

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

Oticon Intent: Audiological Innovations

Presenter: Virginia Ramachandran, AuD, PhD

- [Virginia] Hello, my name is Virginia Ramachandran, I am head of audiology for Oticon, and today I'm gonna be talking about Oticon Intent: Audiological Innovations. Oticon Intent is our new state-of-the-art product, our new premium product, that has just been launched and we are gonna be speaking today about the audiological components of Oticon Intent. So I won't be getting into some of the details about the product per se, such as colors and styles and receiver sizes and so on, those are in some other trainings, but really focusing today on the audiology of Oticon Intent. As I said, I am head of audiology for Oticon and therefore I do receive a salary. I also have a couple of other disclosures: as an author and editor for Plural Publishing, where I have a financial relationship, and non-financial relationship as past president of the American Academy of Audiology.

And my email is on here if you need to get in touch with me at vira@oticon.com. And as a disclaimer as well, the opinions and assertions presented here are the private views of myself and/or Oticon and should not be construed as official or as necessarily reflecting views of the American Academy of Audiology. So the learning outcomes for today are: after this course, you should be able to describe the components of Oticon Intent's MoreSound Technology features, you should be able to explain the benefits of applying the Audible Contrast Threshold to the programming of MoreSound Intelligence 3.0 in Oticon Intent, and you should be able to explain the patient benefits that are accrued from Oticon Intent.

So before I get going into what's new, I do wanna be mindful of the fact that what we do from Oticon's processing isn't just starting from scratch. We don't sort of throw out old things just to put new things. Really our technology is built upon past innovations so that we're growing in a particular direction. We've had the brain hearing concept that we've been working with for many years and all of our innovations are centered around that idea that you hear with your brain and that we have to be really focused on what the brain is doing with sound so that the patient can have a meaningful outcome.

We started back in 2016 with OpenSound Navigator and the idea of an open soundscape, that open sound paradigm, which is designed to really best support those brain hearing ideas.

When we had another major innovation, 2020, with MoreSound Intelligence, we ended up using architecture via a Deep Neural Network embedded on the chip, and the Deep Neural Networks are really designed around how the brain naturally works. And we used that strategy to be able to increase access to sounds, and to create contrast between speech and other sounds. So again, giving the brain access to what it needs to be able to function well. And what we're gonna be talking about today is another innovation on this paradigm, which is the use of 4D Sensor technology, really tapping into user behaviors so that it can better inform the instruments about what the person is intending to do. And so this is not just a strategy of, you know, how do we process sound?

It's a strategy of what does the brain need and how do we give the brain what it needs? So, as I said, we're gonna be talking about Oticon Intent and our tagline for this is, "Engage in life like never before with the world's first user-intent sensors." It's important because really engagement is what we're looking for for patients when we think about all of the challenges that they face and that they come into us with, it's really about relationships, communication at some level, and so engagement is what we're hoping to achieve and we're providing a tool to help enable that via these user-intent sensors. Now what I'm gonna be really focused in on for most of this is our MoreSound Intelligence 3.

0. So we've got different components of the signal processing, MoreSound Intelligence is about cleaning up the sound signal, creating access to speech sounds and contrast of speech sounds with other types of sounds in the environment, maintain access to those sounds but ensuring that speech becomes contrasted with those sounds in a

way that creates a really useful experience to the patient. The innovations that we're gonna be talking about are possible because we are introducing a brand new platform called Sirius. Sirius is enabling many different things, including our Deep Neural Network 2.0, which I'll be talking about in a little while, the 4D Sensor technology, which is something I'll be talking about very quickly here. But then there's also things like the expanded frequency bandwidth, being able to communicate with our new intelligent miniFIT Detect speakers and being able to engage the future-proof Bluetooth LE Audio.

And so we're really expanding both social engagement and digital engagement through these options here because of what the platform enables us to do. So let's dive into really what's the problem that we're trying to solve with this latest iteration. It's a nuance to a problem which is that, you know, we wanna make sure that people have access to sounds in different environments and we personalize our fittings right now. But a lot of the adaptive features of the hearing aids, historically all adaptive features of hearing aids, were based on the acoustics of an environment. And you can have different people within that environment that have different needs over time. So it's not just about programming the device specific to each unique individual's needs, but then it's about how does the device respond to what their intentions are in the moment.

And that's the piece we haven't seen before. And so if you look at this image here, you've got three people, James, Valentina and Dave. So James... Well, let me step back for a moment. They are all in the exact same acoustic environment and so depending on how their hearing aid is programmed each of those devices is going to respond to the acoustic environment. So if the acoustic environment is simpler, we might not engage those MoreSound Intelligence features quite as much as when the environment is more difficult. And so we wanna get the right amount for the right person at the right time. So only depending on the acoustic environment then the individual needs of the moment aren't necessarily captured.

So you have James here, he is walking around and so when he's in this complex acoustic scene he's going to need more awareness of the sounds around him, especially if he's navigating a room. If you have Valentina there, she's in a group with multiple people and so she is following a conversation, but that conversation means that there's turn taking, there's different people who might be picking up the conversation, and so she needs both a level of awareness of the surrounding speakers but also enough focus that she can have meaningful dialogue in that complex environment. Then you have Dave over there. Dave is having an intimate conversation with somebody, just one person, and so in that environment where you're having that high degree of communication need, like in an intimate conversation, you want the most support possible from MoreSound Intelligence.

And so being able to figure out how can we understand what their individual needs are in the moment is what we're really trying to achieve here. And, again, the reason that we're trying to do that is because we wanna support the brain as best possible. And it's not about just being able to hear, or listen, or even to comprehend, it's about being able to communicate and engage. And being able to communicate and engage is a higher level process. When we as humans try to communicate with others we're not necessarily just taking in auditory information. Sometimes we are, but we also may be taking in visual information and looking for nonverbal cues. We may be aware of what's happening around us in the environment, right?

And so we're not just using that. So, for example, you might be able to hear me right now and you're able to see what's on the screen, right? So you're using two pieces of information to better understand what I'm trying to communicate. In the same way the hearing aids can use more information we've historically given it. So they've always had sensors in the form of acoustic sensors within the hearing aids, but by adding in motion sensors and conversation activity detectors, we can add in very useful

information to, again, help inform the hearing aid more about what the individual user's needs are to facilitate communication. So what we're really talking about here is capturing listening intentions using the sensors.

So we've got an accelerometer that has been added in with the Oticon Intent, and the accelerometer is able to pick up information on both head and body movements. Excuse me. So that we can understand when somebody is turning their head, when somebody is nodding, when somebody is walking or running, we can actually take in that information to help inform what the person's intentions are. And the reason we say what their intentions are is because these are behaviors that they automatically engaged in from a communication perspective. And that's what's really key here is that taking and capturing those natural behaviors is really correlated with what the intentions are. If they're trying to walk or run we can infer that they're not trying to have an intimate conversation at the same time, right.

When they are, when there's a conversation happening and they are moving their heads and standing still, we can infer that they are having a group conversation. And so it's really about capturing what their intention is in the moment and hence the name Oticon Intent. So the 4D Sensor technology is then combining information about the acoustic environment, it is picking up conversation activity, so not the conversation itself but just is a person engaged in a conversation, and then it's also picking up body movement and head movement using the accelerometers. And it takes these four pieces of information and then uses this combination of information to tell MoreSound Intelligent how assertive it should be in applying help and noise.

So we can kind of walk through this for each of those listeners that we just saw in the other picture. We saw Dave. As I showed before, he is having a high level of conversation engagement, that's his need. And in order to do that he's naturally gonna behave in a certain way, which is that he's going to be having little or no head

movements. If he does, it's probably a nodding type of behavior. But he's not like, you know, turning his from side to side in order to sure other listeners. And so the hearing aid is able to infer, based upon the acoustic complexity, the fact that there's a conversation happening and that he's not moving very much, that he is intending to have an intimate conversation and therefore a high degree of support is needed from the system.

Then you take Valentina, she's monitoring, that's her need in that moment where she's having a conversation with multiple people. In that case she will be naturally using head movements because it's normal for humans, when they're communicating, to have eye contact with the person who's speaking. And so you're gonna be facing the person who's talking, and so as the conversation partners change over time their head movements will indicate that this is in fact a situation where Valentina would need to be pretty focused but also have environmental awareness around her so that she can follow changes in who is speaking in that conversation. And so it determines that the support need is for both speech clarity and awareness.

And then you have James, who was our person who was walking around in the environment. And you could recall that he has a need that moment of high awareness of sounds around him. Evolutionarily it's really just advantageous that if you're moving around that you're able to hear the sounds around you. And so that user behavior would be that he's maybe walking or running, the accelerometer picks up that information and tells the system that we need a heightened awareness of the surroundings and therefore it's going to prescribe the appropriate level of support and MSI for that. And so you can see this little image here trying to help illustrate that. That over time even the same person in the same acoustic environment may have different needs and so the 4D Sensor technology puts together all of that information about the user's natural communication behaviors to adjust the level of support that's provided by MoreSound Intelligence.

So it's taking a process that was already occurring via MoreSound Intelligence but adding nuance to it so that you get this range of additional support to go increasing or decreasing the level of contrast. And in this case it was measured to be a five dB output SNR range that we can get, sorry, up to a five dB output SNR range within a given acoustic environment. So let me give you yet another example just to try to illustrate this point a little bit more. Here you have the increased environmental complexity on the bottom. And so this is, as you go to the left you have a decreasing signal to noise ratio. So starting out on the right, you have a really good signal to noise ratio in the environment and as you go to the left, things get harder.

What happens in the hearing aid right now is that you have the output SNR enhancement increasing as the environment becomes more complex. So we need more support from the help and noise system when the environment becomes more complex. And in fact it does do what it's supposed to do, and we can measure that by the output SNR of the hearing aids. So the blue line here is what Oticon Intent is doing with the 4D Sensor technology turned off. Okay? So this is how the hearing aid would react just based on the acoustic environment. When we add in the other elements of that 4D Sensor technology, we can increase the range of support based upon the user's needs.

And so now you have... The default would be that middle line, but we've got these other two lines above and below. And you can see that with the 4D Sensor technology, if somebody were engaged in a situation where they needed a high level of awareness, 'cause they were walking around, the 4D Sensor technology can decrease the level of help and noise. And if they needed additional support because they're having an intimate conversation, then it can increase that level of support so that it can be flexible to the user's needs in the moment. So MoreSound Intelligence is really the entirety of all of this processing of cleanup. The sound signal we've always... Sorry, we've had

since Oticon Real, our new wind and handling stabilizer working on the front end of that to reduce wind and handling noise by momentarily halting the input to the microphone where it would have the greatest of the noise sources as needed.

Then we're adding in the information from the 4D Sensor technology. So some of that is coming in from the microphones, the acoustic sensors, so that's the acoustic environment as well as the conversation activity detection, and then we've got the accelerometer sensors to help understand head and body movements. And that is going to steer the dosage of the medication, if you will, that MoreSound Intelligence is providing. So the 4D sensors are part of MoreSound Intelligence 3.0 and they're really responsible for the dosage of help that you're getting. Now the medication, or the support and the help, that you're actually getting comes from the Spatial Clarity Processing and the neuro clarity processing. And if you've been familiar with MoreSound Intelligence in the past, these should be pretty familiar terms.

The Spatial Clarity Processing consists of the virtual outer ear, which is a pinna model, and then the intent-based spatial balancer, which uses inverse directionality as needed to support the user in complex situations. The Neural Clarity Processing is our unique noise reduction system and that is based on our onboard Deep Neural Network, which has been trained to create contrast between speech and noise using a filtering strategy. Now the really cool thing is that the Neural Clarity Processing, which is made possible by the Deep Neural Network, is now going into its second generation. So this is a really big step forward. We've learned a lot through that first iteration of the Deep Neural Network and in the second iteration we were able to take those learnings to create a greater diversity of sound samples and to analyze those sound samples in a more precise way.

And that allows us to use even more noise suppression than we have been able to historically. Because we're more precise about removing that noise, we can be more

assertive in how much noise we can remove. So I'll go into that a little bit more in just a moment. But before that, I do want to take a step back into that Deep Neural Network and the idea of deep learning. There is a lot of artificial intelligence and deep learning around us in our world today, but how exactly it gets applied into the hearing aid processing can be a little bit of a mystery. And even if you've heard it before, sometimes it's easy to forget, so... So if this is new to you, I hope that this is worth your time.

And if it's a review, I hope it's worth your time as well. But basically the idea of deep learning is that a machine is learning patterns and those patterns can then be used to predict future behaviors, is really what it comes down to. And it's really modeled off of how the human brain works, that you have patterns that are picked up and those patterns basically reside in the neural architecture in the brain and when stimulus comes in, it gets processed by the brain and we don't know exactly what that processing looks like, but you get an outcome, and the person learns from that outcome, there's a consequence to an outcome, and that consequence then gets fed back into the brain, which basically updates what it knows, and the that it has based on that feedback that it receives.

And so there's this learning process that occurs. So with deep learning there is an output, and that output and that learning process has to be very carefully controlled. So we teach the system through this process of deep learning. And then when the hearing aid is available to the user, it has already learned those patterns. So it's not actively learning while it's on your patient, it has learned these patterns ahead of time so that now when it encounters new stimuli it can process those on its own. So that's what we mean by deep learning. So just an example of how this deep learning process would occur, you take two sound samples, one is target sounds, so this is clear speech that has been recorded, and then you take noisy sounds in isolation.

And then you can take those and put them together to create a speech-in-noise sample. So then you're really working with two samples here, you want the target sound, which is the clean speech, and you want the stimulus, which is the speech-in-noise. And you take that speech-in-noise and you put it into the system that is the Deep Neural Network. And that Deep Neural Network weights certain variables of that speech-in-noise signal and it emphasizes certain sounds compared to others and it provides an output. And that output is then evaluated to say, "Did we filter that speech-in-noise in a way that makes it sound like the target sound?" So it compares the output with the target sound to see how closely it has come to accurately filtering out the speech-in-noise to resemble the target clean speech.

It can actually mathematically compare these two and then it can use that information to update the relationships to create a new, more appropriate pattern. So let me just give you an example here. This is a photo of a dog and if I was gonna do some deep learning here to draw a picture of a dog, it would at first be random because those connections haven't been trained, haven't been reinforced with an outcome. So the first time it's done it would be pretty random. And so we would compare this random outcome to the image of the dog in the photo and say, "Okay, well that wasn't quite it, what can we do to get it closer?" And so the second time, based on what it had learned from the first, it comes a little bit closer.

So compared to that first drawing, the second one isn't really much like that dog, but it's more like it than the scribbles, right? Then the next time it updates even more and it gets a little bit closer to what that dog looks like and it uses that information to update. And then it comes with another image that's even a little bit closer still. Still not quite there, but we get more and more feedback until you end up with an approximation that's pretty darn close to where you started. And that's really the concept behind deep learning, is that you are progressively getting closer to what your target is. However, in the case of speech-in-noise it's really a little bit more like instead of proactively drawing

something, it's about taking away the right information to uncover what's beneath that is that clean speech.

So if you take this spectrogram, it's pretty noisy to begin with, what you're asking the system to do is filter out sounds so that you end up with this cleaner signal. And this would be like the the clean target speech. And so that's gonna be your reference, whereas the stimulus that you're inputting would be more like speech-in-noise. And you're training the system so that it gradually is able to do the spectrogram on the right versus the one on the left. As I said, this is really in line with what our brain is doing. And so that's partly why it's called deep learning is that it's really very much modeled off of the human brain architecture and the cognitive processes that are similar to learning.

So just give you an example, let's say that this father is with his little baby and maybe that baby sees my little Yorkie dog who's very, very tiny and fluffy. Maybe he sees that and his little brain sees this small animal and thinks, "Oh, that's a cat." And so he points at it and he says, "Cat." And then his father says to him, "No, that's a dog." And so now he's gotten feedback and he's able then to take that feedback and his brain will actually update its connections and strengthen some connections and prune some others. And so that the next time that he encounters a different type of dog his brain is actually able to use this new information and these new connections to accurately say that it's a dog.

And then once again he'll get feedback from his dad, "Yes, that's a dog." Right? And so that feedback will then further strengthen those connections for future use. So that's really the part of the idea here is that it's just really modeled on the human brain and what we're actually doing from day to day. Now specific to training the steel network over the first one, we actually have a greater variety of sound samples that are involved in training the Deep Neural Network 2.0. So we learned which sound samples were

most effective through this first iteration. And this is a completely newly trained Deep Neural Network and the diversity of sounds has been increased as far as the input.

And so that's that little part that's labeled A there is all these different types of inputs. Now the part B, which is the Deep Neural Network itself, that part B is actually the part that has remained the same. So the architecture and structure of the Deep Neural Network remains as it was. Part C, where the output is analyzed, has actually been updated. What was done in the past was, as I said, there was this mathematical calculation of how much did the output resemble the desired target speech? And it was done in by dividing the signal up into 24 channels and analyzing the correlation within each of those channels. What we've been able to do with the DNN 2.

0 is to process that output in 256 channels rather than the original 24. And that means that we have a much more precise understanding of how well the output matched the target and we can then give much better feedback through the system, which is that label D. So there's this backward propagation of information of feedback to the system to strengthen the connections most appropriately. So with this greater precision, we're then able to do more aggressive noise reduction because we're much more confident in the results. And so you can see that we have Oticon Real and this sentence that was processed, it was, "The street was very busy." This is the English version of the sentence, it was actually done in German and I can't say it so we're gonna go with the street was very busy.

But you see on the left how that processing was done with Oticon Real and where you've got that contrast, which was pretty good contrast to begin with, where you can see you have your speech cues standing out, but there is some noise still in there. On the right side you can see with Oticon Intent that you have a greater contrast between those dark spots of silence and those speech indicators of these orange and red lines there. So we've, for MoreSound Intelligence, we've improved both the dosage via the

4D Sensor technology, how much noise reduction and help-in-noise should you have? And then we've also improved the medication itself via an improvement in the Deep Neural Network and the new iteration of the DNN 2.

0. So we've got a lot of good outcomes that have come from this and I'm gonna be able to show you those in a moment. But before I get into that, I do wanna kind of go back to this notion of why it's important that we personalize this care so much. And you could see with that 40 sensor technology, how we're really living into this idea of precision medicine. That you have most medical treatments are designed for the average person, a one size fits all approach, which can be successful for some patients but not every single patient. And so with precision medicine, our goal is to target the right treatments to the right patient at the right time.

We can add a lot of nuance to the dosage of support that we're giving to people by understanding their intention. That is a big step toward becoming more precise in what we're doing for patients with the hearing aid technology. But we've been able to also add in another new strategy to the Oticon Intent, which is the use of audible contrast thresholds. Now I wanna make it very clear that the Audible Contrast Threshold, or ACT, is an independent test through our diagnostics colleagues. They've created this test and it can work with any whatever you fitting out there. And on the other hand, you absolutely do not need to do Audible Contrast Threshold tests, or ACT tests, in order to fit Oticon Intent.

It's really just a new tool that's available that can help further personalize your fitting so that you're giving, again, the right amount of medication and to the right patient at the right time. So what ACT does is it's measuring the quality of hearing. So the audiogram of course measures the quantity of hearing, your threshold in decibels, ACT is a way to measure the quality of hearing by being a hearing-in-noise test. And so we consider this hearing-in-noise test opportunity to be an opportunity to support the brain better

by giving the right amount of support with our instrument. So some highlights of the ACT test are that it provides an objective diagnostic evaluation of performance and noise, it's language independent, because it is not speech stimuli, it's a noise stimulus, but it's a spectro temporal noise stimulus, meaning that it has fluctuations that are similar to types of fluctuations that you would see in an ongoing speech signal.

And so the theory is that if you have this type of stimulus, what level of contrast do you need to be able to distinguish that from ongoing background noise and can we correlate that to actual speech performance and noise? And the answer is yes, that we have been able to do that. That was again done by our diagnostics partners and it's was able to be validated according to an aided hearing-in-noise test, the aided HINT test. And specifically an ecologically valid HINT test. It provides an above thresholds true-to-life assessment, prediction of behaviors for the aided HINTs test. But as I said, ecologically valid. And that was why we consider it to be a true-to-life.

It's quick to perform, it's approximately two to three minutes. It is performed using an audiometer, and I do have a list of which audiometers it is compatible with. And it is then personalizing the level of port as needed within the Oticon instruments when it is used with Oticon Intent or Oticon Real. So ACT in a nutshell, you have pink noise samples that are presented binaurally through headphones and interspersed with these noises you'll have these spectro-temporally modulated sounds, they sound like sirens, it's like a European type of a siren sound. If you're curious to take a listen to this there's other trainings on ACT that are out there and you can hear some samples of that also through Interacoustics websites, or the E3 websites, and so on.

When the patient hears that siren like sound in the midst of the ongoing pink noise, they press a response button. So you can imagine that this is very similar experience to having them the softest sound that they can hear and pressing a button when you're doing pure-tone audiometry. This is a similar type of procedure. They're hearing a

noise, they hear the siren sound within the noise, they press the button when they hear that siren sound. And that siren sound is adjusted automatically by this audiometer so that you are presenting levels going in a staircase method, very much like our modified Hughson Westlake procedure, that when a person actually presses the button, 'cause they've accurately heard that siren sound, then you go down two levels and then you go up five when they do not, or, sorry, up one level, or down two levels and up one, very similar to the down 10 up five procedure for Hughson Westlake.

And with the results, which is what's the contrast level between the softest siren sound that they can hear and the background noise. That result is going to be anywhere between negative four and 16. Now one of the things about the test stimulus is that it is actually based on the audiogram. So it's the audible component of the Audible Contrast Threshold test refers to the fact that you actually do have to do the audiogram first and then every aspect of that spectro-temporally modulated noise signal is presented at a super threshold level. And this is important because this is not something that we concurrently do with calibrated speech-in-noise tests that are out there right now and therefore this goes a long way to helping to predict what your aided speech-in-noise test is gonna be because you are in fact, a bit like with a hearing aid you would be providing audibility across the frequency range.

So this test stimulus does the same thing. So that's the audible component of the Audible Contrast Thresholds test. As far as compatibility, ACT is currently available for the instruments that you see here, we've got a total of eight audiometers that it's available with. If you have any questions about that, I would reach out to your diagnostic partners, find out if your instrument's compatible. And, as I said a moment ago, you're doing this... Sorry, my slides got a little out of order there for a moment. But our test stimulus is, you can see here, meant to be audible. So if you were having a regular uncompensated stimulus that was broadband in nature, like a speech-in-noise test, that would resemble something more like the gray line that you see.

And if your hearing loss and your thresholds are that light blue line then some of the information is gonna be below threshold and therefore inaudible. When you are compensating again to adjust for the hearing loss, like you would with a hearing aid, then you'll actually get what you see in that darker blue line as a stimulus, and that is the X type stimulus that is actually compensated for with the hearing loss. As I said, the level you get is somewhere between negative four and 16. It's reported in Normalized Contrast Level. So dB nCL. And dB nCL, the concept is a lot like 0 dB HL. Right? That you have a reference to a group of young, normal hearing listeners.

And so in the same way the ACT value and the contrast level is normalized to a group of young, normal hearing listeners and that normal threshold is 0 dB nCL. And so the range around that 0 dB nCL, of negative four to plus four, is determined to be normal. Then you would have a mild contrast loss between four and seven dB nCL. You would have a moderate contrast loss at seven to 10, and then a severe contrast loss between 10 and 16. And so that gives us a sense of as your ACT value gets worse, that you're gonna have more trouble in noise when you look at the prediction of aided speech-in-noise. And in researching this, just to give you a sense, the ACT score is very impactful in predicting the aided testing.

When you look at pure-tone average alone, we all know that pure-tone average alone cannot predict aided speech-in-noise very well, but it is important. And so you can see that almost 40% of the variability comes from the pure-tone average alone. But can we get better? And the answer is yes, that when you add in the ACT score you can get an even greater prediction. Or if you do the ACT score alone, I should say. If you do Age alone, it's a very poor predictor. But when you put the three of them together, the ACT score, the pure-tone average, and the age, you can get a really good estimate of the person's aided speech-in-noise ability on an ecologically valid HINT test.

Now that is all about ACT, which again is a separate independent test. When we look at what we can do with ACT in Oticon Intent, be able to then help to use this with MoreSound Intelligence to ensure that you get the best starting place for programming MoreSound Intelligence to match, as best as possible, the patient's needs. So there are many iteration of how we can program MoreSound Intelligence, but conceptually, let's say there's a low, moderate or high level of help that you might be providing, as compared to providing no help, or off. If you have three different groups of people, good speech-in-noise ability, moderate speech-in-noise ability, poor predicted speech-in-noise ability based on their ACT levels, we can actually measure what is the level of help that's needed for each of these different groups to be able to provide optimal support?

Now we don't wanna provide too much help-in-noise, simply because it's a deviation from the natural sound processing, right? And any deviation from that can have a negative impact on sound quality, and we wanna make sure that we're giving the brain the most natural signal possible. So we wanna use just the right amount of support-in-noise, and that's really what ACT can help us to understand, provide an evidence-based support here. So the research that was done was looking at measured speech-in-noise for those three groups. As I said, you've got good speech-in-noise, moderate and poor. And this gray line here is going to show the level of performance for normal hearing listeners.

And so if you have people with normal hearing, this is what their performance would be in an ecologically valid HINT setting. And so our goal would be to allow these different groups of users to get a performance level that reaches into that normal hearing level, right? To give them performance to where they can at least function in the same situations as their normal hearing peers, and in doing that allows them to engage. So let's look at this group with good performance. What we find is that the different levels of help that MoreSound Intelligence could provide, we could actually get that group

into the level of normal performance using the low MoreSound Intelligence settings. When we look at the moderate group, where they had moderate difficulty, we were actually able to get into that group using a medium type of setting.

And then for those people who had the poorest performance, we were actually even at the very highest level just short of that normal hearing performance. So in that case, we needed, if they were really needing to get to that level of performance, as normal hearing listeners, we might need some additional help on top of that. For example, remote microphone technology, right? So it's really just helpful to us to understand what's the best starting place in the programming for different groups based upon their hearing-in-noise performance. You can again use different types of tests to predict that somebody will have more difficulty when they are aided based upon other hearing-in-noise tests, or, better, sorry, other speech tests, speech-in-noise tests.

But we don't necessarily have the evidence to support exactly how we should program the hearing aids as the starting place for that person. ACT, in addition to some of the advantages that we saw a little bit earlier with the language independence and so on, ACT gives us the ability to be able to use that information to program the hearing aids as a really good starting place. So that is one of the additional ways that we're trying to be more targeted and precise in how we deliver technologies, the right dosage at the right time for the right patient. 4D Sensor technology and ACT allow us to come closer to that than we've ever been before. So lastly, what I'd like to do is take you through some of the clinical evidence for Oticon Intent, and we've got several different studies that were done.

So I just wanna make sure that you have an understanding of how we assess these and what you might be able to expect for your patients. So there were several studies, as I said, first as a technical study using head and torso simulator to really understand what the acoustic output of Oticon Intent is. Then we did a clinical study to understand

what the brain's behavior looks like using Oticon Intent. We also did a speech comprehension task and a speech understanding task and got some subjective feedback as well. So let's first look at the technical study. And the purpose was really just to understand the level of support that's provided for Oticon Intent in a comparison to Oticon Real.

And so what we did was we had speech coming from the front, male speaker at 65 dB, maskers and steady state noise coming from the back at 100 and 260 degrees. That intensity varied. And the head and torso speaker, or head and torso simulator, excuse me, the HATS mannequin was in the middle. The instruments put on the HATS mannequin were either Oticon Intent or Oticon Real with a closed fitting and they were programmed to a standard N3 hearing loss using linear amplification with all of the features at default otherwise. Oops. And we investigated the ability of the hearing aids to create contrast, and then we did an analysis of the Speech Intelligibility Index with that output to understand the availability of speech cues.

As I said with the Speech Intelligibility Index, we can take that signal to noise ratio and really understand how much of the speech is audible to understand what the access to speech cues is for each of these products. We did this with Oticon Intent with the 4D Sensor technology on, with the 4D Sensor technology off, and for Oticon Real. And what we find is first, what's that output signal to noise ratio? Compared to Oticon Real, Oticon Intent with the 4D Sensor technology off demonstrated an improvement. And then what we did in the other situation with Oticon Intent with the 4D Sensor technology on, since it is on a mannequin we actually forced the settings, using a research computer, into the lowest level, where they would be navigating the room or the highest level, which we need for intimate conversation.

And you can see that you get a range of performance, but that in, and in both cases it's higher output SNR than we were able to see with Oticon Real. When you convert that

and apply the SII to it, what you get is an increase in speech cues because of the improved output SNR. On the left there you have the 4D Sensor technology off. So what you're seeing is a 12% increase in access to speech cues just based on the new platform alone, and the new DNN 2.0. So by itself that new DNN with the platform already results in an increase. Then if you were to assume you were in an intimate conversation and add in the influence of the 4D Sensor technology, directing the dosage toward that level, you'll end up with 35% more access to speech cues compared to Oticon Real.

So we're really increasing access to speech and that then is necessarily gonna have some good benefits to the brain. So let's look at those benefits as well. We actually can understand whether Oticon Intent is providing the right amount of support to the brain depending on these user intentions by using some neural tracking methods such as EEG. And this is really important because, you know, some of our tests really only get at comprehension and, as we saw earlier, it's not just about comprehension, it's about listening and communicating and engaging. And so these are higher level activities that in fact do happen in the brain and so we need to have more than just speech comprehension measures, or sorry, more than just word recognition measures and so on in order to assess whether this is effective at the level of brain hearing that we're looking for.

So by using the EEG methodology, we actually used this on 30 experienced hearing aid users, mean age of 70 years, mild to moderate hearing loss, and they use the default prescription settings in the hearing aids, they used the two settings of 4D Sensor technology, on versus off, and they were fit using the VAC+ fitting rationale. And in this study their hearing speech coming from the front, there's a bunch of noises coming from the back and to the sides, both noise and then interrupting noises, and the EEG is being recorded, and actually what happens in this situation, it's a test paradigm that we've been using for a little while, you can actually follow the response

of the envelope of the ongoing target speech, or the ongoing background noise, and you can correlate that EEG activity with that.

And, based on the strength of that correlation, you can understand how salient that signal was to the brain because of how much is locking onto that signal. So you can see what the trials kind of looked like right here where you've got ongoing speech coming from the front, and then from the back speakers you've got noise that's happening, it's actually at a 0 dB SNR, so very difficult type of situation. And then there are also environmental sounds that occurred that would be similar to what you might see in a household or social situation. So cutlery, children playing, tools, et cetera. And so there's sounds that they might want to hear and are gonna be impactful to their attention in the midst of ongoing background information as well.

And so what we would hope would happen in a situation like this is that when you are in a situation where you would need more access to those types of sounds, if there's gonna be an ambulance around, you wanna be able to hear it well, right? If you're walking, or moving, and are in a situation where you need that high degree of awareness, you would like to see in the brain that there's a high level of awareness. And so you would wanna see a higher correlation to those background sounds when the hearing aids are behaving as if you should be picking up more sound around you because you need that awareness. On the other hand, if you're in a situation where you are in an intimate conversation, you would wanna minimize the brain's reaction to those sounds, you would want the brain to be able to exclude those sounds more.

And so we would want the right level of processing to support that. But in spite of how much access the brain has to those background sounds, we would always want the brain to have consistent access to speech. So it's really how much awareness we have of the environmental sounds that we want to vary and we do not want to vary the access to speech. And this is indeed what we ended up finding with these neural

tracking experiments. So on the left side you can see the speech targets, which is the neural tracking to the speech stimulus, and there's no significant difference whether they were in the 4D Sensor scenario, where they were navigating the room, which is the darker blue on the bottom, or if they were in the intimate conversation setting, which is that lighter blue at the top.

No difference for speech, which is exactly what we would wanna see. Then when you look at the sounds in the environment, and their awareness of the salience to those sounds, you can see that it was minimal. It was at its minimal level when it was in the intimate conversation mode. Then when you went to the navigating the room mode, where you want more awareness of those sounds around you, you actually can see that the brain had a stronger response to the environmental sounds in that situation. So we were actually able to show, using this neural tracking, that the desired outcome was occurring because of the processing. In the second study we were actually looking at speech comprehension.

And in order to do this we actually had to create a novel speech paradigm because we wanted to be able to really evaluate the interaction of the 4D Sensor technology for with a patient. And so in this case, again, we had 30 experienced hearing aid users, mild to moderate hearing loss, and we were evaluating 4D Sensor technology on versus off. And what we did here was worked with the Technical University of Denmark in their audio visual immersion lab and created a virtual audio visual scene. And so the participant would actually be seated and they would hear... Well, there would be 15 avatars surrounding them horizontally. And four different stories were simultaneously being told from four random avatars within among those 15.

And they were told to listen to a particular story, had to go and figure out which of the avatars was telling that story and then sit and listen to that avatar tell the story and then answer a question about the story. So it was really a listening comprehension task. And

as you can tell, this is a really difficult situation because you've got four different speakers, there's actually background noise happening at the same time and they're trying to do a comprehension task. And so this is really kind of that monitoring and then focusing in and how do the 4D Sensors respond to that? What you can see is that the speech comprehension in this very difficult task was improved when the 4D Sensor technology was on compared to when it was off.

And we actually saw a 15% improvement in their ability to comprehend that speech that they had to first localize and then listen to. So comprehension being a really high level task, not quite at the level of communication, but extremely important for communication. And then we had this third clinical study, speech understanding and sound quality. And so there were two parts to it. One, the objective part and two, the subjective part. Here we had 25 experienced hearing aid users, mild to moderate, or mild to moderately severe hearing loss. And we used Oticon Intent in its default settings with the 4D Sensor technology. We also did an omnidirectional setting and then Oticon Real in its default settings. And we used a test called the OLSA test, which is actually a German multi speaker test.

And the person evaluates what the person respond... Or repeats speech back. And so it's really evaluating their speech recognition, but it's while they're engaged in a group conversation. So let me show you what that looks like. So the person is listening to three different speakers, and any one of those three speakers could talk at any time. And then there's also background noise. And so whenever a sentence started with the name Kirsten, the participant was instructed to repeat the last word of all the sentences from that talker. But then another talker would start a sentence with Kirsten and that meant that they had to switch their attention to a different speaker, listen to that speech and repeat the last words of those sentences.

And then somebody else would say, Kirsten, and they would then have to orient to that person. So very much like trying to show what it would be like in a complex situation with overlapping speech. So there is an overlap on the sentences for about one second. They were supposed to keep their heads still so that they could be able to monitor and the, yeah, so you can see kind of at the bottom here too, how those sentences would overlap with one another. So very much a monitoring task. And we compared, as I said, the default settings, Oticon Intent versus Real versus Omni. And you can see here that Oticon Real over Omnidirectional resulted in improvement in speech recognition.

And that Oticon Intent also increased over Oticon Real in terms of speech recognition. And that was for the speech coming from the front and then off to the side as well. And then in addition to the quantitative measurements, we actually did subjective measurements for sound quality, and the sound quality was really grouped into an overall rating of sound quality. They were asked to rate the nuances and details of the speech and they were asked to rate the comfort. And so they were listening to different types of settings. There was a forest scene, speech in quiet, speech in cafeteria noise, speech in traffic noise, speech at the train station, speech at the pub. And what you see here is a table of those ratings for each of the different sound scenes.

Overall, to summarize, those Oticon Intent outperformed Oticon Real by up to 10% better sound quality, up to 13% more nuances than sound scene, and up to 10% higher listening comfort. So with that, I want to wrap up this training and I wanna thank you for your time and attention. I hope this gives you some background into the strategies for Oticon Intent and I would welcome you to reach out to myself, or to your trainers, or managers if you have other questions. My email address is vira@oticon.com and I wanna just thank you for your time and attention.