

*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](mailto:customerservice@continued.com)*

## **Enhancing Human Auditory Function with Artificial Intelligence and Hearing Aids**

**Presenter: Dave Fabry, PhD**

Hi, I'm Dave Fabry, Starkey's chief hearing health officer, and the topic of this talk is enhancing human auditory function with artificial intelligence and hearing aids. I'm going to frame this discussion over the next hour around these three main topic areas. Our focus is really first and foremost to focus on what will provide the best hearing performance in terms of sound quality and speech intelligibility for hearing aid users. But we've also been focused on using artificial intelligence in terms of health and well being because we know that hearing care is healthcare. And I'll talk about recent advances that we've made in the past year in this topic area.

And then the third is an homage, really, to chat, GPT and OpenAI in the ways that hearing aid users may be able to use voice commands, to ask queries and automatically triage whether it's a question about how to operate their hearing aids, whether they want to control their hearing aids, whether they want to go out and search on the Internet. That whole idea of using natural language processing as a virtual assistant is a burgeoning area of opportunity for us in the hearing aid industry, and I'll talk about some progress that we've made in terms of collaborations and partnerships in that final stage. First, better hearing. And really the top three drivers, or top four drivers, rather, of hearing aid performance and sound have been, according to the marketrack series, relatively constant over the last decade or so. First and foremost, natural sound quality and speech intelligibility is the highest standard, if you will, also ensuring that minimizing background noise in noisy listening environments.

That's often where people first notice that they're having difficulty and then also preventing loudness discomfort in loud sounds. We never want hearing aids to be uncomfortable. We want to use as much of the residual auditory area in terms of the frequency domain and in the amplitude dimension, but never over amplifying. The fourth one is the one that's often overlooked, and that is the ability to tell where a sound is coming from. As we've seen, hearing aids move from being primarily monaural fit around the world to now binaural.

In the majority of cases, the patient really expects that to be able to close their eyes and tell not only when a sound is present, but where that sound is coming from. So as we're focused on this, I want to level set by talking about the fact that we've seen in the hearing aid industry tremendous disruption in terms of over the counter hearing aids and third party payers and disruptions from low cost providers. But a reminder for those of you who are working with patients dispensing hearing aids is that hearing aid satisfaction and benefit has never been higher. 85% of patients queried in this study that was published by Aaron Piku indicated that they were satisfied with the physical aspects of hearing aids, and 81% indicated that sound quality of those hearing aids was providing them with delight in terms of noise reduction, naturalness, clarity, feedback, et cetera. The opportunity for growth, however, is that understanding speech in challenging listening environments where noise is present or multiple talkers are present gets a little more challenging.

Seeing that as we went from one on one conversation, about 83% of people satisfied, down to large lecture halls where there's noise and reverberation, getting to 72 or even 70% in the classroom where there may be multiple noise sources interfering with the teacher. When we talk about modern hearing aids and the use of artificial intelligence, I think it's important to remind ourselves that most modern hearing aids have used acoustic environmental classification for about the last decade to monitor the listening environment where the patient is wearing his or her devices and adapt by incorporating features like noise management, directionality, wind, noise suppression, specific time constants, or gain adjustments for music. All of those situations really allow for the patient to wear their hearing aids, go throughout their daily business, and they just adapt according to the different environments, employing those features only as appropriate through machine learning acoustic environmental classification. But the issue is, how does that classifier know what the end user, what the listener wants to hear? I want to spend a couple minutes talking about this AI based sound classification as it exists here.

I've just illustrated seven different acoustic environments where a hearing aid user might expect or want to wear their devices. And I'll just give a framework, if you will, in terms of a machine learning acoustic environmental classification system showing how sound is input, and a feature extraction system that in this example, uses five temporal and spectral and timing features to be able to sort out, if you will, the different acoustic environments, like a deck of cards, into the different suits and the different cards, the different scenes, and then this post processing that then optimizes for performance in those scenes by applying directionality, noise management, etcetera. And the success of any environmental classification system is how accurately can it sort out when the devices are being fed different quiet, noisy, musical, windy scenes into the appropriate pile, the way a human would classify those environments, if we simplify this into a cartoon, if you will, that in this case only uses two features, as shown here, on an x and a y axis. Let's say one's a temporal feature, the other a spectral feature for those seven different scenes that I'm now fed in and see how well and how successfully even using only these two features, we separate those scenes into distinct entities. You see, there are four clusters here, three that are uniquely different from the others, even with just a two feature classification system, and then four that are lumped in the middle.

And it really is the question of what causes a machine learning classification system to make errors. And I would say it's really life. And shown in this screen, what we see is from New Orleans on Bourbon street, where there's a cacophony of noises, including that carriage going by, the guitar player, the trombone player, a bunch of revelers, all interfering with the ability to understand speech. Speech can be both the stimulus of interest if you're engaged in a conversation, and background noise with all the people walking on the street. Music similarly can be both a stimulus of interest that I want to stop and enjoy, or if we're engrossed in a conversation that I want to suppress.

So really, what I'm saying with this is that sound is personal in the sense that everyone has a unique auditory ecology, as Stuart Gatehouse first coined this term a couple decades ago, now before his

untimely death, that the unique aspects of that listener intent of what they want to hear at any given moment really personalizes that sound. And this really began our journey with Edge mode back in 2020 when we launched a feature that we talked about at the time, enabled a patient to put AI at their fingertips because they could double tap into an acoustic scan that would move beyond acoustic environmental classification to give, really another opportunity when a patient is having difficulty to say, whatever is in front of me right now is my stimulus of interest. And if I double tap, it does an acoustic environmental scan to relook at the different types of acoustic categorizations that may take place while applying offsets that initially when we introduced edge mode, would increase the audibility of the sound in front of the individual. Breaking this down, similar in the workflow as I showed for machine learning classification. Now what we have is the user requesting assistance by double tapping, or now pressing in an app that environmental detection, acoustic snapshot, if you will, to look at what's in front and what's around, and then to optimize for that sound.

The user then, in turn can listen immediately to those new parameters. They don't even need to be connected to an app. It truly is edge computing, hence the name that can be applied by the individual with only the hearing instruments and a double tap if they want or through the app, and then the user can accept or reject those parameters that have been included to try to optimize and improve clarity for the sound in front of them. What we found was the development of this led to spontaneous user acceptance in terms of not only noisy listening environments, but later in 2020, after we launched this product, many of the patients that I was working with said, hey, you told me to use it in noise, but I'm finding that when we were all wearing masks, it crisped the sound up to use the words of one of my patients. And so we really felt we were on our way with a feature that could combine machine learning classification with listener intent.

To say, what's in front of me right now when I apply this acoustic scan is what I'm interested in hearing. We continued, however, with that journey on a next gen edge mode that provided that optimization, and

we started to look at and compare how well would patients prefer and do with edge mode versus the acoustic environmental classification? And I'll show you some results of that initial study that we did. This was back with evolve as we first launched this with edge mode versus the normal or automatic memory. And what you see is that on is better or much better for the majority of those patients with also reduced listening effort as well.

And this was for receiver in the canal devices and custom devices. We wanted to continue to push the envelope in terms of how well Edge mode could do to personalize and optimize to specific listening situations, even in lieu of specialized programs that could be used by the end user in different listening environments. Our goal was to assess whether the automatic or now personal program, the automatic environmental classifier plus situational use of edge mode would be preferable over the AEC program as well as a personalized program. In this case, I'll show you in a moment for an automobile or car memory. So, first of all, looking at an edge mode with this evolved feature, we saw that sound, speech, clarity, listening, comfort, and overall preference favored edge mode versus the AEC, or we called it normal program at the time.

And then when we looked in and saw that there was a preference for clarity, comfort, and overall preference versus the normal or personal program, the AEC automated program. Looking at the chart on the right, what we saw even was more interesting to us was with the professional delivering a car memory. It still showed that edge mode had the same or better in terms of the clarity, the listening comfort and overall preference versus an optimized dedicated hearing aid memory. So, continuing down this road to personalization optimization, combining AEC plus individual preference brought us up to where we are now. And what I'm going to report on is the results of a collaboration, a multi year collaboration that we've had between Starkey and Stanford University, where we looked at as well speech recognition in noise measures between the Kwixsin and a number of different speech and noise measures when AEC was on versus the edge mode program was on.

And I'm going to show you just a brief snippet of this. Stay tuned. We just recently presented an expanded version of this collaboration at the American Academy of Audiology meeting, and we hopefully will be submitting that for publication soon. But what we see here is edge mode on and edge mode off for CNC words in noise. And what you see is that for multi talker Babel at plus five decibels, signal to noise ratio is generally higher.

You look at the right hand side of the curve in terms of CNC words with edge off on the x axis and edge on on the y axis, and you'll see that the majority of the subjects in this study fall above the diagonal, indicating that they were doing better when edge mode was on versus when edge mode was off. These findings really suggest that for CNC words, and also for quicksin, the s and R loss is suggesting that in this case now with s and R loss, the SNR loss is lower by about a decibels when edge mode is active. And you think that's in addition to the AEC classification, automatically engaging directionality in a very challenging diffuse noise environment. So a DB can translate into connected discourse into ten to 12% more, a noticeable and statistically significant difference, improving that SNR loss and also providing improved CNC word recognition. What I want to show on this next slide is now the data tells one thing and one part of the story, but in addition, I want to show you one of the best parts of my job is that in addition to working on the features that are coming out in our hearing aids, I also get to work with patients.

And I want to show you two patient scenarios here. The first one from a dairy farmer from Minnesota. They've given us permission to share their stories. I had fitted him with the product, not really realizing that he also was frequently called into meetings on the dairy council, in noisy restaurants and noisy lecture halls to communicate with other representatives from the dairy council. And I had shown him only in cursory fashion how to use edge mode.

But I had also given him some manual programs. He came back and I was doing a virtual follow up appointment with him via telehealth and listened to what he says about how he had gone and explored with edge mode and found a that it really helped him in ways that I didn't anticipate what happened yesterday. We were seated for lunch in a kind of a noisy area, and then a band started and I'm like, holy smokes, which mode should I be in? And I just went, oh, well, let's do this edge mode. Because, and it was absolutely weird for those of us who struggle with hearing to all of a sudden have everything just go, hey, I can hear.

It's just crazy. And to be able to focus on a conversation, I'm no longer having to be that one, sitting there nodding as everybody's talking around, trying to be a nice smiley face. So I've said for a very long time that you don't know your patient. For those of you who are working clinically, until you know your patient, I made the cardinal sin of assuming that he wouldn't want or have the inconvenience of having to apply edge mode, or multiple programs for that matter. But he instantly went into a situation.

He remembered that I had explained how edge mode was to be applied and found the benefits because he was willing to put his fingertips to work in terms of that Aihdenhe and listener intent beyond the acoustic environmental classification. I contrast this with another patient of mine who is a Hall of Fame NFL player and also while he was finishing his NFL career, went to law school, became ultimately a Supreme Court justice for the state of Minnesota. Certainly I was thinking he might want to engage and frequently be in challenging listening environments, but his experience is completely the opposite of what you saw in the last patient. Listen to what he wants out of his use of hearing aids in everyday situations.

You know, that's been one of the blessings is stick them in and basically forget about them. That's awesome. I have to remind myself when I go to take a bath or a shower to take them out. So first of all, I've said that's the best compliment that we can ever receive, is if a patient forgets and jumps in the

shower with their hearing aids. It's that natural.

But notice that his expectation out of hearing aids, of course, is that he wants to hear with great sound quality and intelligibility. Never have it be uncomfortable to provide benefits in noise environments. And then he doesn't. He wants a seamless experience other than that. So he, justice Page remains my big challenge because I would go back to him and say, well, try edge mode.

And he's like I'm hearing great with just the personal program and so, but I still want him to get those remaining, let's say 20% of situations where the, the acoustic environmental classification won't optimize or personalize. I always want to be focused on providing the best sound quality and speech intelligibility in every environment. And when we combine that intent with the acoustic environmental classification, we can get that additional performance benefit that you saw in the lab. But he wants a seamless experience. So when we went and continued developing these features, we came up with edge mode plus to enable him to, or any end user to not only choose that best sound so when they press the button or tap their ear, that they can have best sound, which provides offsets to improve clarity.

And then we added two additional more granular ways for end users to engage with the artificial intelligence aspect of this, by either enhancing further being even more aggressive with those offsets to provide improved clarity or to reduce noise. These are available on our top two tiers of the devices, the 24 and 20 series, but they can identify their listening intent, whether they want to enhance speech or reduce noise, as well as for those who just wanted to have the best sound, the combination of the two with a little bit more clarity in the original implementation of this. So we went and delivered this feature and found that when activated by the user, the hearing aid now classifies the listening environment and applies specialized changes to the device beyond that available in the default hearing aid settings and based on the listener's interest or intent. These classification and adaptation schemes were derived via

both machine learning and using a deep neural network to train these devices with our accelerator engine, in which these models are trained on a large number of real life sound samples and refined with input from both clinicians and listeners alike. And I want to report briefly on the results of three studies that looked at what we call now edge mode plus.

With that additional granularity for those who have the capacity to select improved clarity or improve comfort in challenging listening environments, we'll look at study one, which was that edge mode plus on speech understanding and noise, and you see a pretty garden variety sort of n two or n three kind of hearing loss, a presbychusis loss, noise induced loss, with twelve experienced hearing aid users fitted with Genesis AI 24 ric devices with edge mode, and with the estat 2.0 fitting formula, the speech understanding noise conditions you see here, with speech coming from the front, noise from a diffuse array around them. And then we've assigned, for purposes of this study, Edge mode plus enhance speech in this case with restaurant noise at 63 decibels, and we're presenting at an individualized level for speech to achieve 70% speech understanding with the default hearing aid processing, that personal program, if you will. And what we found in this case was speech recognition in noise resulted in significantly better speech recognition in noise than in the default hearing aid settings, even in that very complex diffuse environment. So beyond what you can already achieve with directionality and noise management, when we went into maximize speech mode in this case provided an additional roughly eleven and a half percent improvement in recognition that was significant at the 0.027 level. So that's consistent with what we saw in the Stanford study that was done independently.

We did this study internally and saw a similar degree of benefit in this case in this very challenging diffuse noise situation with speech coming from the front. The second study I want to report briefly on is on listening effort, if you will. In this case, again, 20 experienced users, five females, 15 males, range with a mean of 71 years, and then again using estat fitting formula with AI 24 RIC devices. In this case, we were also enhancing speech and looking at the effort to listen in those cases using the adaptive

categorical listening effort scale and with using the english matrix test modulated noise, and then had individuals rate the listening effort across an array of signal to noise values with a fixed noise level. And in this case, what we saw at a high level was that edge mode indeed improves listening effort, reduces listening effort compared with the automated approach.

And as an on demand feature, edge mode plus can provide significant benefits not only in terms of speech recognition and noise, but also in reduced listening effort, if you will, for those that need additional assistance in noisy or challenging listening environments. Moving on to the third study, I want to talk about the fact that there are indeed individual differences for edge mode plus in terms of benefit, and what we found was in the most part more significant losses in older individuals. This is painting with a fairly broad brush, but the important thing is that there are individual differences, and future work is looking at how it is that we can characterize who might be a better candidate for use of this on demand feature beyond just having them try it and by being able to predict on the basis of whether they have more significant loss, more s and r deficit potentially, and other factors that may be used by clinicians to be able to predict which individuals are most likely to benefit the most from this edge mode plus feature. I continue to be reminded then of my justice page patient who was resistant to the initial edge mode. I presented Edge mode plus to him and he said, no, I'm still hearing great with these devices.

They're just seamless. I don't have to think about it. But then we wanted to continue on this path for Justice Page and others who want that seamless implementation of AI in combination with machine learning, classification and listening intent. And so our automatic edge mode plus is the latest attribute that we've just recently introduced. In the last couple months, in this case, we've had many clinicians and patients say, with the original edge mode or even edge mode plus.

If it's so great, why can't I just have it on all the time? We had been somewhat hesitant because we thought those offsets for clarity or for reducing noise might be too aggressive in real world situations to

be applied continuously. But when we took it to the lab and began testing on some individuals, we found that there was a large group of individuals who wanted to just be able to use edge mode in an automatic fashion the way that they would the personal or AEC classifier. So for those who are still comfortable using situational edge mode, where they tap in and tap out as they go from one challenging environment to the next, they can still do that. But now for those who wanted it applied all the time, where it continues to adapt either more aggressively, enhancing speech or reducing noise throughout the day, now they have that feature on the latest implementation of edge mode automatic, and the preliminary results indicate over 80% of patients report benefits from edge mode plus automatic in everyday listening environments because of the convenience.

Went back to justice page one more time. He tried it, and this was his initial report indicating his satisfaction. And also the fact that now the new reality is that it's so seamless, he kind of forgets it's on until he remembers how much better he's doing. Okay, justice page. So you've had Genesis devices now for, gosh, I think, since last year.

Yeah. And, but then, and we talked about the fact that you gave me the greatest compliment that we could ever get, is that you just like to put them in your ears and go throughout your day and not have to engage with them. To the point that when we talked about the edge mode feature as we introduced it in the product, you had to engage with it regularly. And probably you said you were hearing well enough that you really didn't feel the need. But now we've updated to the edge mode automatic, and I think you've had a chance to play with it a little bit.

Well, now I've got the other side of that issue being it's automatic. So I don't think about it. Right.

You know, other than the initial use of it now I don't think about it. So I can't tell whether it's working or not. Yeah. Your initial feedback, you said you thought, I think when we discussed it, right after you updated the firmware and you did it from the app and you said things sounded more crisp and clear.

Clear.

And now. And now you kind of take that for granted. You're hearing well, and now we fallen back in. As long as you remember to engage the edge mode automatic when you put them in, you're finding that now you like to just leave it there and go about your day. That's where I keep it.

That's awesome. Okay. And you do find that you are hearing well, notwithstanding any challenging listening environments. Yeah. Okay.

That's good. So what you could hear justice page talk about in that sample was really now the new normal is that it's always on. But when he first got it, he said, oh yeah, I didn't realize that I was using automatic mode. But I do find that I'm hearing sounds clearer than I had before. And now the challenge is it's become so seamless for him that he forgets that it's on.

And one nice thing is to every now and again, be able to tap out of that and go back into the automated environment to remember that it's those 20% of the time when you have really challenging listening environments where there's a bunch of talkers around or noises or music as a noise to see indeed, if they're performing well in this seamless fashion. One pro tip, I would say, on the basis of justice page and others that I've tried this with, is that it may well be that if patients want to just go into edge mode plus automatic full time, have them begin with best sound, because that's applying offsets that increase the aggressiveness of noise management and apply the improved offsets for audibility for the sound in front of them relative to the AEC program all the time. But I have had a couple cases where if they wanted to be on clarity, enhance clarity full time, that that may be a little too aggressive for everyday environments. That's why we built the personal AEC program. And for those who are going to use it all the time instead of the personal program, I'd recommend that they start out with edge mode automatic on best sound and take it from there.

Now I want to return back to the top four drivers briefly where we talk about those natural sound quality, quality speech intelligibility comfort for loud sounds, minimizing background noise. And I want to talk about the fourth one, the ability to tell direction, that one that's often overlooked. As I mentioned earlier, in many cases, when we're talking about sound and the hearing and balance system, we're focused on the sound going into the cochlea. But we got to remember that the vestibular system is attached to the hearing imbalance organization, and it provides important information about movement and the patient's orientation in space. We've certainly, through our fall detection feature, and I'll come to that more in just a couple moments, found that having embedded sensors near and in the ear, near the balance organ is important.

But also we found, and we've included motion based optimization that can use those same sensors that can track footsteps, exercise and standing can also improve sound quality and improve spatial awareness in these latest products. I'll give you three examples. The first one is improved wind noise reduction when the sensors detect that the patient is moving and outside that less wind will be evidence, as seen in the chart, when we switch into omnidirectional mode, when they're moving to have the wind going across one diaphragm instead of two microphone diaphragms. And you see that wind noise in addition to the wind noise enhancement, starts at a better baseline with using omnidirectional versus directional microphone inputs. In addition, improved spatial awareness shows a statistically significant improvement in the signal to noise ratio needed to achieve soft but audible targets for target signals that you see listed here, when a person is using that motion based adaptation for those different sounds relative to when that motion based optimization is not active.

And then the final one is people are active and they want to wear their hearing aids when they're exercising, when they're playing tennis, as in this example. And we found that when activating that motion sense feedback management, we get 10% greater gain margins to improve feedback cancellation during physical activity. So the motion sensors we've done for health and wellness in the

past through encouraging better cardiovascular health by exercising and getting their steps in and musculoskeletal strength to move, but now also applying it in the acoustic arena as well, to be able to improve feedback management, to reduce wind noise, and to help not only identify that sounds are occurring in the environment, but potentially to help locate in space where those sounds are, which is that fourth driver of expectation for hearing aid benefit. Now, I want to transition a little bit from the speech and noise and better hearing, if you will, results that we just shared into the health and wellbeing. And I alluded earlier to the partnership that we've had with Stanford over the past several years, and I want to talk about an important project that transcends better hearing, but thinking about hearing and balance.

As many of you know, we were the industry's first and still only manufacturer to introduce a fall detection feature back in 2019. And we've heard from the market and from patients and providers and families alike that it's provided peace of mind for family members to know that their loved ones can live independently in their home and not worry that they're going to become the old woman from the commercial in the eighties that fell and couldn't get up. But the reality is as well, a fall detection feature that can alert up to three trusted contacts with a text message, even show on a map where they were, when they were wearing their devices and suffered a fall. That's a great feature. But if the patient broke their hip when they fell, it begins a downward spiral of health that in many cases, falls don't cause death, but as a result of consequences of the fall.

Like my mother, a couple years later, she passed away, not due to the consequences directly from the fall, but that downward spiral in health. So we began thinking about ways that we could potentially, in the long term, prevent falls before they occurred. And that's really the genesis, pun intended, of this project and partnership and collaboration between Starkey and Stanford, where we wanted to look at whether hearing aids can assess fall risk. We also reported on this at the American Academy of Audiology meeting in April in Atlanta. This year, falls are the leading cause of injury related deaths in

older adults, resulting in over 31 billion in healthcare costs per year in the US alone.

And then looking can AI driven changes in signal processing improve the signal to noise ratio? That's the study that we already talked about in the past one. But I'm going to focus on this issue of fall risk. Our goal is to determine whether hearing aids can provide similar ratings to that of trained observers on a fall risk assessment, and I'll expand on that more in a moment. First of all, I talked last year at this conference about the fact that we have our hearing aids embedded with additional sensors that can detect velocity and rotational changes when the wearer is wearing binaural devices, then connected to their smartphone, that fall detection feature is made possible.

What we wanted to do on this partnership was expand on that to begin to look at how we could do risk assessment for falls using the CDC criteria. If you use the QR code here, you can get additional information about the CDC and their look at identifying 95% of people who are at risk of falls with three questions do they feel unsteady while walking? Do they worry about falling, or have they fallen in the past year? If a person answers affirmatively to at least one of those three questions, it identifies 95% of the people who are at risk for falling elevated versus the general population. So we began with the goal of using these simple three questions for those working clinically who said, I don't want to work with balanced patients in depth.

But those three issues could identify a person on the basis of those three questions whether they're at risk for falls. The method of this study was to look at three different segments that assess balance, strength, and functional gait using the scoring, if you will, through a hearing aid versus a human score. Three observers for a large cohort of participants, 250 patients aged 55 to 100 years of age, with a mean age of 78 years, roughly two thirds female, one third male and then using the CDC stopping elderly accidents, deaths, and injury protocol, which I'll expand on more in a moment, this algorithm accurately moves beyond those three screening questions to identify whether a person has balance

gait or strength risks related to falls. And these limitations of the steady algorithm, pretty good sensitivity and specificity in discriminating low and high fall risk. It's not specific in predicting mortality, risk, or other risk factors.

Want to be very clear on that. Some of the advantages, however, it's very simple. Coordinating between clinical and community based practices. Our real goal, to give the patient in the long term, the capability to assess these measures in the convenience of their own home with the accuracy that can be achieved in a clinical implementation of this. That makes it then more suitable for widespread implementation.

To summarize the four stage balance test procedure, we look at monitoring whether a patient's standing with their feet side by side for 10 seconds and having an observer notice whether they swayed in a large fashion. The IMUs, the inertial measurement units, can pick up very slight sways, and we wanted to look at seeing whether there was good agreement between a human observer and the algorithm scoring whether a person was swaying when they stood with their feet side by side, whether they were in the instep, as you see in the second set of footprints, whether they're tandem also holding for 10 seconds, or I balancing on one leg or the other again for 10 seconds. All of those really are specifically focused on balance. The second area that we talked about is in terms of that strength, if you will, and the CDC on the steady protocol has a 32nd chair test, which, like the observer, like the cartoon shows here, no arms, arms on the side of the chair. To help themselves up, they have to be able to get up and sit back down as many times in a 32nd interval as possible.

And that really gets that a strength factor. And then the last one, the timed up and go gets at a gait component of balance, which is they start in a seated position, get up, walk 10ft, turn around and come back down. In each of these cases, we have a human observer doing scoring on those measures, and then we have automated scoring on this. And in the case of the timed up and go, it's a number of

seconds from the time they're seated to go back and seat again after going 10ft out and turning around in the stand sit test, it's just the number of times that they get up in that time interval. And then the balance test is how steady they are with their feet moving from those four different seated positions to get at the balance issue.

And then for these to get a pass fail for the balance test, the stand count is how many times they did that. And then the timed up and go is a time. Well, let's get to the results briefly to look at how accurately and how in close agreement, there is rather between the human score and the automated scoring. Now, one thing that I'll preface this with is we make the assumption that the human observer is the gold standard, but we also know that humans make mistakes as well. But in this case, what we saw was excellent agreement with 93% scoring agreement for the side by side foot position, for the slight instep back on one side.

And you see doing here with a little bit of wiggle. On this screen we have 91% scoring agreement and on the next, 186 percent for the tandem. This gets a little bit harder with the 2ft like that to get agreement between the human and the automated scoring. And in the overall, the 1ft position where they're up on 1ft, there's 90% scoring agreement. So excellent, really.

Agreement between the human and the automated scoring for all four of those, the worst one being the tandem, showing a discrepancy between the human and the automated scoring. If we look at the second one of the 32nd chair stand test, what we see is the diagonal. As the data fall along the diagonal here for this 32nd chair stand test, that the more closely they fall that way, the better agreement is. And we found excellent agreement with an  $r$  squared of 0.93, which is statistically significant at the 0.001 level. In this case for that 32nd chair stand test on the next one.

The one thing that we did see in those rare occasions where there was a discrepancy between the human and the automated scoring, the observers report nearly one more stand than the hearing aid

automated scoring in this case. This was, as I say, a statistically significant agreement between this.

But talking about whether the human or the machine is the gold standard, what we found was that if the machine and the test in empirical sense only gives you a full count if you sit and go to stand and then sit back down again, humans were more likely to give a partial sit credit for full count on a stand sit in that case. And what we found was when we lengthened the interval slightly to be able to include that final seated position, we found better agreement, but some debate over whether the machine was actually more accurate than the human in this case, and we were just giving credit for finishing the sitting part of the stand sit procedure in the last one. We also saw a strong relationship between the hearing aid and observer on the timed up and go in this case, and the hearing aids couldn't calculate the time for some of the position.

We had 220 individuals that the machine and the human could both agree on the get up and come around. But for 30 of the participants, the hearing aids couldn't calculate the time. And this was thought to be an error between the individual not being able to operate the app. And we know in many cases that some patients can't use a user app for user control on their hearing aids. And we found similarly some patients, a small percentage, maybe 10%, can't actually do this task.

And so in that, I think it's just a simple referral back in for those on that subscale, in to see a professional for this. For 220 of the 250 patients, there was great agreement on the timed up and go at 0.94, with a probability of less than 0.5. In terms of their difference between these two, the agreement between the human observer and the machine was quite high. If we look at a high level summary on this in terms of the remote assessment limitations, for those who are interested in the longer and additional information that we published or we're going to be publishing and we just presented at the American Academy of Audiology meeting, you can see that the next stage of this study, we looked at what happened when the clinician was not in the room with the patient, but rather remote on a screen, counting the patient while they were doing the testing in their home. And then the third stage would be

in an unsupervised fashion when the patient was in their own home doing the four stage balance test, the timed up and go and the Stan sit test, and do that from a progressive nature.

We really wanted to err on the side of very slowly moving from supervised with a clinician in the room to a clinician remote via telehealth to the patient in an unsupervised fashion. And we really see promising results from this initial stage. We've completed the second stage and are really on the way to the third stage of providing a way so that patients could do the steady protocol assessment in the convenience of their own home. With then the long term goal be the possibility that they could use this steady assessment in combination if they identified a weakness in Gaitley, a weakness in strength, or a weakness in balance with training exercises that ultimately could be used to improve their steady scores, that they could monitor in between clinic appointments and working with the audiologist, with the physical therapist, with the physician to improve their balance and potentially, and their balance, strength and gait to potentially lower the risk or potentially lower the risk of a fall before it occurs. So with the remaining few moments that I have, I want to talk about the virtual assistant aspect of this.

We've certainly seen the, as I mentioned, the chat GPT phenomenon has grown to the point that natural language processing for being able to ask for questions about if you've gone out and used chat GPT and ask something about audiology or about hearing aids, you can get answers that collect from a variety of sources from the net. One of the things that we've looked at with the virtual assistant for some time now has been we want, in the case of individuals who have limited dexterity or neuropathy, to be able to control their hearing aids using their voice or a double tap in combination with their voice, doesn't require fine motor control, but a tap here or a button press in the app and they can actually say, raise the volume control on my hearing aids, and it will actually change and control their hearing aids. Change to an outdoor program, and it'll change programs using voice commands. As a virtual assistant, the cool thing from my perspective is it will also the virtual assistant will also intelligently triage to say, how do I change the volume control? And it can populate in the user app on their

smartphone to show them and guide them as to how to change the volume control to open that memory, how to change the volume in both ears or one ear and walk them through and even give them examples in the app.

It can also then intelligently know when to go out and search on the web to look for a question as to what the weather is going to be today and know where they are and tell them what the weather forecast will be. And that virtual assistant. It will triage between being an electronic manual to going out and searching on the web, or in the third case, to be able to actually control their devices using only their voice without having to manually manipulate on onboard controls or app controls to change their hearing aids. The area of opportunity also came about in the past year when I was fortunate to work with the CEO of Intel. He has a hearing loss and I fitted him with hearing aids.

He quickly saw the potential for those with hearing loss to be able to interface. And what he talks about is the AIPC generation, and I'll let him, Pat Gelsinger speak for himself. But what he did was quickly developed a demonstration, a proof of concept demonstration, of this collaboration between Starkey and Intel on their innovation day last year. And I'll just play this brief video and let him describe it in his own words. We've talked about the superpowers.

This idea of these five technology superpowers just invading everything, compute, everything becomes a computer connectivity, everyone and everything is connected. Infrastructure, the unlimited scale of the cloud combined with the unlimited reach of the intelligent edge, simultaneously addressing latency and higher bandwidth. And AI, this intelligence everywhere, being able to take this data and compute and algorithmic breakthroughs, software, writing software at scale. But the last of these sensing and breakthroughs in low cost, high resolution sensors, I believe are just opening up other pathways to bring technology into our everyday lives, bringing more data and capabilities. And we're seeing advances in automation, processing, robotics, giving machines human like capabilities where even our

disabilities become digitally enhanced strengths or superpowers.

One of my favorite sounds is hearing my granddaughter call me Papa Papa Pat. And if it were not for my hearing aids, and I have a family that almost every one of them has lost their hearing, I might not be able to hear that in the future. And I believe in this area of the superpowers of sensing, we have so much more to do. And with that, I'm super excited to find a soul mate on this journey. You saw it with rewind, leveraging core ultra and Openvino, and transforming lives and improving accessibility.

But until recently, PCs couldn't connect to hearing aids like mine because the traditional Bluetooth simply used too much power. Well, recently, Bluetooth low energy audio that we've worked on with Microsoft is first coming and available since earlier this year. And part of the core ultra platform. And we've been collaborating with Starkey Labs. And these hearing aids are right, like the ones that I'm wearing here, to create a POC for how AI can improve the hearing aid experience.

So we're going to join a call here with Arno on my Samsung Galaxy book. So, Arno, how are you? Hello, Pat. I trust your keynote is going well. Yeah, I trust you're keeping track of all the great things I have to say, Arno.

Yeah, I'm watching from backstage. So now your PC is contextually aware and will automatically switch your hearing aids between ambient aware and focus mode. By adjusting the sound amplification, you can sustain your flow without missing any important information. Okay, so. And I can go up here and I can switch from focus mode, you know, to ambient mode.

And now I can hear the other things going on. And in focus mode, it's just you and I. Arno. Yes, Pat? We won't amplify the background noise, so you can concentrate on our conversation, but we'll make sure you don't miss any important interruption, such if a delivery arrives.

Well, okay, maybe the grandkids are at the door. Hopefully my wife gets it, so we'll dismiss that there. So stay focused just on you and I. But the PC is still detecting. Yeah, your PC recognized the delivery on alerted you, but you choose to ignore it.

And the PC decided to keep the hearing aids in focus mode. Excuse me. Excuse me. Okay, well, okay, I think I'm gonna go talk to who's ever here. Just a second, Arno.

Sure. Hey, Diana. Hey, what's up? Why are you interrupting my demo? Do I need a reason?

It's always a good time to talk to me. Oh, I'm sorry. Thank you for taking the interruption. So, yeah, thanks for taking the interruption. Actually, I think I might have wasted my time waving because I think your computer actually heard the sound of me saying, excuse me.

And it even told you which direction that sound was coming from, so I just wasted a lot of energy. But I'm so glad you came over. Thank you. I think you've got a really great system over there. It sounds like with the head tracking AI, it actually detected when you left the conference call.

And so your hearing aids were automatically switched to ambient mode so we could chat. And then when you go back, it's going to automatically switch your hearing aids back to focus mode so you can be focused again on talking to Arno. So, pretty cool technology there. Very nice. Tracking AI is actually running on the MPU in the intel core.

And another thing we have running on the NPU is the summer and AI summarization technology. So it's actually, as you're talking to me, it's over there busily summarizing. So Arno's still talking and it's summarized. Arno's still talking and it's summarizing. It's condensing it for you.

And so you can just go back there and you'll be able to jump in really quick. You'll know exactly what transpired while you're gone. Very nice, very sleek and. Yeah. Okay.

Can I go back now or are we. Okay, so I know I'm back here now, right? And Diane was helping and we were chatting, but you were talking in French when I was gone. Yes, sorry, Pat. I was explaining the demo in French.

Oh, sorry. So now not only am I getting real time summarization of what I missed when I went out of Zoom, it brings me back to focus mode when I'm back to the PC. And it also did translation from French in real time as well. What do you think? Is this the AIPC generation?

So I love that Pat Gelsinger approaches this from an engineering mindset from someone who's the CEO from a Fortune 50 company who has the means to be able to recognize the challenges faced by those with hearing loss. Much has been given. You know, attention has been focused on oracast and the potential for ubiquitous connectivity. But I think the reminder in that, that video segment really is that you're still going to need signal processing to go from focus mode to ambient mode to do the translation and the transcription and summarization, as he said, as we transition in this AIPC generation, because in addition to being tethered to our cell phones, we're also tethered to our PCs in many cases for teams, meetings and Zoom calls and throughout our workday and play day. And I think that video really gives me a lot of optimism for the way that we'll see not only ubiquitous connectivity with oracast, but also the ways that we can improve the ease of use and alerting people, whether they're in a focus mode or an ambient mode, to hearing and the direction of sounds that are occurring in their environment.

It's really impressive. And look for those features to impact not only telephones, but PCs, televisions, and other ways that we can use this wireless connectivity in the next generation of products for this smart assistant functionality. So, in conclusion, the use of AI has been increasing rapidly, as I hope

you've seen in the last hour, whether it be for providing best speech sound quality and speech intelligibility, whether health and wellness from that fall detection feature and the signature of what happens when a person suffers a fall. But also the way that we're incorporating the steady exercises to try to ultimately provide a fall risk assessment that may, in the long run, help individuals at risk for falls, prevent falls from occurring before they happen. And then finally, that intelligent assistant, as we just saw, provides great opportunities for the ease of use for people with neuropathy and arthritis, as well as those who want to continue to work in the environment and use those features with their hearing aids as an intelligent assistant.

And so with that, we've come to the end of today's session. Look forward to seeing you again very soon.