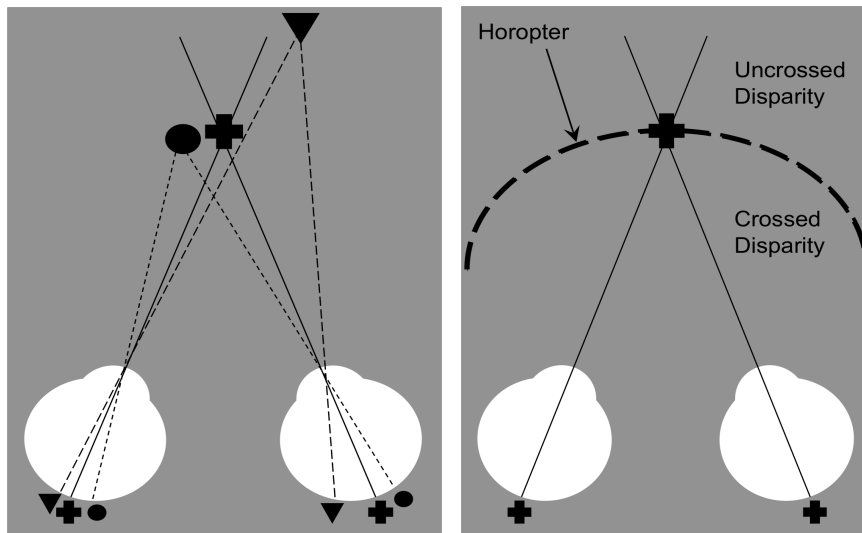




We can see from the above image that the circle and the cross that L and R images fall on corresponding points of each retina, however, for the triangle, half the image falls just to the left of the retina on the left, but a long way to the left of the retina on the right eye. This means they are falling on non-corresponding points and the object has binocular disparity.

The **horopter** is a line of all possible locations where an object's half-image falls on corresponding points. It defines locations at which objects have zero disparity. It is the locus of points in space that yield single vision → anatomical identical points falling on corresponding points on the retina. As shown in the above diagram, objects that fall nearer than the horopter are said to have crossed disparity (think of going cross-eyed), whereas objects further away than the horopter are said to have uncrossed disparity.

But if we have two completely different viewpoints from each eye, how is it that we see one object, rather than having double vision? Because of **fusion**. If the distance between images on each eye is too far away, the disparity is too great then we experience diplopia, which occurs far in front of or far behind the horopter. This usually occurs when there is poor depth information and the eyes cannot converge to the target image, impairing the function of extra-ocular muscles. If objects are within Panum's fusional area (the zone around the horopter where single vision occurs) then there is fusion of the two images. The horopter changes dependant upon your fixation difference.

Fixation is when you look at something so that the image falls directly in the middle of your fovea. If you fixate on an object with both eyes, you converge on it. **Fusion** is the phenomenon that some objects with small disparities are not seen as double, instead the images are 'fused' into one. **Focus** relates to the shape of the lens; if the current state of the lens causes light rays from a specific distance to meet at a single point on the retina, objects at this depth will be in 'focus'.

* 3D Images

In order to simulate stereoscopic 3D images from flat images you need 2 images or photos taken 6.5cm apart from left to right. These images then need to be displayed to each eye. Many methods have been derived for this:

- Mirror stereoscope – devised by Wheatstone. Mirrors are angled at 45 degrees in front of the observer so that the left eye sees the left image and the right eye sees the right image. The images need to be reversed, because they are being reflected off a mirror. Also, it only allows one observer. First ever 3D viewing device.

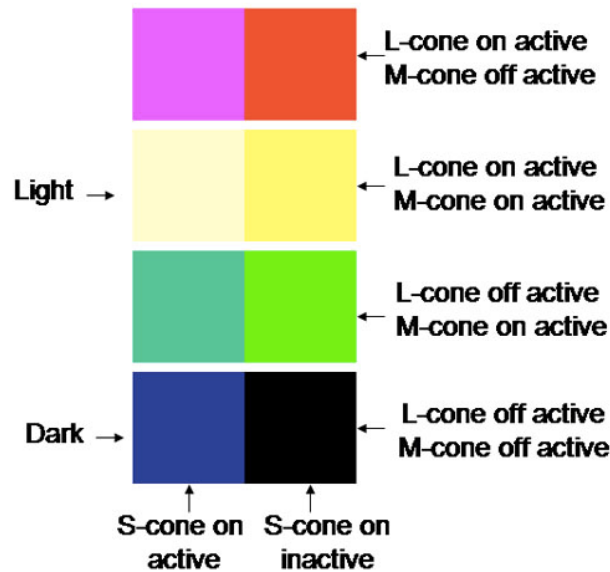
intrinsic (at the edge of an object) or extrinsic (the occlusion of the object by another, closer object). Intrinsic terminators are real edges (so motion signals true object direction) whilst extrinsic terminators are caused by occlusion and are not the real edges of the object (thus motion signals are misleading). Depth cues can tell us which terminators are which, through occlusion, disparity, and shadows. These effects can also occur with 2D images, as before, there is an 180 degree directional ambiguity → all possible motion vectors terminate along one line of constraint. Physiologically speaking, end-stopped neurons process terminator motion as early as V1 and involve hypercomplex cells. V1 simple/complex cells tend to respond to component motion so they suffer from the aperture problem, whereas MT+ cells tend to respond to pattern motion regardless of direction, thus solving the aperture problem.

** Inflow vs. Outflow Theory*

Retinal motion could be explained by object motion or eye movements (eg. Pursuit). We can tell which by comparing retinal motion with some measure of eye velocity. Inflow theory (Sherrington) holds that eye movement signal comes from eye muscle proprioceptors which sense muscle stretch and send signals to the comparator. Outflow theory (Helmholtz) holds that eye movement signal comes from eye movement commands from brain which goes to eye muscle and to comparator. For moving eyes around with an after image: Inflow has no retinal motion but does have muscle movement signal, and outflow has no retinal motion, but does have an efference copy command. Both theories are therefore correct, motion is perceived, like pursuit of a moving object.

Preview from Notesale.co.uk
Page 15 of 30

Trivariant Color Vision



Colour constancy refers to the tendency for objects to generally look the same colour in a wide variety of lighting conditions. To achieve colour constancy, the visual system appears to compare the responses of different colour channels across large areas of the visual field. Neurons demonstrating colour constancy are first found in area V4.

Preview from Notesale.co.uk
Page 20 of 30

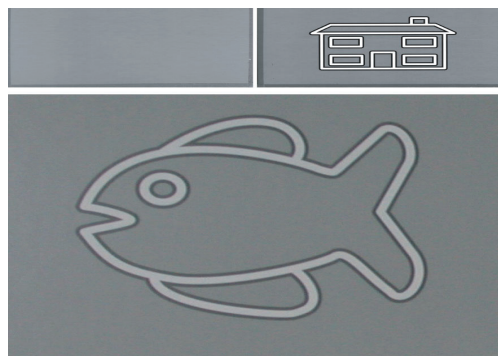
Week 11/12 Development of Vision

How do we measure vision? There are a number of measures that are used to try and determine how it is that infants develop vision.

a) Behavioural Measures:

- Preferential looking: determining where a baby fixates its vision. Although you can focus attention on something you are not looking at, 99% of the time we are focusing on what we are looking at. This measures both acuity (smallest amount of difference between stimulus and no stimulus that is detectable) and contrast (the difference between the amount of light and dark area that is required for a discrimination to be made between the two). Babies show a strong interest in stripes up until the age of 12 months.

- Optotype cards: low contrast cards that require ability to tell acuity to recognise the figure.



- Indirect boredom measure: eg. Thumb sucking. This measures brain activity, whether they are responding to a stimulus. The typical behaviour of a bored baby is thumb sucking, so if they respond to a stimulus by stopping thumb sucking it is an indication they are responding to a visual stimulus.

- Visual Evoked Potentials (VEPs) a measure of brain activity to determine whether there is a systematic signal over and above the random noise, and therefore a response to a stimulus.

* Spatial Vision

There are various aspects of the development of vision. The first is spatial vision; the ability to detect a stimulus over no stimulus. Acuity increases over the first 32 months of life, almost linearly. Contrast sensitivity function (CSF) also measured shows babies are poor at low to medium spatial frequencies (represented by objects that tend to have large areas), but are ok at high frequencies (fine detail). Adult levels of CSF are reached by about 8 y.o, which is a relatively long time compared to most other species. Humans take a long time to develop vision of black and white relative to our lifespan. Vernier acuity is the measure of displacement of a line/grating, ie. can you tell whether the grating consists of a set of parallel black and white bars or, whether those bars are displaced. Acuity and Vernier acuity do not develop at the same rate. Initially, acuity in babies is much better than Vernier acuity, that is, they can distinguish there is a stimulus but not necessarily displacement of stimulus. At about 20 weeks of age we become better at Vernier acuity than simple acuity.

* Feature detection

Is orientation detection innate? By shifting phase and orientation of stimulus we can record the activity in the visual cortex using VEP studies and see that phase change recorded at earliest age but no orientation detection until about 6 weeks. Changes in phase and

Week 13 Face Perception

Faces are important to us as social animals. We seem to be really very good at face perception. We can remember thousands of faces across our lifetime, and most of the time can recognise a face despite changes in many things such as angle of face, when lighting changes, or when they age, change expression, hairstyles etc.. We can also automatically perceive someone's race, gender, identity, emotion etc. We can see and recognise a face even given only minimal information. People are predisposed to see faces even when there is no face.

** Perception and Memory*

Many face processing studies are principally concerned with memory (cognition). Recognition/recall tasks involve a learning phase where subject is exposed to a set of faces, with instructions to remember them, followed by a retinal interval. Then you do a test phase where subject views again, along with never-seen-before faces and asked if you saw those faces in the first task. In order to remember something, we have to be able to encode it first (perception). Some studies investigate this aspect of face processing per se (matching tasks).

Studies that use the same image (for learning and recall phases in memory tasks, or for matching stimuli in perceptual experiments) may not reveal anything about the details of processing faces per se, just perception of faces. Humans are not good at remembering/matching images (of faces or otherwise), because in real life we never see the exact same image twice. With images it can become about image matching rather than face matching.

Are faces special? Are they dealt with by the visual system in a different way to the way that we deal with other objects. The answer is yes, well probably. There is some debate about this and it is not unanimous. One of the things that make us think faces are special is that faces can be recognised and discriminated by children at an extremely early age. Babies are essentially born ready to process faces. Some infants can detect faces over "scrambled" faces within being only 1 hour old. Two-day old infants will track pictures of their own mother vs. other women.

The Composite Face Effect is when half images of two different people are put together (Russel Crowe's eyes and Mel Gibson's mouth). When you jam different parts of faces together, it makes a whole new face percept in your head and you process it as a whole new face. As soon as you split the two half images so they are not an aligned face, you can instantly recognise them as parts. Humans cannot avoid holistic processing of faces. Face processing is more holistic than any other object recognition. The **Part/Whole effect** dictates that when we only have parts of a face (such as a nose) we find it more difficult to identify the face.

Inversion Effects also affect how easily we recognise faces. It affects faces more than it affects any other object stimulus type. The reason for this is that your holistic processing is suddenly disabled. We do not engage all of the same processing as we do not recognise it as a face anymore. An inversion effect is the **Thatcher Illusion**; expressions are more difficult to discern in inverted faces. With inversion, there is poorer matching and recognition for inverted vs. upright faces. When a face is turned upside down you no longer get the composite face illusion, and you don't get the whole part effect.

** Neurophysiology*