

This unedited transcript of a continued webinar is provided in order to facilitate communication accessibility for the viewer and may not be a totally verbatim record of the proceedings. This transcript may contain errors. Copying or distributing this transcript without the express written consent of continued is strictly prohibited. For any questions, please contact customerservice@continued.com

AB Sound Processing and Front-End Processing

Presenters: Rebecca Lewis, AuD, Jacob Sulkers, M.Cl.Sc.

Sam.

Foreign.

Hello everyone. Thanks for joining today's Virtual Intro Workshop series. Today's topics are AB's front end processing and sound processing. My name is Katie Vaden. I work in education and training here at Advanced Bionics.

So we're very pleased that you could join this Virtual Intro Workshop series today. And as a reminder, please sure to download any resources associated with today's course. We have a handout for you which you can download and take notes if you'd like. We have the slide content available and any supporting clinical resources and research for you.

We have some housekeeping so before we get started, here are the risks and benefits for today's course.

Here you can see the three learning outcomes for today focused on AB's front end processing and sound processing technologies.

And the disclosures for today's presenters and hosts are included on this slide.

At ab, we really value your feedback and because of this we've created a short survey that will take you under a minute to complete. It's very brief. Three short questions. We'd love it if you would go ahead and use the QR code or the link that you can find in the chat to access the short survey and have it open so that it's easily accessible once the course is complete.

And with that, let's talk about how we'll spend our time together today. We're honored to be able to work with the following professionals as we created content for this training. So first we will hear from Becky Lewis, the CI Audio Ology Director at Pacific Neuroscience Institute. We will hear from J. Jacob Sulkers. He's the coordinator of the Surgical Implant Program at Health Science center in Winnipeg, Canada.

You'll hear from myself as well as my colleague Lori Hambrick in Education and training.

Now as a reminder, I did want to mention that we've offered this training experience to live so that you would have the opportunity to ask your questions in real time via the Zoom Q and A. We can certainly are going to try to get as many possible answers to you that we can, so feel free to type answers into the Q and A. We hope that you'll be able to learn from other attendees questions as well and participate and engage with the training content via the Zoom polls that we have planned for you and all of this allows you to increase your learning through participation. We have recorded the instructors because they're coming to us and engaging with us from all different time zones, so they were able to record their content at a time that was convenient for them. But we at Advanced Bionics are here with you to execute the course and answer any questions that you may have.

So now it's time to get started and I'd first like to kick things off with Dr. Becky Lewis, sharing information about front end sound processing technology offered by advanced Bionics and sharing her clinical experience. Dr. Lewis, I'll hand it off to you. Hi, I'm Dr. Rebecca Lewis or Becky Lewis and I'm the audiology director of the Pacific Neuroscience Institute Adult and Pediatric Cochlear implant program. And I'm grateful to be here today to talk to you about the front end processing of advanced Bionics cochlear implants. I've been working with advanced bionics products for almost 20 years, so I'm grateful to be here today.

So have you ever used a Google Translate or another translation application? If you are in the live webinar, please type your answer into the chat.

Personally, I've used Google Translate as like an instant menu decoder. So when traveling to other countries making sure I'm ordering the food that I actually want to eat. So you can think of the cochlear implant system and the sound processing pathway as a translator that changes an acoustic signal to an electric signal. Sound processing takes the sound from the world around the recipient and translates it into an electrical signal that can be understood by the recipient. Today we'll break down the sound processing pathway, how it captures, composes, details and delivers sound as a signal that AB recipients can understand.

By the end of this session, you'll have a deeper understanding of AB sound processing technology, including the ability to identify at least three three technologies, front end and or sound processing that are unique to ab. Explain how autosense benefits your patients and recommend sound processing features that provide improved performance for your patients.

So the microphone modes, Comfort features and AutoSense OS operating system work together to meet the recipient's listening needs to provide a powerful hearing experience. And we're going to dive into these now.

So there are four microphones on the sound processor. There's the headpiece mic, the rear processor mics which work together, and the T mic. And what's really special about AB is and Phonak is that they collab they calibrate all the microphones in the same way to ensure alignment between the ears and devices.

Take a look at what happens when a recipient uses the settings shown here. The setting that the programming clinician selects in Target CI will determine which CI microphones are in use. So here's

the omnidirectional the omnidirectional picks up sound from all directions with equal gain. Front processor microphones are utilized in this setting.

For the T mic it utilizes an omnidirectional microphone that's placed at the opening of the ear canal to take advantage of the pin effects for improved speech understanding in quiet and noise and for it to sound more natural. The real ear sound uses the directional microphones designed to simulate the pinna related directional cues that are normally captured by the T mic.

And fixed directional. The fixed directional microphones are using a beam former to improve hearing and noise. It utilizes the front and rear processor microphones for to belt the beamform and the null is at 180 degrees pointed towards the back so that your hearing is focused towards the front.

Ultra Zoom and SNR Boost. So Ultra zoom combined with an additional adaptive phonak algorithm to improve speech understanding and noise by providing additional attenuation of sounds from the back. This is an automatic part of autosense OS and it's selected microphone mode is for the speech and noise program Stereo Zoom, the bilateral beam former that uses four microphones to create a narrow focus for one to one communication and loud noise. And this is the automatic program. This is the selected microphone mode for speech and loud noise program.

And something to note is that this is only for bilateral bimodal recipients. It is using four microphones so from both ears to create that special beam form.

Great. So Sound Relax, wind block and Echo block serve to provide listeners with a comfortable listening experience no matter what the environment. So Sound Relax is an automatic feature that will attenuate sudden loud sounds to provide a more comfortable listening experience. So if dishes were to crash in the sink or a fire truck to roll by, it would be more comfortable. Wind block obviously reduces wind noise, so this is great for outdoor settings or golf.

And then Echo block improves comfort and ease of listening in very reverberant rooms. So like a classroom situation. For example, in the past it was common practice to provide different programs for different listening situations. So think about your cochlear implant recipients or if you haven't seen any cochlear implant patients yet, your hearing aid recipients. What percentage of your recipients want multiple programs?

And here's the important part, they actually use them effectively. If you're on the live webinar, go ahead and type them in the chat.

Great. Thank you for your participation. My experience is that patients like a set it and forget it approach and they don't want to fiddle with programs as often.

So Autosense OS automatically adjusts to select the best microphone mode and comfort features to ensure that recipients hear their best as they change environments throughout the day. Let's take a more in depth look at how AutoSense OS works.

So many recipients have described the inconvenience of changing manual programs. Even those who want manual programs have difficulty remembering to change programs and selecting the correct program from different environments. Providing a solution that. Sorry. That automatically manages programs for the recipient is very beneficial.

And this is a recipient who has really benefited from this technology.

Marvel is marvelous. I absolutely love it. Besides autosense, which is great, the Bluetooth streaming is absolutely amazing. There is a music program that automatically comes on when you're streaming music and it makes the sounds. The instruments come through loud and clear.

I can understand lyrics really, really well. The sound quality is fantastic. And also it's the same with taking phone calls. It's optimized for speech perception. When you take phone calls and you all only have to just press a button on the processor and bam, you're right on the call.

Doctor Lewis, thanks for sharing this video with us. What do your patients report as some of the benefits of using autosenseos?

I would say that my patients report no delay in program changes. They don't even notice it. And in fact, one of my patients is a teacher and she goes from a lot of different listening environments. She's in a high school, so she's going to cafeteria, the gym, she coaches soccer outside. And then of course she teaches in the classroom.

And historically she's had to have maybe four to five programs that she shifted through throughout the day. And now she doesn't need to worry about her processor settings at all. She just goes throughout her day and focuses on her teaching. Oh wow, that's amazing. Things have changed a lot in the worlds of both cochlear implants and hearing aids and making things easier for our recipients.

Thanks so much for sharing information about your patients. Will you continue with the presentation and explain how Autosense works? Of course. Thanks.

So AutoSense identifies recipients environments by analyzing incoming sounds every 400 milliseconds and classifies sound into predefined sound classes. AutoSense uses artificial intelligence to calculate the statistical probability that the recipient's listening environment matches 1 of 7 predefined sound classes and activate the appropriate sound class or blend of settings.

So, with multiple features available with AutoSense, the classifier can create over 200 distinct possible settings based on the recipient's specific listening environment. AutoSense classifies inputs as either acoustic or streaming. As you can See [here](#).

So acoustic environments are classified into seven different sound classes. In addition to classifying acoustic environments, AutoSense OS is the first classifier in the industry to classify streamed audio signals. So input is classified as media speech or media music, and the appropriate parameters are adjusted to optimize the listener's audio streaming experience. So as the recipient goes about their day and moves from a quiet house, to listening to a podcast in the car, to streaming music while on the walk, AutoSense can classify these unique types of environments and automatically optimize settings.

So depending on the listening environment, AutoSense features will blend in and out. The core of AutoSense OS is based on the seven programs that you see [here](#).

The three exclusive programs in the top section, Speech and Loud noise Music. Speech in the car are exclusive classes, meaning that these classes cannot be blended with others and will be automatically selected when the recipient is in one of those specific environments. We know that many environments are complex and don't fall into just one category. So there are four non exclusive acoustic programs that may be automatically activated, either exclusively or as a blend, when it's not possible to find complex real world environments by one acoustic classification.

So let's take a look at some examples of the classifier changing as a recipient changes environments. So in this situation, we're at the zoo. There's a lot of activity happening for this family. There's kids running around, there's a zoo train running by, there's lions and tigers roaring. Looks like this family's feeding some llamas.

So in this case, the classifier chooses a blend of speech and noise and comfort and noise for this complex environment.

Now we see a recipient in a new restaurant in town. On an opening night, in the restaurant, there's tables that are packed. The restaurant is open so diners can watch chefs as they work, and our recipient is having a conversation at the table with friends. Think about which program the classifier might choose.

So while in this very loud, lively restaurant environment, the classifier chooses speech and loud noise.

At home, while we're relaxing and listening to music, the classifier chooses the music program.

And in the classroom, there's a lot going on. The teacher is talking, children are shuffling papers, other children are chatting while walking down the hall to change classes. And there is some reverberation on the concrete walls. So what program or blend of programs do you think the classifier will choose?

The classifier chooses a blend of speech and noise and comfort and echo.

Autosense OS is the first in the industry to classify the difference between streamed media that is Speech and streamed media that is Music. So these streamed inputs are optimized for speech clarity during streaming of conversation, dominant programs like audiobooks or podcasts, and sound quality when streaming music rich programs. So streaming signals are delivered in stereo to bimodal and bilateral recipients. This results in quality streaming experience when watching tv, listening to music or streaming podcasts. Additionally, streaming programs can automatically detect the presence of streamed input from Bluetooth phone calls, so for the Phonak Partner mic or the Roger microphones, making it easy for recipients to use a variety of streaming devices.

So as recipients start and stop streaming, the processor automatically or seamlessly transitions into and out of the streaming programs without the need for the recipient to push any buttons or make any changes to their hearing instruments. And in my opinion, this is a big game changer because it allows patients that normally would be too intimidated by this technology that required multiple steps and now they can actually just use it because all they have to do is turn it on and it'll automatically switch the mention as mentioned, the Sky CI processors use AutoSense Sky OS 3.0 developed specifically for the listening environments children experience to ensure that kids and teens have the optimal listening experience in their unique listening situations. So this might include classrooms, libraries, the playground and when listening to music.

Let's take a look at some data about Autosense.

So this figure shows speech recognition scores in quiet and in noise with AutoSense OS 3.0 set to on and off. So as expected, performance was equivalent in quiet and showed performance and showed improvement in noise. So here you see. Sorry, Here you see 46% improvement in noise with AutoSense on showing that the benefit provided when AutoSense 3.0 adapts to the user's listening environment. And to note, this study was done at a plus 5 signal to noise ratio.

So this is a that's a very challenging listening environment.

A multi center study was completed by independent researchers at five different research centers. This study evaluated speech recognition results for AZ bio and quiet and in noise using the previous Naida Q90 technology with AutoSound, the Marvel CI with AutoSense and when using the Marvel CI set in omnidirectional mode. So on the far left we see similar performance in quiet with auto Sound on Naida Q90 in blue and AutoSense in the Marvel CI in green. What's most exciting is the results found when the listening in noise. So in the middle graph this shows the statistically significant improvement in noise with Autosense on the Nanita M90 in green over the performance with Auto Sound with the Naida

CIQ 90 in blue.

Not only did we see a significant improvement during this study, but recipients have provided very positive feedback about AutoSense. Even those who didn't prefer automated features like Ultrasoom in the past liked autosense. So the graph on the right compares Marvel CI with Autosense ON versus using an omnidirectional mic. So Autosense not on we see a statistically significant improvement in noise with autosense in Marvel CI versus omnidirectional mode. Individuals perform much better in noise when using AutoSense technology.

Outcomes from the study really demonstrate the effectiveness of AutoSense OS and its ability to automatically identify challenging listening environments and optimize settings in the manner that improves speech recognition ability. I would also like to share a quote from one of the research subjects who describes their experience with the Marvel CI. So they wrote, I feel like my hearing with the new model was better. Car noise was reduced. Wind noise did not seem to be as much of a problem.

I was in several restaurant situations and felt I could hear well. I tried it in a small group setting and always had to use my Roger microphone and I tried it without the Roger and was able to hear all of the conversations. That was a really nice experience.

Other studies performed in Geneva and Hanover looked specifically at streaming capabilities within the AutoSense OS. So you can see the results on the Geneva study here showing that patients significantly preferred streaming on the phone with use of their Marvel CI devices compared to previous generation Naida Q devices on all measured tests. The results of the Hanover study showed that patients preferred listening to music in both ears compared to one ear and that streaming to both ears was preferable to streaming with both ears in the free field condition. The automatic detection of streaming devices through AutoSense OS makes this process seamless and simple for recipients.

So here are the results of a study evaluating AutoSense sky os. So nine pediatric bilateral CI recipients were assessed using pediatric AC biosenses in quiet and in varying noise levels in an omnidirectional program and in sky os. So in gray we see the results from the Omni program and in blue the AutoSense Sky OS program.

The investigators saw a 46% increase in speech recognition between programs and I just want to highlight that this was done at a negative 5dB signal to noise ratio.

Additionally, there was a significant increase in listening and clarity ratings with the sky os.

So you can see that the combination of multiple microphone modes, comfort features, and sophisticated AI technology. And AutoSense OS allow for an incredibly powerful sound processor from advanced bionics. And in the next portion of today's webinar, you'll learn more about the sound processing technology that occurs to transform the acoustic signal into an electric signal.

Dr. Lewis, thanks again for doing such a great job of explaining AutoSense and helping us understand how it's beneficial to patients. Before we move on and hand this over to our next presenter, would you share how you explain what your patient should expect when using AutoSense? Sure. Honestly, I mostly tell patients that they might not even notice AutoSense working. And that's the point.

It's doing the hard work for them to make hearing feel natural and effortless. Or I'll tell them it's sort of like cruise control for your hearing. You're constantly, it's constantly working in the background so you can focus on living your life and not fiddling with the processor. That's fantastic. I love that cruise control analogy, especially now when cruise control is even fancier than it used to be, where it will actually adjust to the car in front of you.

So I think that that's such a great analogy for autosense. Yeah. Thank you so much for having me. All right, thank you, Dr. Lewis. Now, in week one, we discussed an independent study from Vanderbilt that highlighted how audiologist recommendations can shape patient choices.

And that technology is one of the primary factors that drive CI candidates choice and how enhanced patient education could improve their confidence in their decision making and potentially their retention of using a cochlear implant. So considering that technology is important to candidates and that counseling from you as their audiologist is key, what would you say to help a CI candidate or recipient understand the benefit of autosense? Go ahead and type your answers into the chat.

Okay, I see a few answers coming in and you're right. Set it and forget it. That Autosense is easy to use. It manages the settings for patients in every environment. All of these are good reasons why autosense is going to help your patients hear well in their everyday environment.

I really appreciate Dr. Lewis sharing information about our front end sound processing, including her clinical experience. And now we're happy to have Jacob Sulkers provide information about sound processing technology from advanced bionics. My name is Jacob Sulkers. I'm the coordinator of the Surgical Hearing Implant Program in Winnipeg, Manitoba, Canada. We currently serve around 6 to 700 cochlear implant patients and another 2 to 300 bone anchored implant patients and I'm going to be walking us through sound processing today.

So the sound processing pathway is the step after the front end processing that you've already reviewed. So I'm going to be talking about some signal analysis, the mapping and a little bit of stimulation to get us on our way.

Really what I mean by sound processing is how is sound? So how are acoustic signals picked up by the device and changed into electrical pulses that are then delivered to the electrode array and onto your

hearing nerve and up to the brain? So I'm going to start off by talking about various aspects of sound.

So I'm going to start off with timing cues, so temporal details. It's information on the time of the acoustic signal coming in.

So in this whole processing pathway, I am starting with the signal analysis phase. So we're going to start by talking about timing cues. We look at this waveform, it is the sentence, where were you last year, Sam? So on the X axis it's time and along the Y it's the amplitude of the signal. Now why are timing cues important?

It gives us a lot of information about the actual acoustic signal that we're trying to listen to. So timing cues can give you information on word and syllable structure. So long words, multi syllable, short words. It gives you information on the prosody or the stress. So how you're talking about, it gives some information on the manner of articulation.

So is the sound you're listening to a stop? Is it voiced or unvoiced? So is it a B or a P? It gives information on voicing. So as an F versus an S, the strength of those sounds and the length of them.

It also gives you some information on nasality. So M versus N. So these are all clinically important things to be aware of because they do relate to clarity and naturalness of speech. Now what I'm going to spend a little bit more time discussing is how stimulation rate relates to temporal detail and naturalness and clarity of speech. So when you look at a normal hearing ear, it uses both the overall envelope, so the shape of sound over time and the fine structure, so the very, very quick moving information inside to make sense of a speech signal. So what we want to talk about is the ability of a cochlear implant system to pick up sound accurately.

So the sampling rate, picking up sound at a quick enough rate to accurately represent it, and then the stimulation rate. So how quickly is that sound presented to the auditory nerve? So with a low

stimulation rate in a conventional system, you can see the blue bars are pulses of sound, and they give you a general outline of the acoustic signal. So if the black is the actual envelope, it does give you some of that sound. But to have accurate, more natural sound quality, you want to be able to provide not just the overall envelope of the sound, but some of the fine structure within it.

So if we look at a system using a higher rate of stimulation, what you can see is the little stimulation bars are really small, and they really reflect a lot of quick peaks and valleys in that speech signal. And it's going to give you a better version of the actual acoustic signal. So you're going to get some better clarity and better naturalness with that higher rate of stimulation.

So the reason temporal details are so important is they do give you an accurate representation of the sound, and it's going to result in better hearing performance.

So that's one aspect of sound is the time. So basically, sound over time, how is the system capturing it and delivering that to your nerve? The next aspect of sound that I want to talk about is spectral details. So it's frequency information in the sound processing pathway. Again, it falls in that same signal analysis section.

So where it's looking at sound over time, and it's looking at the frequency makeup of that actual sound, using the same sentence as before, we can see a spectrogram. So where were you last year, Sam? Across the X axis is again time, this time on the Y axis, it's frequency. So from low frequency up to high frequency, the hotter the color, so the more red the color, the more information there is in that frequency region. And you can see down low, you get some sort of, it's like the, the fundamental frequency and then some formants up over top going across the sentence.

Where were you last year, Sam? So up at this high section, the high frequency is the S of Sam. The importance of spectral characteristics. So the importance of accurate frequency information comes

down to a lot of the clarity of speech. So vowel identification relies on clear ability to hear the frequency information involved in those vowels, the pitch of vowels.

So whether it's an oo or ooh, it's the same vowel, the same structure, but different pitch, the overall intonation of your voice, the consonant place of articulation. So if it's a front, a mid or a back sound, so is it a T, is it a T or a K at the back of your throat? The spectral characteristics also provide a lot of information on the depth and the naturalness of voice.

So how advanced Bionics deals with frequency information and how it provides some good depth of that information is through the use of current steering. In a typical conventional CI system, stimulation occurs at the location of the electrodes. So the black bars here are physical electrodes. And you can see that when you stimulate sound, you can stimulate at one electrode or the other and it will send stimulation to the neural elements that are around that physical contact. With the advent Of High Res Fidelity120, Advanced Mechanics introduced current steering.

And what that does is it allows for the perception of more channels of sound in between the physical electrode points. So again on this slide you can see the black bars are physical electrodes. But you're able to provide current in between the physical locations. So you're able to steer current and reach discrete so neural elements in between the physical contacts. The reason it's called Fidelity 120 is the system has 16 electrodes which results in 15 pairs along the entire array.

In between each of those pairs there are eight potential stimulation sites. So let's say between electrode one and electrode two there's actually eight little steps along the way. And you do the math and it works out to 120 virtual channels of sound. Now this diagram or this video gives you a representation of how current steering works. So the black squares at the top are physical electro contacts and these little fellows at the bottom are the neural elements.

So you can see as stimulation is delivered from one contact to the other, you can actually steer the location where it's hitting those neural elements to provide even more channels of sound than there are physical electrodes.

In terms of what does that mean and how does it actually sound? You can see on the left the conventional sound processing. So there's really discrete levels at all the electrode contacts. And in the high res example you can see there's much more frequency information spread across what would theoretically be that person's cochlea. And in terms of the sound quality, you can notice quite a striking difference.

So I will play the sound and let you listen. Windows provide light and air to rooms. Windows were once covered with crude shutters. Later, oil paper was used for window panes. Windows provide light and air to rooms.

Windows were once covered with crude shutters. Later oil paper was used for window panes. Glass windows first appeared in ancient Rome.

Huge difference, right? So the first one sounds quite robotic. I've listened to it enough that I know what it's saying, but it takes a little bit more effort to try to focus and understand what it's saying. Whereas in the second example it sounds a lot more natural. For me it's easier to understand.

And one of the biggest differences is in the second one there's a baby crying, whereas in the first one I don't pick that up at all. So the increased spectral resolution, so the more frequency information you can provide, the more natural and clear sound can be when it comes to listening to music. And what kind of differences are there? Again, just even visually comparing the conventional and the high res, you can see there's much more diversity of sound with the current steering. And we'll give this music a little listen to and see if you can notice a difference.

Welcome to the Hotel California.

Cool. So for me, I. I mean, it's a huge difference. I could, I probably would have been able to tell you what sound it was. With the conventional processing, I could at least recognize enough of the voice to tell you what it was. But when it switched over to the high res, the 120, the depth of the sound, the naturalness of the sound, the instrumentation and the vocals are all much more natural to my ear.

And that's all thanks to spectral information. So frequency information. So the better and more detailed frequency information we can provide, the more natural and easier to understand the information will be.

So in some of the spectral details, it contributes generally to sound quality and speech intelligibility. It's going to help patients hear better in noise. The more accurately you can hear and differentiate sound, the better you will do in noise. It's going to help with appreciation of music, identification of environmental sounds, and understanding of tonal languages.

So a very quick knowledge check based on the information you've heard so far. How is advanced bionics able to steer current between electrodes?

I'll give you a minute to type your answer into the chat, then we will all agree it is a B answer. It's independent current sources. So the ability to stimulate each electrode independently and turn them on and off at varying degrees allows you to steer that current and come up with far greater virtual channels than physical electrodes.

The third aspect of sound that I want to discuss today is intensity. So the loudness. So we talked about temporal cues, so the sound over time. We talked about spectral details, which is the frequency

information over time. Now we're going to talk about loudness over time and how do we deal with that?

So back to our sound processing pathway. This is falling in more into mapping. So how are we actually taking the acoustic signal and mapping that on to the client's electrical dynamic range into their system?

Going to start by talking about input dynamic range. So if we think of a Typical functioning ear, it has a very wide range of sound between what is perceived as very soft and very loud. An average, I think in typical hearing is about 100 decibels. So it's this giant range that encompasses almost all the sounds in our environment. Now, the input dynamic range of a cochlear implant system is how much sound that is captured by the microphones is actually delivered and presented to the system.

So the microphones pick up more sound than is necessarily going to be processed for every single user that we have. So a narrow input dynamic range, which is available in some systems, you can see that the, there are some sounds, especially soft level sounds, that'll be inaudible. There are loud sounds that will be compressed and then inside that dynamic range there are sounds presented quite faithfully. Now, the wider you make the input dynamic range, the fewer sounds are made inaudible. So there are still some very low level background sounds that'll be inaudible and there are loud sounds that still get compressed.

But with the wider idr, there are fewer sounds inaudible and it's going to be picking up more sound for patients. Now the reason this is important is it's, it's kind of the whole point is we're trying to provide patients with the best, most accurate sound we can and we want to minimize those sounds that are inaudible to them. So with a wide idr, soft sounds are going to be more audible. So it's going to improve perception of important speech sounds. Because as we talk, not all speech sounds are the same level.

So our voices go up and down, they're dynamic. If a person's voice is sort of dipping below that input dynamic range, those sounds will not be presented to their auditory system and they will not be able to hear them. So having a wide IDR really can go a long way in improving access and giving patients the best option to hear clearly.

I would say clinically, one of the questions we get is what, What IDR should I use? Is wider better? Is narrow better? There was some research by the group at Hearts for Hearing looking at IDR and its relation to speech perception of soft sounds. The short answer is the default of 60 DB or AB devices is appropriate and will result in the best performance for most patients.

Less than 60dB is not ideal, making the IDR wider than that. So going up to 70 or 75 or 80 is appropriate in some cases. And I know that for me personally, just based on my own clinical experience over the past 15 years, we tend to default at around a 70dB IDR for new patients. But there are reasons and ways to change that over time that we don't need to get into today.

That is everything to do with sound thus far. So we've talked about time, we talked about frequency and we talked about intensity over time. So the loudness over time. And now I'm going to talk a little bit about signal analysis. So actually cleaning up the sound that is picked up.

So if we think about the sound processing pathway, we're going to go back into signal analysis. So this is where we are dealing with cleaning up the sound as much as possible to make it as understandable as we can for our patients.

Clear voice is a noise cleaning feature. So its job is to improve the signal to noise ratio on each channel within the system. So it's trying to separate distracting noises from important speech to make communication easier. It is proven to be superior in noise, which I'll talk a little bit about in a few slides from now.

Now again, we're going back to looking at a waveform over time. You can see the red signal in this picture is noise. It's a steady state background noise. It goes on for a few seconds and then a speech signal will kick in. So what I want you to listen to is the difference with clear voice versus clear voice off to see if you can tell a difference.

So we'll start with no clear voice.

We faced four hours of steady work. So for me, I can hear the person. The, the noise is there and sort of interrupting in the way, but you can, I mean they're definitely sort of pretty equal in level. Now if we turn clear voice on, what you're going to notice is you still hear the noise. So the noise will kick on in the same.

There's an analysis window where clear voice is assessing the sound and then you'll hear clear voice basically kick in and then see if you can notice how it affects your ability to understand the speech.

We faced four hours of study work, right, so that you hear the noise in and then it really kind of drops down and it lets that voice shine through. It doesn't completely eliminate the background noise. So you still do get a sense of being in space and being aware of your surroundings. But it improves that signal to noise ratio, allowing for better understanding of the voice of speech.

I'm going to spend a little bit of time talking about why clearvoice is superior. So if you look at this graph, it's showing you the percent correct for speech and noise for both steady state noise and multi talker. Babble. It's just been aggregated together. The control in gray is Clear Voice disabled.

And with blue, that's Clear Voice turned on. And what you can see is that for clear the various clear voice settings for low, medium and high, there is an improvement in the percent correct. With Clear

Voice enabled and the presence of background noise. In terms of the medium and high settings, it's actually a statistically significant improvement in speech recognition. And if you look at the actual values, it is a clinically significant difference.

Like it's at least a 10% or more change with the presence of background noise. This is a first of its kind. So to be able to demonstrate improved speech and noise understanding with a digital noise reduction feature is fantastic. And AB was the first one that actually demonstrated that ability over time. What it does is clearvoice has shown a clear benefit for providing clear speech in clearer speech in the presence of background noise.

Thank you for describing that research to us and making it so clear. As you mentioned, ClearVoice is available in both low, medium and high. So which setting do you routinely choose for your patients? We tend to start with the default value, which is medium. I find that for most patients this tends to provide the benefit noise that you're looking for without any noticeable detriment in noise.

Occasionally we will have patients that notice that clear voice medium results in what they feel is a bit of a drop in volume and noise, so we may move to low. Occasionally. You do have patients that are really in extreme noise or really dislike any presence of background noise, so we'll often give them an option of going to clear voice high. I would say for us clinically, it's rare that clear voice high is the the setting in their main AutoSense program. If anything, we often will give them a manual noise program with clear voice high.

We do have some occasions where we also completely turn clear voice off, and that tends to be for people who work in noise who need to hear that noise. So if they work in a shop or work with machinery where they they really need to hear steady state noise accurately again, we'll often provide a manual program with clear voice on low or disabled to let them sort of have have the technology when they need it and have it turned off when they don't. Thank you for sharing that. It's really helpful to hear your

experience when we have these three different settings because I think often newer clinicians see that medium is the default and will start there, but then are also curious about, well, when should I change it? And those were really great Examples.

So thank you.

Just before I move on we'll do a knowledge check. So which environments do you think clear voice would be beneficial? So if you use it in a car, at a wedding reception, sitting near an air conditioner, if you just type your answers in the chat, I'll give you a minute.

And yes, the answer is all of the above. Any presence of background noise, Regardless of the situation, we have found that clearvoice can really be beneficial to make it more comfortable and to provide better clarity of speech.

The next feature that I'd like to discuss is soft voice. Soft voice is another sound cleaning feature available in advanced bionics devices. Its job is to improve audibility of low level input. So you can see in this picture it is a SA presented at 33db SPL. So it's a soft voice.

You can see some energy clustered down low with the ah and then the S and back to the R. But you can notice on many of the channels there's just a very low level noisy input. Soft voice is a feature available to clean up that unwanted low level input. So how does it work? You can see that when soft voice is turned on, that low level input is very much cleaned up and allows that speech signal to really stand out. So what it's doing it is removing low level noise just to make it easier to detect low level soft sounds.

So it's again, we're just trying to allow people to hear the acoustic signal as accurately as possible regardless of the volume of the input.

And thank you again for your great descriptions of the technology. When you are working with your recipients, do you leave soft voice on an initial activation? I do. So especially for new clients and new patients, we tend to leave soft voice on again just to provide them with the cleanest sound possible.

Great, thanks for sharing that. And are there any situations in which you are finding that you want to turn it off? Anything you can think of off the top of your head? We have, we have had some patients who have upgraded to new technology who are used to a certain way of sound. So sometimes you'll find with upgrade clients they don't particularly like having soft voice on.

Counterintuitively, every now and then you'll get someone who reports that with soft voice turned on they hear more noise. It's likely that they're actually hearing more environmental noise. Oh, right. But it's a, I would say it's a rare case, but usually upgrades are the times that we, we inform patients about it and ask if they notice a difference on or off. But with new clients we definitely leave soft voice enabled.

That's great. Change is hard for everybody and I think that's one of the biggest, bigger challenges when you are doing upgrades is people get new devices because they want the better opportunity to hear. And often that does introduce new sound processing technology that results in a change for them. So I've heard that as well is that on occasion people are just wanted to. They want the new technology, maybe they want the new features like Bluetooth streaming, but they want the sound to be more similar to what they're used to.

So thank you for sharing that.

The last aspect I'm going to talk about. So we're still in that sound processing pathway. We've gone through the front end processing, through the signal analysis we're starting to map. We actually need to

deliver the sound to the patient's auditory system. So we're going to be talking about paired versus sequential stimulation.

And I know with the students we've worked with and new clinicians we've worked with, this is often a big question of what do you do and what's the difference and how do I start? So I'll do my best to walk us through the differences between the sequential and paired. So currently AB is the only company that offers both. You can do sequential, which is just means that one channel of sound is being stimulated at any given moment in time. AV also offers paired.

So there are two stimulation sites at any given moment in time. They both offer unique benefits to patients. Some patients may have a preference one way or another and we can talk a bit about that as well. And we as clinicians might have our own preferences as well. In terms of options.

You have a lot of so it's good for sequential stimulation, AB offers Optima S, Fidelity 120 high RES S and CIS prepared stimulation. Again you have Optima P, Fidelity 120 P, Iris P and MPS. As a new clinician or as someone programming new patients, you are going to stick with ideally the latest technology. So you're looking at Optoma S. So we discussed Fidelity 120 and the use of current steering earlier and how it can give you much more frequency resolution. Optoma S is based on fidelity120 but is more efficient so it uses less battery power.

So there's if you're going to be using current steering, we would very strongly recommend using the Optima programs.

They are all of these options are available across all of the various electrode types. So you do have it for for new patients and new activations we always start with the latest technology, so we'll start with an Optima program. Patients do have a preference, sometimes one way or another. And that's part of the joy of activation. And early follow up is finding out what provides the most accurate sound for those

patients.

And you mentioned that you start with Optima And a question that I get very, very often is there are both sequential and paired stimulation options. So which one would you start with? S or P? That's a good question. I tend to start with Optima S. So this is mostly honestly from clinical experience.

I have found that most patients have a preference for Optima S. If I think back when I sort of first started this job, we used to give everyone choice and we still do. Typically what I do is I start with the default of Optima S and then we do either a one week or a one month follow up where we'll offer the other version. So if we started with S, we'll offer an Optima P. At that point, some patients do have a very clear preference or they will find one version is much more clear or natural or both. If a patient has a preference, we always go with that preference. So we're not trying to steer them towards an S or a P. If a patient doesn't have a preference one way or another, I tend to leave them on the Optima S program.

Again, that's sort of my clinical bias based on my experience over time. The nice part about these options is switching between them to offer them as different programs or different options for patients is very straightforward. There is some adjustments to your M levels, but it's quick and easy to do. So I tend to start with Optima S but give patients the option of trying a P and then going with their preference over time. Great, thank you again.

Your clinical experience is so valuable. I appreciate you sharing that.

In summary, I appreciate everyone taking the time and what we're really trying to look for is we're talking about a combination of technology from Phonak and Advanced Bionics where the front end processing of Phonak and all their years of experience combined with the sound processing technology of Advanced Bionics. The whole goal is to create effortless hearing. We want patients to be able to put on their device and use it every day with minimal to no adjustments. And that's what the technology is

allowing us to do. Thank you, Jacob.

As we previously discussed, we know that technology matters to CI candidates and recipients. Patients want technology to fit seamlessly into their life. So think about all the technologies from today's discussion which ones would be most important to them. So if you're joining us in the live session, go ahead and type your answers.

All right. Good answers. I see some good responses rolling in. You're right. Patients want to be able to hear their best no matter what activity they're engaged in during the day.

And AB technology is designed to do just that. It's designed to specifically help them hear their best in any environment that they find themselves in.

Okay, I'd like to wrap up by thanking all of you for taking time to join us during today's webinar. And again, as a reminder before you leave, we would love your feedback if you didn't access it earlier. Here is the course evaluation QR code. The web link is in the chat. Before you leave, if you would take just a few seconds to give us your response to three short questions, we would really appreciate your feedback.

I'd like to wrap up by thanking Dr. Lewis and Jacob Sulkers for their time and expertise and for sharing their clinical experience with us today. And thank you all for taking time to learn more about AB technology. We look forward to seeing you next week when we discuss programming and target CI.