like that.

So my suspicion is that it's the kind of image that pushes the system into a funny operating point. And that really is critical to actually underlying this. But really, I would say we actually don't really have a very good deep theoretical understanding of why these kinds of illusions exist. Any questions about that?

AUDIENCE: So basically how do you make the image not appear? It's not like using your perspective, if you, say, zoomed in four times?

JOSH You should try it. That's a great illusion lab. I can tell you from having played around with these things in my

MCDERMOTT: youth, there will be effects that relate to that. I don't know how easy it would be to relate these to actual receptive field sizes that you would find at different stages of the visual system. That would be interesting to try. But yeah, it's certainly a very powerful effect that it doesn't happen at the fovea, and it happens in the periphery. So yeah, you should play with it.

Explaining illusions with receptive fields implicitly involves convolution. So again, the idea is that the convolution kind of simulates the response of this big set of neurons that all has kind of similar receptive fields. We think that that's effectively what your retina is sending to the rest of the visual system.

Another really important thing to know about the visual system is the way that it's organized with respect to space. And so this is a diagram that shows the visual system looking at a screen, essentially. So the idea is that you're fixating that cross. And there's a left hemifield and a right hemifield that are visible. So the colored part of the screen is like the extent of visual space that is visible to your eyes. And it's colored red when it's on the left and blue when it's on the right. And then this shows you where those points in the image end up being represented at different stages of the visual system.

And so what do we take away from this? Well, first of all, both eyes see both hemifields but to different degrees. And you can just-- if you fixate here, you can verify this yourself. But then something crazy happens. So you get the optic nerve exiting the eye, and then the optic nerve fibers kind of split up. So the fibers that correspond to the left hemifield all go to the right side of the brain, irrespective of which eye they come from. And then the fibers that correspond to the right hemifield all go to the left side of the brain.

And so there's a right LGN and a left LGN. That's the part of the thalamus. And so those are segregated according to which side of space we're dealing with. And then those project to the left and right primary visual cortex, so really kind of striking segregation. So what this means is that once you leave the eye, stuff that's happening on the left side of space is represented in a different bit of the brain than the stuff that's represented in the right side of space. And of course, this is all relative to the fovea. So you move your eyes, and what counts as left and right is going to change.

So information from each hemifield projects to the contralateral LGN. But the information from the two eyes stays separated until you get to the cortex. So there are different layers of the LGN that correspond to the left eye and to the right eye. And there's an example of that that's shown here. This is an actual picture of the LGN, so this little bit of brain. And you can see that this is a microscope image, where the cell bodies have been stained.

And you can see that there are these different layers. There are six layers of the LGN. The four outer layers are called parvocellular layers. And as we talked about last time, they receive input from the midget ganglion cells in the retina. And then the two deeper layers are the magnocellular layers. And they receive input from the parasol ganglion cells.

So there's, again, this really kind of precise beautiful anatomical organization where you have these two distinct types of cells in the retina, midget and parasol, both of which come in the on-center and off-center receptive field varieties. And then they project to these distinct regions of the LGN, where different layers would get input from either the left eye or the right eye. So the eyes remain segregated.

So we've got these two types of cells in this way station in the visual system. What's the difference? So parvocellular cells are notably color opponent. So that means that they receive inputs from more than one class of cones and are typically kind of computing the difference between those different cone types.

So we have red-green opponency and blue-yellow opponency. And again, we'll go into that in more detail when we talk about color, so I'm not going to dwell on it. But they also give sustained responses to stimuli. So they tend to respond over extended periods of time, whereas magnocellular cells, by contrast, are achromatic, so they tend to not have information about wavelength, and they give transient responses to inputs.

So these were physiological differences that were observed. And several decades ago, here in this department-- you can see this, Department of Brain and Cognitive Sciences, Massachusetts Institute of Technology-- this piece of research was done by Peter Schiller. So Peter is an emeritus professor. And I did a UROP with him a long time ago, which was a lot of fun. So he was a pioneer in the visual system. And the middle author on this paper is Nikos Logothetis, who's also now a very famous neuroscientist. There's lots of stuff in vision and MRI. And he's now working in Germany.

So what did they do here? So this paper used a method that's no longer really in very widespread use because we now have more precise ways of doing similar things. But in this era, it was very common to actually investigate brain function by lesioning parts of the brain. So that means typically sometimes you go in, and you cut out parts of the brain. Sometimes it means that you would inject acid to kill the cell bodies in a part of the brain. And then one could look at the effect on behavior. And this was one way that you could infer the role of different parts of the brain in function.

And so in this particular paper, they lesioned regions of the LGN. So remember, this is the structure of the LGN. We got these six layers 1, 2, 3, 4, 5, 6. The top four are parvocellular. The bottom two are magnocellular. So they went in and injected-- I think this was ibotenic acid into particular regions. And this is one of the LGNs in one hemisphere. And this is the other in the other one.

So they injected a ibotenic acid into a particular part of the parvocellular layers in one hemisphere and a particular part of the magnocellular layers in the other hemisphere. Before they did that, they measured the receptive field locations. And this is the receptive field locations plotted in visual space. So this is fixation. And so the region of space in which there was a parvocellular region is here. And the region of space where there was a magnocellular lesion is here.

And so now, the idea was that they could present stimuli within these regions. And this was done to monkeys, macaque monkeys, a common animal model in vision. They could present stimuli in those regions and then ask whether the monkey's ability to perform some kind of discrimination or other task on the stimuli would be affected by these two different types of lesions.

And so these are the findings. So they measured the ability to discriminate a whole bunch of different types of basic visual properties. And this is showing the performance of these animals on these different discrimination tasks. So the y-axis here is plotting proportion correct. And there are three types of positions here. N is, I guess, normal. P is the parvocellular lesioned region. And M is the magnocellular lesioned region.

And what you're supposed to take away from this is that in the region where there was a parvocellular lesion, there's a massive deficit in color discrimination and some deficits in what they call texture and pattern discrimination. By comparison, in the region where there's a magnocellular lesion, these other things are normal. But motion discrimination is greatly impaired. And there's also some other deficits. But those are the most striking ones.

And so this kind of provided some very compelling evidence from segregation of function, the idea that there are these distinct populations in the LGN that are carrying information that's very critical for color perception and maybe fine form perception on the one hand and then motion perception on the other. And this starts all the way in the retina. We've got these midget and parasol cells in the retina. And those project to these distinct regions of the LGN. And those project to other stuff, as we'll see. But there's a pretty striking segregation of function. Any questions about this experiment? Yeah.

AUDIENCE: Do you what the explanation is for what's going on between shape one and shape two, why parvocellular is really detrimental to discriminating between square and circle but not square and triangle?

JOSH I'm not sure. Yeah. Yeah. That's a good question. I don't know. Any other questions? Yeah.

MCDERMOTT:

AUDIENCE: Do you think that the curvature is like [INAUDIBLE] like [INAUDIBLE] linearity ever since?

JOSH It's possible, yeah. Yeah, it's possible. Yeah, it's not what I would have predicted, necessarily. Yeah. The
MCDERMOTT: difference between-- so the parvocellular cells tend to have smaller receptive fields. And so yeah, maybe somehow detecting the curvature is more dependent on that or something. I don't know. Hard to say.

So that's just a summary of what we saw there. So lesions in the parvocellular layers produce deficits in the discrimination of color, texture, fine shape, and pattern. Lesions in the magnocellular layers primarily resulted in difficulties in detecting motion, as well as also discriminating flicker. That's from another related study. And so the conclusion here is that the magnocellular cells play a specialized role in transmitting information about motion. And that, I think, has held up pretty well over the years.

So we've been talking about the LGN, the lateral geniculate nucleus. That projects to the primary visual cortex, also known as V1. The visual cortex, in general, is divided up into lots of different areas, as we will see. But V1 is the largest. And so we're going to talk about that next.

So onto the rest of the visual system, in particular, the cortex-- so again, this is where we're at. We're going to talk about primary visual cortex. One of the most salient features of the visual cortex, and again, a very fundamental organizational principle in the visual system, is what's known as retinotopy. And that refers to the fact that the receptive fields that you measure in individual neurons in the visual cortex are organized spatially across the cortex.

So this is an example of an experiment that measures this in humans. So this was measured with fMRI. So this was an experiment where the person is in an fMRI scanner. They're looking at stimuli that will vary in eccentricity. Remember, eccentricity is distance from the fovea. So imagine you're looking at annuli that gradually get bigger very slowly. And then after the fact, you can analyze the data. And for every point in the cortex, you can estimate the eccentricity at which the stimulus produced the biggest response. And then you color code that.

So this is a flattened map of the cortex. And what you're supposed to take away from this-- and so this is the-- it's a diagram that shows you the relationship between the color that's being used on this map and eccentricity. So the point is that yellow is the fovea, as well as the furthest eccentricity. And then red is the parafovea. Green is a little bit further out and so forth. And so you can see these stripes in the visual cortex, from yellow to red to green to blue. So there's a map of space that's laid out on the cortex.

So this is in humans. This is a really famous experiment that kind of showed additional evidence for this. This is in macaque. This is, again, a really old school method that nobody would use anymore today because we have more refined methods. But what was done in these experiments is a monkey was injected with radioactive glucose. And then-- so they're anesthetized, of course. So the monkey's anesthetized and unconscious. And the monkey's made to look at the screen.

So the idea is that you look at that screen. The visual cortex, bits of visual cortex that are processing the image, are active, taking up glucose, which is radioactive. And after the monkey's looked at this for a little while, the animal is sacrificed. They take the brain out and lay it down on some kind of photographic plate. And the idea is that the parts of the brain that are radioactive are then going to show up in an image. And so this is the resulting image of a monkey that looked at that.

And so what you can see-- so that the image that the monkey was looking at was a set of rings that were flickering bits of checkerboard, and they're spaced out eccentrically. And then so it's kind of like there's polar coordinates here. And then you have these spokes that kind of emerge from the wagon wheel. And you can see that those polar coordinates kind of get translated into what looks like Cartesian coordinates, roughly speaking, in the visual cortex.

But the other thing to take away from this is that you can see that the spacing of the rings is not uniform. The ones that are very close to the fovea are very closely spaced. And then as you move out to the periphery, they become much further apart, whereas when you look in the visual cortex the spacing is pretty uniform.

So the stripes on the screen don't translate to a similar spacing in the visual cortex. And this represents a phenomenon, a very important phenomenon, known as cortical magnification. And the essential idea is that there's much more cortical real estate that is devoted to the center of gaze than to the periphery. So the representation of the center of gaze has kind of blown up in the cortex.

And so the consequence of that is that those rings that are very closely spaced on the screen end up being spread out in the cortex, whereas the rings that are very far apart on the screen, because they're out in the periphery, end up being kind of a little bit closer together in the cortex.

So what's that all about? And why does it happen? So we can get some insight into this from looking at the retina. So these are two cross-sections of the retina. So this is the ganglion cell layer. That's the photoreceptors. Remember, you've got all these cell layers in the retina. And you can see here, this is what's called the parafovea. So that's right next to the fovea.

And you can see there's a lot of ganglion cells, because remember, near the center of gaze, the image is sampled very, very densely. There's a ton of photoreceptors, lots of ganglion cells, small receptive fields, high resolution. This is out in the periphery. And you can see that the ganglion cell layer has far fewer cell bodies. And that's because out there, there's not as many photoreceptors. Other receptive fields tend to be bigger. The sampling is much coarser. Resolution is worse.

Now, this kind of has to be set up this way because the retina is physically constrained by the fact that the image is formed on the retina. And so your receptors kind of have to be physically at the image location that they're sampling. And so if you want to have really dense sampling at the fovea, you just have to squeeze everything into the fovea.

So unlike in the retina, in V1, the cell density is relatively uniform. And as a consequence of that, the mapping of space becomes very non-uniform. So you take all those ganglion cells from the fovea, and those get kind of spread out over the visual cortex. And then you end up needing much less cortex for the periphery.

So that's the idea of cortical magnification. So it's a consequence of the foveal sampling organization of the retina that we were just talking about, where we sample very densely at the center of gaze. And in the cortex that translates into a very large region of representation. Questions about that?

So cortical magnification, that's the big idea. Now, there's a lot of other stuff that you encounter once you get to the cortex. So the most famous one is a novel feature of cortical neuronal responses called orientation selectivity. It's what it sounds like. It refers to the fact that the neurons become selective for the orientation of simple stimuli, like lines and bars. So when the line is vertical and placed in the receptive field, you get a brisk response, whereas when it's horizontal, you don't. So this is probably one of the most important discoveries in the history of neuroscience, arguably.

And so we'll talk a little bit about the basis of this. So what might the receptive field of an orientation selective neuron look like? So again, I have a video here of an experiment where the experimenters are recording from a neuron in primary visual cortex, flashing light up on a screen. So the movie shows patterns of light that they flashed up on a screen. And you would be listening to the response of the neuron. Again, we don't have sound today, so I'm going to skip this, and we'll do it next time.

But what you would see is that you would see evidence for a receptive field that I think looks something like this one. I can't remember for that exact example. So essentially, there is an elongated region of space within which light increases the firing rate of the neuron and then kind of adjacent regions of space where light actually decreases the response of the neuron.

So you get these receptive fields that look like this. And so the concept here, it's exactly analogous to those center-surround receptive fields that we were just looking at. It's a spatial filter. So you take that filter, and you compute the dot product with respect to the image. And that predicts the response of the neuron pretty well. Same idea, it's just that these are not circularly symmetric. They now have this kind of oriented structure. And they often kind come in these different flavors and are often described as even or odd functions. Remember, an even function is symmetric about 0, so $f(-x) = f(x)$. And an odd function is antisymmetric.

So these neurons were discovered by Hubel and Wiesel, neuroscientists who worked for many years at the Harvard Medical School and won a Nobel Prize for this and related discoveries. And these neurons were dubbed by them simple cells. And that will be contrasted with complex cells that we'll see in just a second.

So you see these when you get to primary visual cortex. And Hubel and Wiesel also had this fairly influential proposal for how you would build a neuron with this type of orientation selectivity. So remember, primary visual cortex gets its input from the LGN. In the LGN you have center-surround receptive fields, just like you have in the retina. And the idea is that you could create an orientation selective neuron by setting a neuron up so that it would get inputs from a set of neurons that have center-surround receptive fields that are spatially arranged in an organized way.

So the idea is that each one of these things corresponds to one neuron. And so if you took all of these and all of these and all of these, and you wired them up so that they would all provide input to the same neuron, you could think of the response of that neuron as being a sum of the response of these neurons. And the receptive field would loosely be approximated by the sum of those receptive fields. So that was their idea.

So in part because this discovery of orientation selectivity has been such a signature discovery in neuroscience, there's been lots and lots of work trying to unpack the mechanisms underlying orientation selectivity. And I would say it still remains fairly controversial. But I think it's pretty clear that aspects of the original model proposed by Hubel and Wiesel are likely to be correct.

And one piece of evidence for that comes from this paper here from back in 1995. This was from the lab of Clay Reid, who at the time was at the Rockefeller University and then was at Harvard Medical School for a long time. Now he's at the Allen Institute. And so the idea here was to try to understand the way in which orientation selective receptive fields are wired up by measuring at the same time from neurons in the LGN and neurons in visual cortex.

So these are super hard experiments. So you have to have an electrode that is positioned in the LGN and an electrode that's positioned in V1. And then you're trying to locate neurons that have receptive fields that are spatially aligned because the question was, well, if you find receptive fields that are kind of at the appropriate spatial position in the LGN, would they then be connected to neurons in V1 that would have the appropriate receptive fields? So that was what they were trying to get at.

So how does this work? Well, you're recording from the LGN and from V1. These are measured receptive fields. Think of these as like a spike-triggered average. So this is your estimate of the receptive field or the filter. So the LGN has a center-surround receptive field. This is the V1 neuron, which is kind of oriented. So this is an odd receptive field. So there's a lobe here and a lobe here. This is the actual measured spike-triggered average. And then for the purposes of the analysis, they would fit models. This would be a difference of Gaussians. This was like a Gabor function, which we'll talk about in a second.

So they would record from these pairs of neurons. And then the question was, are these neurons connected? And so one way that you can actually get some evidence for that is to look at the relationship in time between spikes that are fired in one neuron and spikes that are fired in the other. So the idea is that if the LGN neuron is providing input to the V1 neuron, that if the LGN neuron spikes, the V1 neuron should be likely to spike just a little bit afterwards.

And the way that you actually detect that is by measuring something called the cross-correlogram. And what this is, is you record from these two neurons. And every time you get a pair of spikes, you measure the time lag between the spike in the V1 cell and the spike in the LGN cell. And you just make a histogram of that.

So what does this show us? So this is 0 here. And this shows that there's this big peak here right after 0, so just a millisecond or two after 0. And that provides evidence of what's called a monosynaptic connection, that the LGN cell's directly connecting to the cell in V1. So that provides some evidence that one neuron is providing input to the other. So now the question is, can you explain the receptive field of the orientation selective neuron in terms of the receptive field of the cell that is providing input to it?

And so the way that they address that was by asking whether the receptive field of the LGN neuron was appropriately positioned relative to the receptive field of the V1 neuron. So in other words, was the excitatory center kind of overlapped with the excitatory lobe of the LGN neuron?

And this is shown here for this one example. So you see the circle here that's fit around the neuron. And then they're replotting that circle around the excitatory lobe. It's a little hard to see here because the lighting is high. But that's there. And then so this is just one example right. So to see whether this relationship kind of held over lots of these pairs, what they did is they recorded from a whole bunch of these pairs and then shifted and rotated the receptive fields of all of the oriented receptive fields so that they were aligned and shifted and oriented the LGN receptive fields in the same way.

So this is a summary figure of the data from this study. So you have all these pairs of neurons. They're now all displayed in this canonical kind of orientation. And so the question is, are the excitatory or inhibitory centers of the LGN receptive fields kind of lining up with the excitatory and inhibitory lobe of the LGN receptive field? And the argument is that they do. So these red circles predominantly fall here. And then the blue circles kind of predominantly fall here.

So this provides some evidence that that kind of original notion from Hubel and Wiesel, that you could explain an orientation selective receptive field just by kind of wiring up combinations of center-surround receptive fields from the LGN. It provides some evidence that that's possibly a reasonable account.

And more generally, this kind of method is something that you could potentially apply elsewhere if you wanted to understand how different kinds of structures in the brain are composed from their inputs. So you establish that neurons are connected and then look at the relationship between the response properties. Questions about that?

So we've got orientation selectivity evident in these neuronal responses. How is this organized? And it turns out that it's organized in a very orderly way. And this gets at another kind of central organizational concept, really, across the entire brain, which is the notion that the cortex tends to be organized in columns.

So what's shown here, it's a schematic result of an experiment where an electrode would be lowered, either in this direction or this direction, and at every position in the brain, you stop and characterize the response properties of a cell at that position. And the line segment that is drawn here represents the orientation preference of the cell at that position.

And so what you're supposed to take away from this is that when the electrode-- this is the cortex. You can think of it as a big sheet, kind of a folded sheet. So when the electrode is inserted perpendicular to the cortex, you have very similar orientation preferences all the way down, whereas when you insert the electrode at an angle, you can see that the orientation preference changes. But you can also see that it kind of changes in a smooth fashion.

And so this hints at two ideas. One is that the cortex is organized in columns, so within a column, the neurons tend to have fairly consistent properties. And then the other thing is that the columns tend to be kind of spatially organized in a particular way so that you have a smooth map of orientation over the cortical surface.

So here's another figure that shows this a little bit more clearly. So this is, again, an electrode that's being kind of lowered through the cortex at an angle. And so this is the track distance in millimeters I don't know why this is so low resolution. But you can see the effect really clearly. And this is the preferred orientation of the neurons that are encountered. And so you can see that this is very smoothly over the course of the electrode track. So as you move through the cortex, the dominant preferred orientation changes in this continuous fashion.

So this is the result of an experiment where it's an imaging experiment, where you look down at the cortex and measure the orientation preferences at each point along the cortex. And so you can see the columns here. So the color here denotes the preferred orientation, I should have said. And so you can see that there are these colored blobs that are distributed all across the cortex. These are orientation columns.

So important concept here, columnar structure, you see that pretty much everywhere in the brain. Orientation columns are a really robust example of that, and then also maps, so maps are really common in the brain. So variables that are very important to be coded, you tend to see tuning for that very continuously.

So we've got this orientation selectivity. One important idea is what's shown here. And that is that just given a single neuron, it's very difficult to actually estimate what the stimulus is. And in this particular example, what you're supposed to note here is that these three stimuli are all designed to evoke the same response from this filter.

So you can see that this is a neuron that's got a vertical receptive field. So the orientation that gives you the best response is vertical. But you can also manipulate the response by changing the image intensity or the contrast. And so here the orientation is changed, and the contrast is increased by just the right amount to give you the same response. So the orientation of the stimulus is confounded with the contrast in this case of an edge.

So this is a really important idea. So this ambiguity between these different stimulus dimensions can be resolved by employing a population code. And so this is the idea. So now, instead of having one neuron that we're measuring the response for, we've got a whole population. The preferred orientations are indicated by these line segments. So this is the response across a population of neurons to this stimulus. This is the response across that same population of neurons to this stimulus. And here's the response across that same population of neurons to that stimulus.

So what can you see is that as the contrast increases from here to here to here, the peak response kind of goes up. And consistent with what I told you, we've selected the orientations and the contrast here such that for this particular neuron, the one that has a preferred orientation of vertical, you get the same response to the three stimuli.

However, you can also see-- and this intuitively makes sense-- that if you were somehow able to measure the peak response over the population, that would tell you what the orientation of the stimulus is. So here the largest response is given by the vertical neurons. Here the largest response is given to the ones that's a little bit off vertical.

So the idea of a population code is that you have the ability to measure the responses from a population of neurons and perform some computation on that to estimate some quantity that you're interested in. And in this particular case, the computation would be to estimate the peak of the response across the population and then to estimate the orientation to correspond to the preferred orientation of that peak.

So where does this end up being useful? One is to help us understand some funny things that can happen to our visual system. So this is going to be a demo of the tilt aftereffect. So we're going to temporarily change your brain. Don't worry. It will just be temporary. And you will then switch back. But what I want you to do for now is look at this stimulus. Everyone can tell if this is vertical. So what we're going to do is adapt you to this stimulus.

And so what I mean by adapt is I'm going to ask you to stare at this for about a minute. So everyone start staring at this. And what I mean, I want you to look at the gray circle. And you can look anywhere inside the gray circle, starting now. So, as you are looking at that stimulus, what's happening is that there are neurons that are firing in your brain. And they're going to fire now, and they're going to fire later. But as they fire, they will adapt.

And so what we think happens is that the response decreases over time. And then that will end up changing the way that things subsequently look, again, for a very brief period of time. Don't worry. So keep looking inside the gray circle. And it's actually kind of important that you move your eyes around inside the gray circle. Otherwise, there's some other funny stuff that will happen, which we can also talk about. So move your eyes around inside that little gray disk just a little bit longer. And then when I give you the go ahead, you're going to then move your eyes over to the test stimulus. And what you should observe is that the test stimulus, the part in the middle, will no longer look vertical. It should look tilted a little bit opposite to the test.

Look over here. Does it look a little bit off of vertical? Yeah. That is the tilt aftereffect. So we adapted your brain to that leftward grating. And subsequent to that, a perfectly vertical grating looked a little bit rightward. It's temporary. So in another-- it's probably even gone now. But in another 20 seconds, it will be gone.

And so like I said, the idea behind after effects is that there are thought to be due to a decrease in the responsiveness of neurons after prolonged activity. And it's unclear whether we should think about this as a feature or a bug. You can think of implementation level explanations of this, which is that neurons use glucose. And then when they fire for a very long time, the glucose kind of is expended, and there's not as much of that or some other metabolic resource. And thus, the response of the neuron decreases.

Other people think that the adaptation actually is of functional importance, that it actually kind of helps your visual system get into some optimal state for the current environment. That's currently debated. But adaptation does happen. And what we will talk about when we resume next time is how we can explain what just happened to you in terms of adaptation and population codes, so to be continued.