

Introduction to Brain & Behaviour
Professor Ark Verma
Assistant Professor of Psychology
Department of Humanities & Social Sciences
Indian Institute of Technology Kanpur
Week 7 Lecture 33: Language Comprehension - 1

Hello and welcome to the course Introduction to Brain and Behaviour. I am Doctor Ark Verma from IIT Kanpur. This is the seventh week of course, and I will be talking to you about Neuroscience of Language Comprehension.

(Refer Slide Time: 0:27)

Perceptual Analyses of the Linguistic Input

- It seems that the brain uses the same mechanisms to understand spoken language and written language, however, there are some key differences.
- When individuals attempt to understand spoken words, the listener needs to decode the acoustic input. The result of this analysis is then converted into a phonological code, and the lexical representations of the auditory word forms stored in the mental lexicon would be accessed.
- Further, after the acoustic input gets translated into a phonological format, the lexical representations in the mental lexicon that match the auditory input can be accessed (lexical access), and the best match can then be selected (lexical selection).
- The selected word includes grammatical and semantic information stored with it in the mental lexicon.

The first (st) first step in understanding or comprehending language irrespective of modality, whether you are listening to the language or you are reading the language, let us say in form of a book or something needs to be the perceptual analysis of the linguistic input. It seems that the brain uses the same mechanisms to understand spoken language and written language. However, there might be some key differences between the two modalities. For example, when individuals attempt to understand spoken words, the listener needs to decode the acoustic input, which is basically the input in the sound format.

The result of this analysis is then converted into a phonological code, and the lexical representations of the auditory word forms stored in the mental lexicon will now be accessed. So, the idea is whenever you are listening to something being said what you do is, you basically try and convert whatever you are hearing into a sort of a phonological code which is basically combination of phonemes that have formed these words.

Once you have activated this basically what you do is that you take these take these particular code and you match it with a similar code or a matching code which is present already present in your auditory mental lexicon let us say. Okay? So, as in the last lecture I said that the mental lexicon could be same for both modalities, or it could be different.

So, let us say if there is in the mental lexicon whether it is the same lexicon for both or a different auditory mental lexicon, this mental lexicon will also have these words stored in this sort of a phonological board. So what happens is when you are listening to some language, basically you try and decode it in terms of this phonological code, you analyse it in terms of this phonological code. And whatever the phonological code of input is matched with the already stored phonological code of the word form in your mental lexicon.

What you will do now is you will try and match this two things, which is basically the process of lexical access we talk about. Okay? So, the lexical representations of the auditory word forms stored in the mental lexicon can now therefore be accessed. Further, after this acoustic input will get translated into the phonological format, as I just described, the lexical representations in the mental lexicon that match the auditory input, can be accessed. So, this is the process of lexical access.

So once you have converted what you were hearing into this phonological code, and you have searched for a matching phonological code from your mental lexicon, both of these things will be compared. That will be lexical access. Obviously when you are doing this comparison, a perfect match will be available but there might be some partial matches also. So all of those words will be activated as well. Now basically because you want to repeat back the same thing, you want to select something that is best that is the best match out of the several activated alternatives.

This again selection of the best match is referred to as the process of lexical selection. The selected word includes now once you selected this particular word and you have you know access this entry from the mental lexicon, its entry in the mental lexicon will include both grammatical information and semantic information stored in the mental lexicon.

(Refer Slide Time: 3:53)

- This information helps to determine how the word can be used in a given language.
- Finally, the word's meaning (store of the lexical-semantic information) results in the activation of the conceptual information.
- The process of reading words, shares only the last two steps of linguistic analysis (i.e. lexical and meaning activation) with auditory comprehension, but because it is a different modality, it differs at the first few steps.
- In reading, the perceptual analysis requires the reader to identify the orthographic units from the visual input. Then these orthographic units may be mapped onto orthographic word forms, or translated into phonological units which would activate the phonological word form in the mental lexicon as in auditory comprehension.

This information helps to determine how the word can be used in a given language. Okay? So, finally once you have activated all of this and you have basically you know you have used this particular word now what you have to do is, basically you have to access the meaning of the word. At the level of lexical access, we are generally only talking about the word for activation. When you are hearing the word, it is the phonological form activation.

When you are reading the word, it is the visual form activation. Now, once you have done the lexical access and lexical selection, and you have accessed the information about what the word form is about, what are its syntactical aspects, some sense of semantic aspects as well, what you then do is, you basically activate the word's full meaning which is basically the lexical and semantic information. And this accessing the word's full meaning, will result in the activation of the conceptual information.

Now the process of reading words is slightly different. In the process of reading words, only the last two steps of the linguistic analysis is basically you know will be common. Basically the lexical access and activating the meaning part. The initial steps will be different because we are talking about a different modality all together. Okay? So, in reading what actually happens is that the perceptual analysis will require the reader to identify the orthographic units, that is the script, that is visual form from the visual input.

Then this orthographic units may be mapped on to orthographic word forms or translated into phonological units. So what you can do is, in reading there is an extra step because you have to go from the visual to the phonological code and then again follow the same series of steps

that we talked about in the auditory word form comprehension. Okay? So, in reading what you first do is, you do the perceptual analysis.

This perceptual analysis basically requires the reader to identify the orthographic units from the visual inputs. Suppose you are reading a word caterpillar, so you have to identify C A T E R P I L L A R. Basically once you have identified this you are going to look for a similar matching visual form from your mental lexicon. Once you have basically activated both of these words, then these then what you can do is, or say for example even before that, once you have as soon as you have activated the visual part you will need to convert that visual part into a phonological part, because access to mental lexicon is from sound units. Okay?

So, these orthographic units may be matched together, match on to orthographic word forms, or alternatively you can translate them into phonological units which will activate a phonological word form in the mental lexicon, as in auditory comprehension. So, either you can directly go from the visual, form pick out a visual entry from the mental lexicon and proceed further. Or you go from visual form to phonological form and then you access the mental lexicon. That is basically what is happening.

(Refer Slide Time: 6:55)

Understanding Speech

- A listener of speech is confronted with a variety of sounds in the environment, and needs to identify and discriminate the relevant speech signals from other kinds of “noise”.
- **Phonemes** are the smallest units of sounds that make a difference to meaning, and their number varies across different languages, for e.g. English has 40 phonemes, whereas others languages may have different numbers of phonemes.
- Perceiving phonemes is different for different language speakers, for e.g. the L & R sounds are different phonemes for speakers of English, they are represented by the same phoneme by speakers of Japanese.

Now let us talk a little bit about understanding speech in a little bit more detail. Now a listener of speech is confronted with a variety of sounds in the environment, and needs to identify and discriminate first, the relevant speech signals from other kinds of "noise". So, our environment is filled with so many different sounds. There is a you know there are sounds of animals, there are sounds of living things, motor sounds, fan sounds and moving of furniture sound. There are so many sounds.

Now if you are having a conversation in such kind of a busy environment, the first thing that you as a system will have to do is will have to distinguish the relevant speech sound, let us say somebody is telling you something you have to distinguish the relevant speech sound from all other kinds of noise. Now, a one of the most important units here is the phoneme. Phonemes are the smallest units of sound that actually make a difference to meaning, and their number varies across different languages. Say for example, as I told you also probably in the last class that English has around 40 phonemes. Whereas many others languages whereas other languages may have different number of phonemes altogether.

Now perceiving phonemes is different for different speakers of different languages. Say let me take an example, the L and the R sounds are different phonemes for speakers of English, that is why we have words like say for example lion and rat, where L and R are very different. Say for example, but these two sounds are represented by the same phonemes in Japanese Language and therefore Japanese Language speaker will basically store them as this version of the same phoneme.

(Refer Slide Time: 8:37)

- Infants have the perceptual ability to distinguish any possible phonemes during their first year of life. Kuhl & colleagues, found that initially infants could distinguish between any two phonemes presented to them, but gradually during the course of the first year, their perceptual sensitivities become tuned to perceive the phonemes of their native language.
 - For e.g. Japanese infants may learn to treat L & R as one phoneme, whereas English speaking infants will learn to treat the two separately.
- The babbling and crying sound that infants articulate during the ages of 6-12 months grow more and more similar to the phonemes that they hear in their languages.
- By the time they reach 1 year of age, they no longer produce or perceive nonnative phonemes.

Infants when they are born have the perceptual ability to distinguish any possible phonemes during the first year of their lives. First year of the life children have this very broad you know band of the different kinds of poems they can perceive unless you do. In some sense you could say that this is an evolutionary gift which allows our system to be able to learn any language where we are born in.

We are not really say for example, it is not determined that we are born to English speaking parents, you would definitely only have the you know neural know how to learn English.

Even if you are born of English speaking parents, but you were raised in a let us say a Tamil speaking background you will be able to learn Tamil as well, because you were hearing that from your first year. So, during the first year of life, infants will have this perceptual ability to distinguish any possible phonemes that they encounter during the entire year.

Kuhl and colleagues, actually found that initially infants could distinguish between two phonemes presented to them, but what happens is gradually this immunity sort of fades out. So gradually during the course of the first year, their perceptual sensitivities become more tuned, more accustomed, more adapted to perceive the phonemes of their native language only, and this generic immunity of being able to perceive and separate phonemes kind of fades away.

Let us take an example, so Japanese infants at birth will be able to treat L and R as one phoneme, whereas English speaking infants will learn to treat L and R as two different phonemes. Okay? Now, the babbling and crying sounds that infants articulate or that sounds that infants make during the ages between 6 to 12 months grow more and more similar to the phonemes that they are hearing in their languages. So, what is happening is that the production system is practicing you know creating these sounds and what happens is the sounds that they kind of make in the beginning gradually start to resemble the sound systems in their in the infants native languages.

Now, by the time infants would reach 1 year of age, they are no longer able to produce or perceive non-native phonemes. So, after 1 year of age a Japanese child will not be able to perceive L and R as separate phonemes, they may basically be able to treat them only as one phoneme. Similarly children will not be able to produce phonemes or even perceive phonemes from other language groups which do not belong to them.

(Refer Slide Time: 11:19)

- Recognizing phonemes is an important part of the ability to understand any spoken language. The listener's brain must be able to resolve a number of challenges in order to be able to recognize the phonemes.
- Some of the challenges include: a) variability of the signal, b) the fact that phonemes are co-articulated, and need to be segmented from each other in order to be identified discretely.
- One might ask, whether and how individuals identify the spoken input given the variability in the signal and the segmentation problem?

Yeah, so that is one of the important ideas. Now, so that is about the variability of the signal. Also within the same speaker, given different emotional states we will basically be able to you know we will basically be able to produce the sounds differently. Suppose when you are very excited, you produce sounds in a different manner, versus when you are very calm and relax. When you are to impart good news your sound will be different, when you are supposed to give poor bad news your sound will be slightly different.

Now the fact that also one of the other challenges is co-articulation. We do not say Pa, Ba, Ca, Ta, all of these phonemes separately. We obviously co-articulate them along with the other sounds that we are producing that form part of words. For example, Pa, Uh and La, we do not say that separately if you want to say the word POOL, we combine the three sounds we sort of blend them together and then we speak a joint utterance. That makes it very difficult for listeners to actually be able to segment this. Okay?

So, in order to be able to segment these sounds from each other, that is a bit of a challenge. So somebody who is listening the language has some of these challenges to take care off. Now one might ask whether individuals identify the spoken input, or say for example how they identify the spoken input given the large amount of variability that is possible in the signal, and also the fact that there is segmentation problem.

(Refer Slide Time: 13:02)

- One of the important clues, is *prosodic information*, i.e. the listener derives from the speech rhythm and the pitch of the speaker's voice.
 - The speech rhythm comes from variation in the duration of the words and the placement of pauses between them.
 - Prosody is apparent in all spoken utterances, but it is perhaps most clearly illustrated when a speaker asks a question or emphasizes something. When a question is asked, the speaker raises the frequency of the voice towards the end of the question, and when emphasizing a part of speech the speaker typically raises the loudness and includes a pause after the critical part of the sequence.
- Anne Cutler & colleagues (Tyler & Cutler, 2009) at the Max Planck Institute for Psycholinguistics, Netherlands have showed that English listeners use syllables that can carry an accent or stress (strong syllables) to establish word boundaries.
 - For e.g. a word like *lettuce* with stress on the first syllables, i.e. trochaic word, is usually heard as a single word. In contrast, words such as *invests* has stress on the last syllable and maybe heard as two words and not as one word.

Now, one of the important clues, is prosodic information. Prosodic information basically is the fact that there is the variation between the rhythm and the pitch of speaker's voice. So, what the listener does is, the listener derives this information prosodic information from the speaker's rhythm and pitch and basically uses that prosodic information to understand that whether there are pauses or whether these sounds are starting or ending at some point. Let us take an example, say for example if I want to tell you the word nitrates, then you will be sure that I am talking about a chemical called nitrate.

However it could also be possible that if I just give a pause, if I say like night rays. And maybe you are taking about I am talking about let us say I am proudly talking about a hotel and I am talking about night rates of that hotel. So, just using these pause or just using say for example you know the the rhythm of all of this speak listeners can actually derive this very important information and be able to understand speech better. This speech rhythm comes from variation in the duration of words and the placement of pauses between them, so thus what I was explaining to you.

Now, prosody is supposed to be an important factor and is apparent in all spoken utterances, but it is perhaps the most clearly illustrated when a speaker asks a question or emphasizes something. Say for example, you know where were you? Where, so the emphasis is on where. Or say for example, did you eat this thing? So, the way I am saying this I am placing emphasis on the word eat, okay? So, the prosody is apparent in all spoken utterances, but it is most clearly visible when somebody is asking questions or is placing emphasis on something.

Now, when a question is asked a speaker raises the frequency as I was saying raises the frequency or the voice towards the end of the question, and when emphasising a part of speech speakers typically raises the loudness and includes a pause after a critical part of the sentence. So, you know this are just examples of how people do variations. These might not be universally applied across all speakers and all languages, but I hope you get the idea that the way we speak, the pauses we take, the rhythm changes that we do, actually act as very important clue for listeners to understand whatever we are speaking.

Anne Cutler and colleagues from the Max Planck Institute of Psycholinguistics, at Netherlands have actually shown that English listeners use syllables that can carry accent or stress stress to establish word boundaries. For example in a word like lettuce, the stress is on the first syllables. Lettuce has 2 syllables, lett has more stress uce has less stress. So, this is basically what is referred as trochaic words. And most words in English approximately anywhere between 60 to 75 percent of the words are trochaic words. Okay?

And typically what happens is when you do this, it is usually heard as a single word. In in the other hand if you have a word like invests, you see that the stress has shifted on to the next syllable and basically sometimes what happen if this can cause the word invests to be heard as two words and not as a single word. Okay? So, these are these minute and you know variations in the stress patterns, which speakers utilize to understand what basically is being said. Let us dell a little bit more in detail...

(Refer Slide Time: 17:01)

- Heschl's gyri are located on the supratemporal plane, superior and medial to the superior temporal gyrus (STG) in each hemisphere, and contain the primary auditory cortex, that processes the auditory input first.
- The areas surrounding the Heschl's gyri and that extend into the superior temporal sulcus (STS) are together known as the **auditory association cortex**.
- Neuroimaging studies in normal participants have shown that Heschl's gyri of both hemispheres are activated by speech and nonspeech sounds alike, but that the activation in the STS of both hemispheres is modulated by whether the incoming auditory signal is a speech sound or not.
- Moving away from the Heschl's gyri, the areas of the brain become slightly less sensitive to changes in nonspeech sounds but more sensitive to changes in speech sounds. The posterior portions of the STS of both hemispheres, although mainly on the left, seem especially involved in the processing of phonological information.

Heschl's gyri that are located on the supratemporal plane, superior and medial to the STG that is superior temporal gyrus in each hemisphere, actually contain the primary auditory cortex,

and this is the part of the brain that processes incoming auditory input first. The areas surrounding the Heschl's gyri extend and that extend into the superior temporal sulcus are together known as the auditory association cortices. So, auditory association cortices do the further processing.

Now neuro imaging studies in normal participants have shown that the Heschl's gyri of both the hemispheres are actually activated by speech and non-speech sounds in a similar manner. But that the activation of the superior temporal sulcus in both the hemispheres is modulated by whether the incoming auditory signal is a speech sound or it is not. Okay? So, it is probably where some of the differentiation between speech and non-speech sounds is being made. Now moving away from the Heschl's gyri, the areas of the brain that become slightly less sensitive to changes in non-speech sounds but more sensitive to changes in speech sounds.

As we go further up from the brain, what do you see is that the specialization for being distinguishing non-speech sound reduces, because they are obviously less important to us as opposed to the sensitivity to speech sounds because we need to understand them in much more detail. So, the posterior portions of the superior temporal sulcus of both the hemispheres, although mainly majorly on the left hemisphere, they seem especially involved in the processing of phonological information.

(Refer Slide Time: 18:48)

- Wernicke had found that patients with lesions in the left temporoparietal region that includes the STG had difficulty understanding spoken and written language. However, the lesions were not limited to the STG.
- A study that contributes to a better understanding was conducted by Binder & colleagues (2000). Participants in the study were made to listen to different types of sound, both speech and nonspeech.
- The sounds presented were of different kinds: white noise without systematic frequency or amplitude modulations; tones that were frequency modulated between 50 and 2400 Hz; reversed speech, which was composed of real words played backwards; pseudowords which were actually pronounceable strings of nonwords that contain the same letters as the real word, for e.g. sked for desk.

Carl Wernicke had actually found through his investigations that patients who have lesions in the left temporoparietal region including the superior temporal gyrus had difficulty in understanding spoken and written language. However, the lesions were not limited to the

superior temporal gyrus as we discussed earlier. They spread out to other surrounding (())19:11) regions as well. One important study that contributes to a slightly better understanding was conducted by Binder and colleagues. Okay?

And basically what happened was that participants in this study were made to listen to different types of sound, so speech sound as well as non-speech sounds. The sounds presented were of different kinds. So, white noise with systematic frequency or amplitude modulations, tones which have frequency modulated between 50 hertz and 2400 hertz, reversed speech, which was composed of real words just played backwards, pseudo words were also there which were actually pronounceable strings that contain the same that contain the same letters as the real word, say for example sked for desk.

(Refer Slide Time: 20:00)

- The results indicated that relative to noise, the frequency modulated tones activated the posterior portions of the STG bilaterally. Areas found more sensitive to the speech sounds than to tones, were more ventrolateral, in or near the superior temporal sulcus, and lateralized to the left hemisphere.
- Binder and colleagues' study also demonstrated that these areas not involved in lexical-semantic aspects of word processing, as they were equally activated for words, pseudowords and reversed speech.
- Based on their fMRI findings, Binder and colleagues (2001) proposed a hierarchical model of word recognition.
 - According to their model, processing proceeds anteriorly in the STG. First, the stream of auditory information proceeds from auditory cortex in Heschl's gyri to the superior temporal gyrus. Here, no distinction is made between speech and nonspeech sounds.

Now the results of the study indicated that relative to noise, frequency modulated tones activated the posterior portions of the superior temporal gyrus, booking both the hemispheres. Areas found more sensitive to the speech sounds than to tones, were more in the ventrolateral part, in or near the superior temporal sulcus, and they were lateralized mainly to the left hemisphere. So, it seems that the left hemisphere it has some specialist abilities which utilizes in analysing heard speech.

Now, Binder and colleagues' study also demonstrated that these areas are not involved in lexical-semantic processing of word. So they are not really concerned with meaning but basically word for analysis. Okay? Because they are equally activated for words, pseudo words and reverse speech. All these three kinds of syllabi have characteristics of human speech but they are not they are different from each other in terms of meaning. Say for

example, only the words will have meaning. Pseudo words do not have meaning and reverse speech basically sounds like gibberish.

Now, based on their fMRI findings, Binder and colleagues actually proposed a hierarchical model of word recognition. Okay? So, according to their model what happens is, that processing proceeds anteriorly from the superior temporal gyrus. First, the stream of auditory information proceeds from the auditory cortex in the Heschl's gyri to the superior temporal gyrus. Here, no distinguishing no distinction is made between speech and non-speech sounds.

(Refer Slide Time: 21:34)

- The first region of the brain where the distinction is made is around the mid-portion of the superior temporal sulcus; but here no lexical-semantic information is processed.
- Neurophysiological studies indicate that recognizing whether a speech sound is a word or nonword happens in the first 50-80 ms. Processing of this kind is lateralized to the left hemisphere, where the combinations of the different features of speech sounds are analyzed.
- From the superior temporal sulcus, the information proceeds to the final stage of word recognition in the middle temporal gyrus and the inferior temporal gyrus, and finally to the angular gyrus, posterior to the temporal areas and in more anterior regions in the temporal pole.

Then, the first region of the brain where the distinction is made is around the mid-region of the superior temporal sulcus, but again here there is no lexical-semantic processing. Neurophysiological studies have indicated that recognizing whether a speech sound is a word or a non-word happens in around the first 50 to 80 milliseconds of time. Now processing of this kind is mainly lateralized to the left hemisphere, where the combinations of different features of speech sounds are analysed.

So that you can make sense of okay, these sounds are present, whether they can be combined together to make a meaningful word or not. From the superior temporal sulcus onwards, the information proceeds to the final stage of word recognition in the middle temporal gyrus and the inferior temporal gyrus. And finally it moves to the angular gyrus, which is posterior to the temporal areas and inside the more anterior areas of the temporal pole.

(Refer Slide Time: 22:30)

- A review of at least 100 fMRI studies, by DeWitt & Rauschecker (2012) confirmed the findings that the left mid-anterior STG responds preferentially to phonetic sounds of speech.
- More recent fMRI studies by Blumstein & colleagues, suggest a network of areas involved in phonological processing during speech perception and production, including the left posterior superior temporal gyrus (activation), the supramarginal gyrus (selection), inferior frontal gyrus (phonological planning) and precentral gyrus (generating motor plans for production).

A review of at least 100 studies by DeWitt and Rauschecker confirmed the findings that the left mid-anterior to superior temporal gyrus responds preferentially to phonetic sounds of speech. So, it is established this is the reason which is this is the region which is more involved in processing of humans speech sounds. Now more recent fMRI studies by Blumstein & colleagues, suggest that there is a network of areas involved in phonological processing during speech perception and production, including the left posterior superior temporal gyrus, the supramarginal gyrus, the inferior frontal gyrus and the precentral gyrus, and the job for each of these are slightly distributed.

So, they say that left posterior superior temporal gyrus is involved in activation of sounds, supramarginal gyrus is involved in the selection of sounds selection of correct sounds, inferior frontal is is basically involved in logical planning and precentral gyrus is basically there for implementing general motor plans for production.

(Refer Slide Time: 23:40)

Written Input

- Reading involves perception and comprehension of written language. To understand written input, readers need to recognize a visual pattern, something our brains are very good at.
- Learning to read requires linking arbitrary visual symbols into meaningful words. The visual symbols vary across different writing systems.
- Words can be symbolized in writing in three different ways: alphabetic, syllabic, and logographic: For e.g. most western languages are alphabetic, Japanese is syllabic while Chinese is logographic.
- Irrespective of the specific writing system, readers will need to analyze the primitive features, or the shapes of the symbols.
 - For e.g. the horizontal, vertical and curved lines and elementary features that combine to produce alphabets.

Now, we have talked a little bit in detail about understanding spoken speech. Let us talk a little bit about written input. Now, reading basically as I said involves perception and production of written language. To understand written input, basically readers need to first recognize a visual pattern, which our brains are peculiarly good at. Now, learning to read requires linking arbitrary visual symbols to meaningful words. Say, for example whatever script you might be writing, for say if you did not know the sounds of them and you did not know that they can be combined to create words they will be just arbitrary.

So, learning to read basically requires linking these arbitrary visual symbols from our scripts to meaningful words. And this visual symbols vary a lot across different kinds of scripts, different kinds of writing systems. So, for example so, for example the way the writing system in English works, is very different from the way writing system in Hindi works or say for that matter in Japanese, Chinese or French. So, words need to be symbolized in writing in three different ways.

There are alphabetic ways alphabetic scripts, there are syllabic scripts, and there are logographic scripts. For example most of these western scripts of European languages are alphabetic, Japanese is syllabic, whereas Chinese is logographic. Okay? Now, however we can talk about this in some detail more in further but basically it is not very necessary because irrespective of the specific writing system, readers will basically need to analyse the primitive features of each of the scripts.

However different aspects say for example how is a you know curved line and straight line join to form the alphabet that represents P. And of you add a slanted line under the curved

line, then it becomes R and so on. So, basically to be able to read irrespective of the writing system involved, the readers will need to analyse the primitive features or the shapes of these symbols. So horizontal, vertical and curved lines and all of these elementary features that you are talking about.

(Refer Slide Time: 26:05)

- To understand the process better, one can consider a model forwarded by Oliver Selfridge who postulated that a collection small components combined together could allow the machines to recognize the patterns.
- These small components, or *demons* as Selfridge referred to them, would record events as they occur, recognize patterns in those events, and then trigger subsequent events, according to the patterns they recognize.
- In Selfridge's model, known as the ***pandemonium model***, the sensory input (R) is temporarily stored as an iconic memory by the so-called ***image demons***. Then 28 ***feature demons*** each sensitive to a particular feature, like curves, horizontal lines, and so forth start to decode features in the iconic representation of the sensory input.
- In the next step, all representations of letters with these features are activated by ***cognitive demons***.

To understand this process better, how do we go from features to, letters to you know words eventually, a model was put forward by Oliver Selfridge sometime back, who postulated that a collection of small components combined together could allow the machines to recognize the patterns. Now, these small components, or demons as Selfridge referred to them, would record basically events as they occur, and they would recognize patterns in those events.

Say for example, when I am reading how many times I came across a curved line, a horizontal line or a vertical line or a slanted line. Now, in Selfridge's model, which is known as the pandemonium model, the sensory input is temporarily stored as an iconic memory by the so-called image demons. Then there are 28 feature demons which are each sensitive to a particular feature, like one demon can be dedicated to processing straight lines, the other demon can be dedicated to processing curves, the other demon can be dedicated to processing slanted lines, there might be a demon that looks for whole circles, etcetera.

All of these things basically need to be coded, actually decoded from the iconic representation of the central input. So when you reading something there are these features demons that are encoding each of these or decoding each of these single features. In the next step, what will happen is that all representations of letters with these features will be activated by what are referred to as the cognitive demons.

(Refer Slide Time: 27:46)

- Finally, the representation that best matches the input is selected by the *decision demon*.
- The pandemonium model has been criticized because it is a strictly bottom-up model and does not allow for feedback (top-down) processing, such as in the *word superiority effect*, i.e. humans are better at processing letters found in words than letters found in nonsense words or even single letters.
- In 1981, McClelland & Rumelhart put forward a computational model that has been important for visual letter recognition.
- In this model, there are three levels of representations (a) a layer for the letters of words, (b) a layer for the letters and (c) layer for the representation of words.
- An important feature of this model is that it allows for the top-down influence of information from the higher cognitive levels to influence earlier processes that happen at lower levels of representation.

Finally, the representation that best matches the input is selected by the decision demon. Which is basically what it is doing is, speaking up the representation from the mental lexicon and it is matching with the representation that represent in the input, again just matching it and basically saying okay, this the word that I am reading. And so I have stored the sound of it and I have stored the meaning of it and I know this word. Okay?

Now, in the pandemonium model there has been criticized because it is strictly a bottom-up model, which basically starts from the features and then moves towards the decision process. Now because this word does not really allow for things like feedback, it is less equipped to describe or explain findings like the words superiority effect. Okay? The idea is the word superiority effect is basically that humans are better at processing letters found in words, than letters that are presented in nonsense letters strings or even in isolation.

This is basically I think one of the experiments I have probably talked about earlier was done by (())(28:51), wherein he asks people to remember letters or strings of letters or basically match them to some kind of test input and he found out that people were best at doing this task when the letters were presented as embedded in parts of words.

Say for example, you have to recognize the word the letter O and letter K, I will remember it better if it presented in a in a word or thing like fork or frock for that matter, versus if it was just presented by itself or you know embedded in a line of axis. In 1981, McClelland & Rumelhart going this further in 1981, McClelland & Rumelhart put forward a computational model that has been important for visual letter recognition. Okay?

So, in this model, basically there are again three levels of representations a letter level of you know a feature level, a letter level and a word level. So, there is a bit of a typo in this in this slide. The first says layer for the features of letters. Then there is a layer for letters and then there is a letter layer for words. Now, an important feature of this model is that it allows for the top-down influence of information from the higher cognitive levels to influence earlier processes that happen at lower levels of representation.

(Refer Slide Time: 30:26)

- The model is in contrast to that of Selfridge's model, as it allows for both top-down and bottom up flow of information whereas the latter only allowed bottom up flow of information.
- Another difference between the two models is that in the former, processing can take place in parallel such that several letters can be processed at the same time, whereas in the latter one letter is processed at a time in a serial manner.
- The empirical validity of a model can be tested only based on real-life phenomena or against physiological data.
- McClelland and Rumelhart's connectionist model does an excellent job of mimicking reality for the word superiority effect. The word superiority effect can be explained in terms of the McClelland Rumelhart model, because the model proposes that top-down information of the words can either activate or inhibit letter activations, thereby helping the recognition of letters.

Now, this model is in contrast to that of Selfridge's model, as it allows both top-down and bottom-up flow of information whereas the later is only allowed allowed bottom-up flow of information. So, this McClelland's model which is a parallel processing model basically did not really allow you know it basically allows both bottom-up and top-down flow of information, whereas Selfridge's model is specifically bottom up model.

Another difference between these two models is that in the former, that is the Selfridge's model processing can take sorry, in McClelland's model processing can take place in parallel such as several letters can be processed at the same time. Whereas in the Selfridge's model only one letter is processed at a time in a serial sort of manner. The empirical validity of a model can only be tested based on real-life phenomena phenomena or behavioural data as you already know.

So if you compare these two models, you will find that the McClelland and Rumelhart's connectionist model actually does an you know it does an excellent job of mimicking the reality for the superior for the word superiority effect. The the word superiority effect basically again can be explained better in terms of the McClelland and Rumelhart model,

because the model proposes that top-down information of the words can either activate or inhibit letter activations, by thereby helping the recognition of letters.

So, the idea is that you know McClelland and Rumelhart's model basically does a good job of explaining the words superiority effect. They say that once the letters are been picked up, they kind of you know activation is sent at the word level and from the word level there is more feedback, which helps us remember these letters better. Whereas if you are reading letters in isolation, there is no feedback from the word level and hence you will not be able to store or understand that much better.

(Refer Slide Time: 32:33)

Neural Substrates of Visual Word Processing

- It has been known for sometime that lesions of the occipitotemporal regions in the left hemisphere, can give rise to *alexia* i.e. a disorder in which patients cannot read words, although other aspects of language were preserved.
- In early PET studies, Petersen & colleagues (1990) contrasted words with non-words and found regions of occipital cortex that preferred word strings. This area was named the *visual word form area*.
- Later studies using fMRI with normal participants by Gregory McCarthy and colleagues (Puce et al., 1996) contrasted brain activation in response to letters with activation in response to faces and visual textures. It was found that regions of the occipitotemporal cortex were activated preferentially in response to unpronounceable strings.

Let us talk a little bit about the neural substrates of visual word processing. Now it has been known for sometimes that lesions of the occipitotemporal area specifically the left hemisphere, can give rise to alexia. Alexia is a disorder in which patient cannot read words, although other aspects of language are preserved. In early PET studies, Petersen & colleagues contrasted words with non-words and found regions found that regions of the occipital cortex that were that preferred reading word strings. So there were regions in the occipital cortex that only light-up or activate when word strings are being read.

This area was referred to as the visual word form area. Specialized area that processes script now later studies using fMRI with normal participants which were performed by Gregory McCarthy and colleagues basically contrasted activation in response to letters with activation in response to faces and other visual textures. It was found that regions of the occipitotemporal cortex were actually activated preferentially in response to unpronounceable strings, as oppose to you know visual face, visual textures or faces.

(Refer Slide Time: 33:51)

- In a combined ERP and fMRI study that included healthy persons and patients with callosal lesions, Cohen, Dehaene and colleagues (2000) investigated the visual word form area.
- In their study, as the participants focused on a fixation cross, a word or a nonword was flashed either to their right or left visual field. Nonwords were consonant strings incompatible with French orthographic principle which were impossible to translate through phonology. When a word was flashed on the screen, they were to repeat it out loud, and if a nonword flashed they were to think “rien” (meaning nothing).
- The ERP indicated that initial processing was restricted to early visual areas contralateral to the stimulated visual field. Further activations revealed a common processing stage, associated with the activation of a precise, reproducible site in the left occipitotemporal sulcus, part of the visual word form area, which coincides with the lesion site that causes pure alexia.

Now, in a combined ERP and fMRI study that included healthy persons and you know patients with callosal lesions, Dehaene and Cohen, Dehaene and colleagues investigated the visual word form area. What did they find? In their study, as the participants fixated on a fixation cross, a word or a non-word was flashed either to their right or the left visual field. Non-words were consonant strings basically incompatible with French orthographic principle which were impossible to translate through phonology. Now words are just non-words are typically just assemblage of strings of letters which are not pronounceable and will typically not have any meaning.

Now, when words were flashed on the screen, they were to repeat it out loud, and if a non-word was flashed they have to say think "rien". Rien basically means nothing. The ERP indicated that the initial processing was restricted to early visual areas that were contralateral to the stimulated visual field. Further activations actually revealed a common processing stage, associated with the activation of a precise, reproducible site in the left occipitotemporal sulcus, part of the visual word form area, which coincides with the lesion site and and basically causes pure alexia.

(Refer Slide Time: 35:23)

- Later studies, showed that this activation was elicited visually elicited, only for prelexical forms, but invariant for the location of the stimulus or the case of the stimulus.
- Finally, the researchers found that the processing beyond this point was the same for all word stimuli from either visual field – a result that corresponds to the standard model of word reading.

Now, later studies, have actually shown that this activation was elicited was visually elicited, only for prelexical forms, but invariant but is basically invariant for the location of the stimulus or the case of the stimulus. So, wherever in the left visual field or the right visual field or in small case or upper case. So, finally, the researchers were able to found out that processing beyond this point was the same for all kinds of word stimuli from either visual field a result that corresponds to the standard model of word reading.

Remember we are talking about that once you have done the lexical axis, once you have done the lexical selection, the modality specificity goes then it is basically for all kinds of words, what you will do is you will activate the access the word form activate the you know the form of the word logical or (())(36:13) and you activate the meaning and then integration and the processes take place.

So, in this lecture we talked about language comprehension, basically the initial stages of both spoken language and of written language. Thank You.

(Refer Slide Time: 36:30)

References

- Gazzaniga, M. S., Ivry, R. B., & Mangun, G. R. (2014) *Cognitive Neuroscience – The Biology of the Mind*. W W. – Norton & Company.