

Signia Multi-adaptive Xperience with Acoustic Intelligence™: Holistic AI-driven enhancement of conversations in noise

Signia Multi-adaptive Xperience (MaX) introduces system-level artificial intelligence (AI) in hearing aids. We call it Acoustic Intelligence. Utilizing four deep neural networks (DNNs) that work in synergy to optimize the four key sound processing technologies offered by the MaX platform, Acoustic Intelligence enables highly adaptive sound processing across a broad range of acoustic environments and elevates speech handling to a new level, particularly in challenging group conversations in background noise. This paper details how Acoustic Intelligence works and optimizes the entire sound processing. We also describe the new eWindScreen 2.0 and SoundSmoothing 2.0 features designed to increase comfort without reducing speech clarity when wind noise and sudden transient sounds appear. All the features contribute to Signia's overarching goal of preserving speech clarity across all listening situations. We present a series of technical signal-to-noise ratio measurements that demonstrate the performance benefits of the Signia MaX platform in group conversations in noise. These results show substantial improvements offered by MaX compared to its predecessor platform, and they show that MaX outperforms premium hearing aids from key competitors.

Niels Søgaard Jensen, Barinder Samra, Homayoun Kamkar Parsi, Cecil Wilson

AUGUST 2026

Take-away messages

- Signia MaX is powered by Acoustic Intelligence that includes four separate but highly interconnected DNNs, controlling four key sound processing technologies.
- Signia MaX shows superior SNR performance, outperforming key competitors in a group conversation in noise, reaching an SNR benefit of 3.1 dB compared to the best performing competitor.
- SoundSmoothing 2.0 reduces sudden sounds while also preserving speech better than key competitors, showing reductions of the transient-to-speech intensity ratio by up to 66%.
- eWindScreen 2.0 reduces wind noise while enhancing speech better than before, showing a 74% improvement in the intensity ratio between speech and wind noise compared to Signia IX.

Introduction

Hearing aid technology has improved substantially over the last decades, leading to high levels of general wearer satisfaction. In the recent MarkeTrak 2025 survey (Picou, 2025), 79% of the surveyed hearing aid owners reported feeling satisfied with their ability to hear across all listening situations, compared to 43% of the surveyed non-owners.

However, the MarkeTrak 2025 data also showed that the satisfaction rates drop in challenging listening situations, with the lowest satisfaction rates – for both hearing aid owners and non-owners – observed for the situations “when trying to follow conversations in noise” and “in conversations with large groups”.

One of the key technologies that has been used to address the problems experienced by hearing aid wearers in challenging communication situations in background noise is Artificial Intelligence (AI), particularly noise reduction systems based on the application of deep neural networks (DNNs; e.g., Andersen et al., 2021; Fitzgerald et al., 2025; Ashkanichenarlogh et al., 2025). In short, DNNs are machine-learning models that learn complex relationships in data through multiple layers of processing, enabling advanced pattern recognition and signal processing capabilities (LeCun et al., 2015).

Various DNN-based noise reduction approaches, sometimes paired with conventional directionality, have in recent years been introduced in the hearing aid market. While these technologies have been shown to offer wearer benefits in certain situations (e.g., Moore, 2026), the development in the MarkeTrak ratings over time does not suggest that these new noise reduction technologies – seen across the board – have had a noticeable effect when it comes to hearing aids’ ability to effectively address the group-conversation-in-noise challenge. One reason for this may be that the scope of most of these systems has been rather one-dimensional: To provide noise reduction – and noise reduction alone. In some situations, this may offer a benefit to the wearer – but in other conditions, it may simply not be enough to satisfy the wearer and their needs in the given situation. The wearer may require support beyond noise reduction alone.

With the introduction of the Multi-adaptive Xperience (MaX) platform, Signia has taken a completely different and much more holistic approach when it comes to the application of AI and DNN-based hearing aid technology. Rather than ‘only’ focusing on noise reduction, the scope of the AI application in MaX is much broader. While other systems typically apply one DNN, which has a rather narrow purpose, and which is only active at certain times (either selected by the wearer as a second program or determined by the acoustic conditions), MaX includes four different DNNs that are active and control the sound processing at all times and in all listening situations.

The four DNNs in MaX have different purposes, and accordingly, they are trained differently. However, at the same time, they are highly interconnected and coordinated, and together, they are steering the sound processing in every situation encountered by the wearer – assuring that MaX fully adapts to the situation and the wearer’s specific needs in that situation, no matter what the situation is.

Together, the four DNNs constitute the ‘brain’ of MaX, and accordingly, we refer to the combined system as *Acoustic Intelligence*. With Acoustic Intelligence, MaX moves beyond the traditional sound classification offered by other hearing aids and provides a much more detailed description of the sound environment the wearer is in at any given moment. By introducing AI-based steering of four key processing technologies provided by MaX, Acoustic Intelligence offers end-to-end system-level intelligence that controls all aspects of the sound processing, including but not limited to noise reduction.

In the first part of this white paper, we will describe the MaX platform and in particular the concept of Acoustic Intelligence – how it controls and connects the different key sound processing technologies offered by MaX, and how it supports the wearer and helps them to perform at their best in the different situations they encounter in their everyday life – including even the most challenging ones when being in noisy group conversations.

The second part of the paper is devoted to the SoundSmoothing 2.0 and eWindScreen 2.0 features which are introduced on the MaX platform, and which offer substantial

improvements in the handling of transient sounds and wind noise. Furthermore, this part briefly explains the update of the MaX-Fit fitting rationale that comes with the MaX platform.

Part I: AI-driven signal processing

In the first part of the white paper, we will describe the MaX platform, Acoustic Intelligence, and the four AI-steered key technologies.

The MaX platform

Signia MaX is implemented on a new and highly advanced hardware platform that offers several significant performance improvements, such as 54 times more processing power, twice the memory, and around a 50% increase in efficiency compared to the prior Integrated Xperience (IX) platform. These performance improvements are required to support the new AI technology offering system-level intelligence as well as the other new or updated processing features made available by the MaX platform.

The MaX sound processing builds on the foundation laid out by Signia IX, including Signia's unique Augmented Focus (split processing) and RealTime Conversation Enhancement (RTCE) technology, which introduced multi-stream processing as a completely new way for hearing aids to address the challenges provided by group conversations in noise (Jensen et al., 2023). These features have been completely redesigned in the new structure, with several new features being added to improve the overall performance of MaX.

The introduction of Acoustic Intelligence represents a significant development in hearing aid signal processing. It is a fundamentally new way of applying – and benefitting from – AI in hearing aids. The system-level intelligence steers and optimizes most aspects of the MaX sound processing, and it increases the effectiveness and impact of the processing features. Beyond enhancing the features, Acoustic Intelligence synchronizes them to every moment. The result is improved and superior performance on several outcome measures.

Besides the introduction of the new Acoustic Intelligence technology, one of the most

significant technological enablers of some of the most important audiological MaX features is the binaural communication, which is offered by the new near-field magnetic induction (NFMI) link between the left and right hearing aids.

Since the introduction of the first ear-to-ear (e2e) communication system more than two decades ago, synchronizing and integrating the sound processing between the two hearing aids has been a Signia hallmark, which has paved the way for market-leading solutions within binaural processing. The new version of the ear-to-ear communication (e2e 5.0) provided by MaX has been significantly strengthened. A redesigned antenna has provided a 35% stronger signal between the ears, which supports a more robust and stable connection between the two hearing aids. This allows the hearing aids to work together consistently in real time, and it improves the working conditions for the various binaural algorithms in basically all situations. Furthermore, higher data rates enable the use of high quality (HQ) audio codecs for ear-to-ear transmission to improve the sound quality of both binaural sound processing and Bluetooth Classic streaming. Thus, in combination with Acoustic Intelligence, previously unseen performance levels of binaural processing can be reached.

Acoustic Intelligence

As already highlighted in the Introduction, four separate DNNs constitute the core of Acoustic Intelligence, and they directly control all stages of the signal processing. In particular, these four key technologies implemented on the MaX platform are controlled and enhanced by Acoustic Intelligence:

- RealTime Conversation Enhancement^{AI} (RTCE^{AI})
- Own Voice Processing^{AI} (OVP^{AI})
- Ambient Scene Adaptation^{AI} (ASA^{AI})
- Dynamic Noise Control^{AI} (DNC^{AI})

The four DNNs are briefly described in the text box below. A detailed description of the four key technologies, how they are controlled by Acoustic Intelligence, and how they contribute to the Signia MaX sound processing is provided in the following sections of this chapter.



Four DNNs powering Acoustic Intelligence

Speech Detection DNN is trained on a wide variety of speech signals, in different languages and presented in a wide range of different backgrounds. It allows accurate real-time detection of external talkers and enables a very precise steering of the speech processing, e.g. as performed by RTCE^{AI} and OVP^{AI}.

Own Voice Detection DNN is trained to detect when the wearer speaks. It directly controls the OVP^{AI}, and it thereby replaces the in-clinic training that was required for the previous generations of OVP. The Own Voice Detection DNN also contributes to the steering of the RTCE^{AI} processing.

Noise Analysis DNN is trained to detect the type of background noise experienced by the wearer at any given time, e.g., whether it is speech babble or traffic noise. Besides offering additional input for the control of speech processing (when speech is present), the DNN also steers the DNC^{AI} and thereby the handling of surrounding noise.

Scene Analysis DNN is trained to provide a continuous interpretation of the sound environment and the listening situation the wearer is situated in, for example, whether the wearer is in a car, or whether they listen to speech or music. This contributes, for example, to the steering of the ASA^{AI} that makes sure the processing of ambient sounds adapts to the most important wearer needs in any given situation.

Since the four DNNs serve different purposes, they are designed and trained differently. The DNNs have been trained on a total of 40 million sound samples, which included speech in 11 different languages, six music genres, and a very wide variety of different voices and sound

environments. The processing power provided by the MaX platform allows Acoustic Intelligence to scan the environment 100 times per second, allowing very fast analyses and reactions to changes.

The four DNNs are highly integrated and work as one unified AI unit, continuously interacting and sharing information. This means that it is the four DNNs in combination – the entire Acoustic Intelligence – that steers the four key features (i.e., RTCE^{AI}, OVP^{AI}, ASA^{AI}, and DNC^{AI}), ensuring a fully coordinated sound processing in all types of acoustic environments and listening situations. Thus, there is not a 1:1 connection between one DNN and one feature. It is analogous to the human brain, coordinating and synchronizing the body's vital functions, ensuring that all parts work together seamlessly.

Figure 1 illustrates the implementation and integration of Acoustic Intelligence in the MaX platform. While the diagram does not include all the actual elements, features and interactions implemented on the MaX platform, it illustrates the basic principles of how Acoustic Intelligence controls the four key technologies of MaX. This is shown by the bold red arrows, while the gray arrows indicate signal paths (full lines) and control paths (dotted lines), respectively.

The diagram in Figure 1 shows how the acoustic signals (from the split processing shown to the left) are pre-processed to provide the appropriate input for Acoustic Intelligence. Acoustic Intelligence then steers the four blocks which either directly process the conversation and surrounding streams or control the processing which takes place in the *Continuously Adapted Output* stage that combines the output from the four technology blocks and optimizes the listening experience for the wearer.

In the following sections, the four key technologies and the underlying subfeatures (indicated as red boxes in the diagram) will be further described.

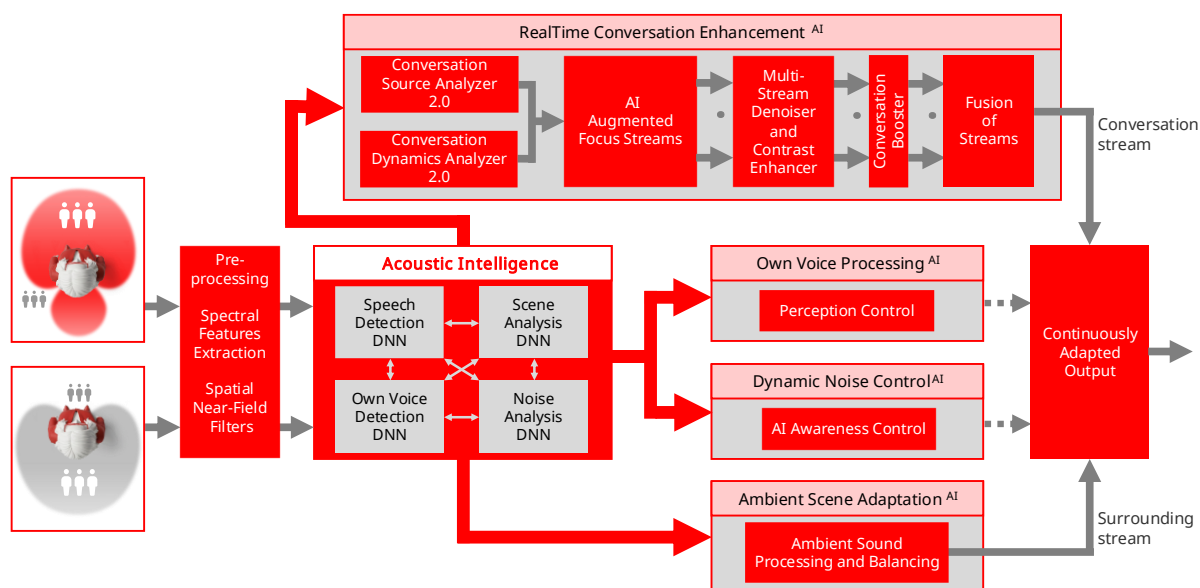


Figure 1. Diagram showing how Acoustic Intelligence with its four DNNs steers the four key technologies implemented on the Signia MaX platform: RealTime Conversation Enhancement^{AI}, Own Voice Processing^{AI}, Dynamic Noise Control^{AI}, and Ambient Scene Adaptation^{AI}. The diagram does not include all the processing features implemented on the MaX platform, and it does not include all the connections and dependencies between the (highly interconnected) blocks.

Four key technologies

In the following, we will describe each of the four key technologies, which are steered by Acoustic Intelligence.

RealTime Conversation Enhancement^{AI}

As one of the MaX platform's key technologies, RTCE^{AI} offers advanced signal processing that is designed to directly address the most problematic situations experienced by many hearing aid wearers: Group conversations in background noise. For the first version of RealTime Conversation Enhancement (RTCE), which was introduced on the IX platform, research showed significant wearer performance benefits, both in lab tests (Jensen et al., 2023; Jensen et al., 2025) and in the real world (Folkeard et al., 2024; Jensen et al., 2024a).

The AI technology available on the MaX platform allows the RTCE^{AI} technology to make a significant leap forward compared to what was provided by the IX platform. Acoustic Intelligence now enables a more advanced acoustic analysis of the conversational situation, and it allows new AI-controlled signal processing features to be added to RTCE^{AI} to further sharpen the speech-in-noise performance. On top of that, the existing features already known from IX have been further

developed to benefit from the multi-dimensional AI support offered by Acoustic Intelligence.

The RTCE^{AI} technology is shown as the largest gray box in Figure 1. As mentioned, it includes elements well-known from Signia IX as well as completely new elements.

The Signia split processing approach, allowing independent processing of sound streams originating from the front (focus) and back (surrounding) hemisphere, respectively (as indicated in the left side of Figure 1), is still an important foundation for the RTCE^{AI} sound processing. It provides the input to RTCE^{AI} – but now with Acoustic Intelligence being added to analyze the input and control and optimize the subsequent signal processing in RTCE^{AI}.

When RTCE^{AI} is active in conversations, updated versions 2.0 of the *Conversation Source Analyzer* and *Conversation Dynamics Analyzer* analyze the conversation setup and determine the location of talkers as well as the conversation dynamics. The analysis has been strongly improved by Acoustic Intelligence where especially the Speech Detection DNN allows for a better and more precise analysis of the given conversation situation.

This information is used to steer the *AI Augmented Focus Streams*. Compared to the RTCE version known from IX, three more AI-controlled focus streams (enabled by binaural beamforming) have

been added to make the system even more responsive in dynamic conversations with multiple talkers. Steered by Acoustic Intelligence and especially the Speech Detection DNN and Own Voice Detection DNN, which controls the strength of the applied beams, the new streams will enable a better isolation (and more augmentation) of the conversation partners. This will make the talkers stand out as more prominent, improving the signal-to-noise ratio (SNR) and thereby potentially improving the wearer's ability to follow the conversation. This is illustrated schematically in Figure 2, which illustrates a group conversation where the conversation partner directly in front of the wearer is the active speaker. In the figure, the activated AI-controlled beam is indicated as [1], showing the narrower focus on the detected talker compared to the broader beam from the original RTCE. The broader beams are still part of the system where they improve the separation of talkers and reduce diffuse background noise between the talkers.

Besides the AI Augmented Focus Streams, two other new processing features have been added to the RTCE^{AI} signal processing. In Figure 1, these features are shown as one block, *Multi-Stream Denoiser and Contrast Enhancer*.

Controlled by Acoustic Intelligence and informed by the operating state of all focus streams, the Multi-Stream Denoiser applies additional binaural noise reduction to further enhance speech quality by suppressing residual background noise. A key functionality is its dynamic estimation of the optimal maximum noise reduction level, computed in real time from the combined influence of the focus streams and Speech Detection DNN activity. This adaptive mechanism maximizes noise suppression while preventing speech distortion and preserving overall sound quality. The attenuation of the background noise is shown as [2] in Figure 2.

The Multi-Stream Contrast Enhancer dynamically adapts to the conversation flow, steered by Acoustic Intelligence, and increases the contrast between speech and background noise by allowing more linear gain when speech is present. This leads to an even stronger focus on the talker, making the talker stand out as even more prominent in the background of noise. This is illustrated as [3] in Figure 2.

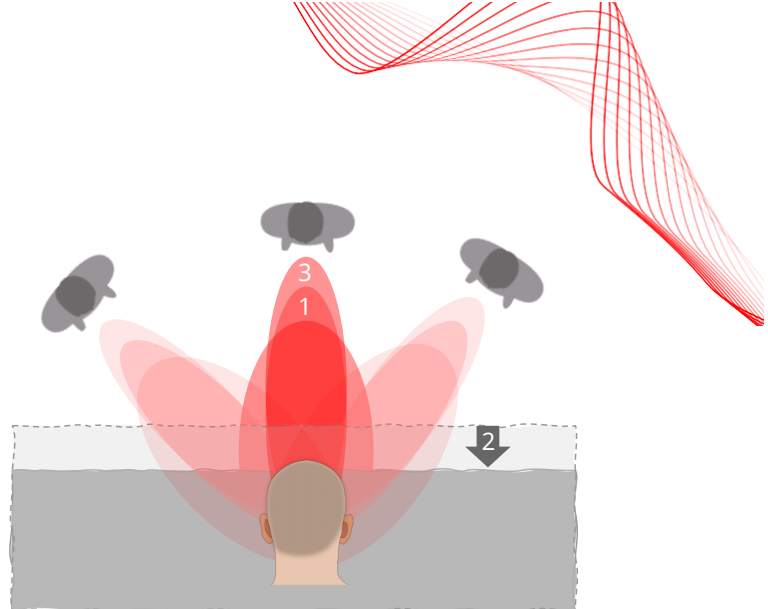


Figure 2. Schematic illustration of the effects of the AI Augmented Focus Streams (1) and the Multi-Stream Denoiser (2) and Contrast Enhancer (3), which are added to the previous version of RTCE known from Signia IX. The three new features all contribute to an improved SNR.

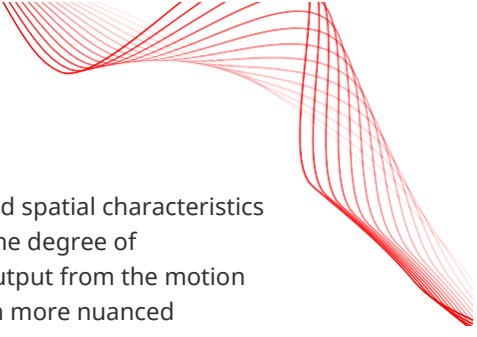
The remaining RTCE^{AI} blocks shown in Figure 1 are known from the IX platform. The *Conversation Booster* controls the enhancement of speech within each stream and optimizes the balance between the streams, and in *Fusion of Streams*, the separate focus streams are merged into one Conversation Stream.

Own Voice Processing^{AI}

Own Voice Processing (OVP) has for almost a decade been a unique feature in Signia hearing aids, leading to improved own voice perception for wearers (Powers et al., 2018). Furthermore, a recent data analytics study of Signia IX wearers showed a substantially lower return rate among those with OVP activated, compared with those whose hearing aids included OVP but did not have the feature activated (Samra et al., 2025).

The purpose of OVP is two-fold: 1) to increase the comfort for the wearer when listening to the sound of their own voice, and 2) to manage the gain when the wearer speaks such that the wearer does not lower the volume of their own voice and thereby puts a conversation partner in disadvantage.

However, one limitation in previous versions of the OVP feature has been that they required dedicated individual training of the system to be completed in the clinic during the hearing aid fitting for the own-voice detection to work. As a result, not all wearers were fitted with OVP, even when their hearing aids included the feature as an option.



With the introduction of Own Voice Processing^{AI}, the need for in-clinic training has become a thing of the past. In Signia MaX, the own voice detection is now managed automatically by Acoustic Intelligence, and in particular by the Own Voice Detection DNN, which has been trained on a large number of non-speech and speech samples, including both the wearer's own voice and the speech of other talkers. The DNN is trained on the output of a set of spatial near-field filters designed to distinguish between own voice and external speech from the front. This enables the DNN to differentiate between the wearer's voice and the voices of other nearby talkers in front of the wearer. These filters are part of the pre-processing block shown in Figure 1, which provides the input to Acoustic Intelligence. For talkers to the side of the wearer, spatial lateral detectors assist in making the decision that the detected speech is not produced by the wearer.

Besides eliminating the need for individual training of the own voice detection system, the new DNN-based detection also provides a very high accuracy of the detection. For example, testing shows that the false-detection rate (the likelihood that own voice is detected when others in front of the wearer speak) is lower than 1% with the new DNN-based detection.

The Own Voice Processing^{AI} block in Figure 1 uses the information from Acoustic Intelligence to determine the appropriate processing of the own voice signal. The own voice signal is primarily captured in the focus (front) stream of the split processing, so the entire focus stream processing will depend on whether the wearer speaks or not.

When the wearer speaks, it will typically result in a reduction of the gain, leading to a more comfortable, pleasant and natural experience of their own voice. The amount of gain reduction and the impact on the additional adaptive processing in the hearing aids will depend on the situation as determined by the Acoustic Intelligence and the other control and processing features shown in Figure 1. This allows Own Voice Processing^{AI} to seamlessly adapt to the situation.

Ambient Scene Adaptation^{AI}

Ambient Scene Adaptation^{AI} allows MaX to adapt to the variety of different sound environments encountered by the wearer. Powered by Acoustic Intelligence, it continuously analyses, for example,

the spectral, temporal and spatial characteristics of the incoming sound, the degree of reverberation, and the output from the motion sensor to provide a much more nuanced interpretation of the ambient sound scene than what is possible with traditional sound classifiers that operate with a limited number of sound classes and a limited number of parameter changes.

The overall aim of the Ambient Scene Adaptation^{AI} is to create a soundscape with an appropriate balance between the different elements appearing, whether it is speech, music, or other types of ambient sounds. Thus, it is not a matter of turning down ambient sounds. It is a matter of finding the appropriate balance between different sounds, not least the balance between the conversation stream and the surrounding stream, which is controlled by the *Ambient Sound Processing and Balancing* block shown in Figure 1.

Dynamic Noise Control^{AI}

The Dynamic Noise Control^{AI} technology implemented on the MaX platform controls the handling of surrounding noise. A central element is the *AI Awareness Control*, which – steered by Acoustic Intelligence – differentiates between different types of background noise and adjusts the processing to give the wearer the level of awareness of the surroundings that is appropriate in the given situation, e.g., depending on whether speech is present, and whether the background is traffic or babble noise.

One of the ways to adjust the noise is by controlling the strength of the effect provided by the split processing and the multiple streams. This allows more transparency and less noise reduction to be provided when walking on a busy street with traffic (where situational awareness is needed for basic safety reasons) than when being in a restaurant with a background of competing speech babble.

Acoustic Intelligence is trained on a wide variety of different acoustic situations and wearer needs, and the level of noise control will therefore adapt to the exact situation the wearer is in and to the listening requirements in this specific situation. This is illustrated schematically in Figure 3 that shows the effects of Dynamic Noise Control^{AI} and AI Awareness Control. The figure shows how the wearers' access to the ambience is adjusted

continuously and changed when going from one type of sound scenario to another. In the depicted example, the wearer is first reading a book in a noisy café. The wearer is not in a conversation but wants to be connected to the sound environment. Accordingly, around 40% of the ambient sound awareness is maintained relative to omnidirectional processing with no noise reduction.

When the wearer enters a restaurant and engages in a group conversation, the situation is completely different, even though the overall sound level is the same as in the noisy café. Acoustic Intelligence combined with Conversation Source Analyzer 2.0 and Conversation Dynamics Analyzer 2.0 (from RTCE^{AI}) detect the environment, recognize that a conversation is taking place, and identify that the wearer is actively participating. As a result, the surrounding sounds are reduced by Dynamic Noise Control^{AI} to a degree where approximately 10% of the ambient sound awareness is retained. This helps the wearer focus on the group conversation while still maintaining some awareness of the surroundings.

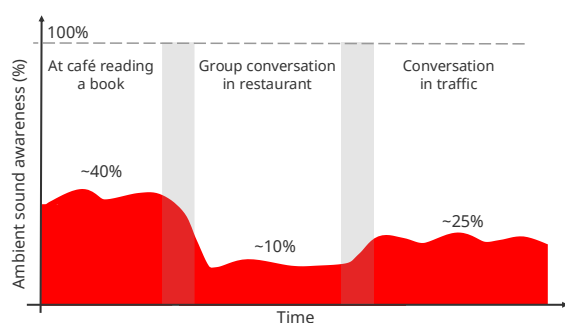


Figure 3. Schematic illustration of how Acoustic Intelligence controls the loudness of the surroundings – and thereby the ambient sound awareness – when the wearer moves between three different sound scenes, all having the same average sound pressure level of the surroundings. The gray areas indicate the transition between scenes. The momentary awareness is plotted as a percentage of the maximum possible awareness in that scene. A value of 100% would be reached with omnidirectional processing and no noise reduction.

When the wearer leaves the restaurant and continues the conversation in a busy street – with a noise level similar to the levels observed in the café and in the restaurant – Acoustic Intelligence will both know that the wearer is in a conversation (with the need to understand speech) and that they are close to traffic (with the need to hear what is going on around them). Consequently, the

access to the ambient sound is increased by Dynamic Noise Control^{AI} from 10% (in the restaurant) to 25% to allow more ambient sound to come through – but without reaching the level in the café to make sure that the conversation can still be followed.

The process illustrated in Figure 3 takes place all the time and in all types of surroundings. Acoustic Intelligence constantly monitors the acoustic environment and provides the information and decisions needed to adapt to the wearer's needs whenever the situation changes.

What the research shows

Following the introduction to the MaX platform, Acoustic Intelligence and the four key processing technologies, we will now present a technical study investigating the performance of MaX in the main target scenario: A group conversation in noise. This study only focused on technical performance. Human performance with MaX was investigated in separate studies reported by Jensen et al. (2026) and Korhonen et al. (2026), see also the Discussion section of this paper.

To assess the speech-in-noise processing offered by MaX in a noisy, complex and challenging group conversation situation, a series of technical measurements were done in a conversation scenario simulated in the lab. The well-established Hagerman procedure (Hagerman & Olofsson, 2004) was used to investigate the output SNR provided by MaX in a comparison with Signia IX and four key competitors, when the hearing aids were worn in a noisy group conversation scenario with alternating talkers simulated in a loudspeaker setup in the lab.

Method

The test methodology is comparable to that used in previous experiments (e.g., Jensen et al., 2024b). When a mix of speech and noise is presented to the hearing aid under test, recording the combined signal (speech and noise) on the output side of the hearing aid, with and without phase inversion of the input signals, the Hagerman methodology makes it possible to isolate the processed speech and the processed noise. By averaging across the isolated speech and noise signals, respectively, it is possible to provide a precise estimate of the actual output

SNR that would be experienced by a wearer in that sound environment.

The measurement setup, shown in Figure 4, was established in a sound-treated room and included a KEMAR manikin (in the center of the setup) and four loudspeakers positioned at a distance of 1 m. Sections of the International Speech Test Signal (ISTS; Holube et al., 2010) were presented alternately from two loudspeakers at 0° and 315° azimuth, simulating turn-taking in a conversation, at a level of 75 dB(A) (measured at the position between the KEMAR ears without the presence of the KEMAR). At the same time, a background noise consisting of a recording made in a busy cafeteria mixed with pink noise was presented from two loudspeakers at 135° and 225° azimuth. The measurements were done at two different levels of background noise, 75 dB(A) and 70 dB(A), respectively. This resulted in overall input SNRs of 0 dB and +5 dB for the test environment.

To determine the output SNR, the hearing aids were fitted on the KEMAR ears. A series of recordings were made with and without phase inversion of the different input signals. Applying the phase-inversion technique on the recordings then allowed estimation of each of the processed speech and noise signals at the output of the hearing aids (i.e., as measured in the KEMAR ear).

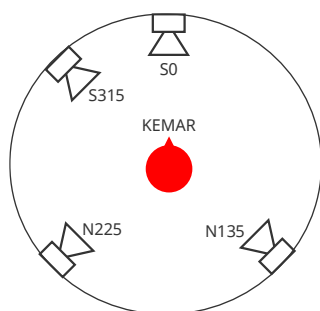


Figure 4. Test setup used for output SNR measurements. Speech (S) signals were presented alternately from the two loudspeakers in the front hemisphere, while continuous noise (N) signals were presented from two loudspeakers in the back hemisphere. The signals processed by the hearing aids were recorded in the KEMAR ears, with and without phase inversion of the input signals, and the Hagerman method was used to generate estimates of the combined speech signals (S0 + S315) and noise signals (N135 + N225).

In the measurements, a conversation scenario was simulated by presenting the speech signal alternately from the two loudspeakers, with the signal being presented from a given direction for 10 seconds before switching to the other

direction. Each recording included two sections for each direction, for a total of 40 seconds. Prior to starting the recording, the sound signals from all four loudspeakers were on for 80 seconds to allow time for all the hearing aids to settle.

The outcome of the analysis (for each pair of hearing aids) is the output SNR, averaged across the entire evaluation time. Due to the test layout where speech was presented from the front and from the left side of the KEMAR, we present the output SNRs of the left hearing aid, which is most relevant to speech understanding due to the better ear effect. In the analysis, we calculated the overall conversation SNR, $(S0+S315)/(N135+N225)$. When presenting the results, a Speech Intelligibility Index (SII) weighting (American National Standards Institute, 2024) is applied to give more weight to the frequencies most important to speech understanding.

Hearing aids

In the test, Signia Pure C&G MaX was compared to Signia Pure C&G IX as well as RIC hearing aids from four different main competitors (Brand A-D), which, at the time of measurement, represented the most recent premium RIC hearing aids offered by the respective manufacturers.

For the measurements, all hearing aids were programmed to a symmetrical, moderate and sloping hearing loss (close to the N3 hearing loss, Bisgaard et al., 2010), using the default setting prescribed by each manufacturer's recommended (proprietary) fitting rationale. With one exception, all hearing aids were tested in their Universal program. The exception was Brand D, which did not offer DNN-based noise reduction in its Universal program but rather in a separate manual speech-in-noise program. Since the other competitor hearing aids allowed for automatic switching to the DNN-based noise reduction in the Universal program, we decided to test Brand D in the dedicated speech-in-noise program, where the DNN noise reduction was turned on all the time. We also measured Brand D in its Universal setting, but since performance was poorer in the Universal program, we only report the result obtained with the manual program offering DNN noise reduction.

To ensure valid application of the phase-inversion technique, features manipulating the signal phase (feedback cancellation and frequency

compression) were deactivated in all the hearing aids. The hearing aids were fitted to the KEMAR ears using closed coupling eartips.

As baseline, recordings were also made in the open (unaided) KEMAR ears. This allowed the open ear SNR to be calculated and used as a reference. Thus, the results will be presented as the SNR improvement relative to the SNR measured in the open KEMAR ear.

Results

The improvement in the conversation SNR for each hearing aid, relative to the SNR in the unaided KEMAR ear, is plotted in the bar graph in Figure 5. The results for 0 dB input SNR are shown in the lower part, while the results for +5 dB input SNR are shown in the upper part.

Figure 5 shows that while all hearing aids offered an SNR advantage compared to unaided listening (as indicated by all values being positive), there were substantial performance differences between the hearing aids. At both input SNRs, Signia MaX offered the largest SNR benefit, reaching levels of 10.5 dB and 9.3 dB, respectively.

In these simulated group conversation scenarios, the MaX performance was between 0.8 dB and 1.3 dB better than the performance of Signia IX. At 0 dB SNR, the MaX performance was 2 dB better than the best performing competitor (Brand A), while at +5 dB SNR, the MaX advantage reached 3.1 dB compared to the best competitor (again Brand A). Two other competitors (Brand B and Brand C) were around 3 dB and 4 dB, respectively, below MaX at 0 dB SNR, while they were at the same level, 3.5 dB below MaX, at +5 dB SNR. Compared to the other brands, Brand D lagged somewhat behind at both input SNRs, showing SNR improvements of only 5.5 dB and 3.4 dB in these scenarios, which are around 5 dB and 6 dB below the performance levels of MaX.

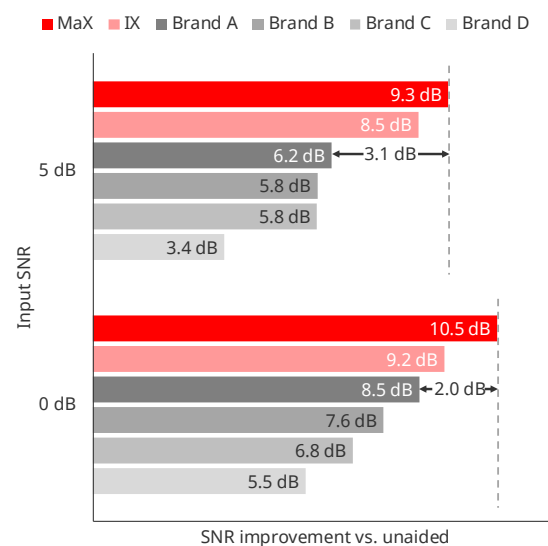


Figure 5. SII-weighted conversation SNR improvement at input SNRs of 0 dB and +5 dB, relative to the unaided KEMAR, measured in the left ear across the two talker positions, for Signia MaX, Signia IX, and four competitor RIC hearing aids. All hearing aids were measured in their default Universal program, except for Brand D that was measured in a manual program with DNN noise reduction activated.

The overall result pattern was the same at the two input SNRs. MaX performed best, showing a clear gap to IX, which showed an equally clear gap to the competitors. The competitors were ranked in the same order at both SNRs, although there were some minor variations in the magnitude differences across the two input SNRs. All hearing aids offered their largest SNR improvement at the lowest input SNR (0 dB).

The clear SNR effects of MaX in the simulated group conversation scenario can be directly linked to the effects of Acoustic Intelligence and particularly RTCE^{AI} with its new AI-controlled features. These updates not only offer a clear benefit compared to the predecessor IX platform. They also provide substantial benefits compared to the competitors, indicating the advantage of the system-level intelligence and multi-stream processing provided by MaX, compared to the feature-level intelligence provided by the competitors.

Part II: Additional new MaX sound processing features

The MaX platform also introduces other new signal processing features, which do not depend directly on Acoustic Intelligence, but which still contribute significantly to optimizing the wearer experience – and particularly the handling of speech – in all types of situations. The new features include eWindScreen 2.0 and SoundSmoothing 2.0, which take the wind noise handling and transient sound protection offered by MaX to new levels. In the following sections, we will briefly describe these features and present the results of technical measurements that demonstrate their performance benefits. Furthermore, we will briefly explain the new MaX-Fit fitting rationale.

SoundSmoothing 2.0

Most hearing aid wearers find transient sounds both annoying and disturbing, and the handling of such sounds in hearing aids is just as important as the handling of more stationary noise (Hernandez et al., 2006). The MaX platform includes SoundSmoothing 2.0, which offers significant improvements in the handling of transient sounds compared to the feature known from previous platforms.

Whereas the previous version only applied transient handling in two broad bands, SoundSmoothing 2.0 increases the resolution substantially by applying transient reduction in each of the 48 processing channels. This means that the gain is attenuated only in the narrow frequency bands where the transient is present, without affecting speech signals in adjacent channels. Furthermore, the e2e 5.0 communication between the hearing aids includes a new fast (low latency) data channel that paves the way for binaural synchronization of the transient handling in the two hearing aids, both when it comes to detection and attenuation. This allows for more stable interaural level differences (ILDs), enabling the wearer to better localize transient sounds due to better preserved ILDs. Tests suggest that SoundSmoothing 2.0 can reduce the RMS localization error caused by ILD variation by around 30% compared to its predecessor, allowing more precise localization of transient sounds, especially in the high frequency

range. Importantly, not only transient sounds are more stable. Other signals also remain more stable during transient sounds due to the synchronization.

The fast binaural communication between the hearing aids essentially allows a “look-ahead” function. When a transient is detected at the input stage of one hearing aid, the communication between hearing aids (which is faster than the filter bank processing time) passes the information over to the output stages of both hearing aids, preparing the output stages to effectively apply attenuation to the entire transient.

To investigate the effect of SoundSmoothing 2.0 on MaX's ability to preserve speech when transient sounds occur, a series of technical measurements was conducted in the lab where the test hearing aids were fitted on the ears of a KEMAR manikin. A recording of a continuous dialog was presented from a loudspeaker at 60° azimuth at 57 dB(A), while recordings of different transient sounds (e.g. the sounds of a knife on a table or a cup on a plate) were presented from a loudspeaker at 300° azimuth at an average level of 64 dB(A), with peaks exceeding 80 dB(A). Recordings of the hearing aid outputs were made in the ears of the KEMAR, and by doing the recordings with and without phase inversion of the speech and noise signals, the aforementioned Hagerman method was applied to isolate the aided speech and transient sound signals, respectively.

Signia Pure C&G MaX, Signia Pure Charge&Go IX, and the (at the time of measurement) most recent premium RIC hearing aids from four competitors (Brand A-D) were measured. All hearing aids were programmed to the same moderate sloping hearing loss (close to the N3 hearing loss, Bisgaard et al., 2010) using the manufacturers' proprietary fitting rationales. Features that could disturb the signal phase were turned off, while all other adaptive features – including features on transient handling – were left in their default settings.

The results of the measurements are shown in Figure 6. The bar graph shows the broadband attenuation of the transients (averaged across all the recorded transients and across left and right ear) relative to the dialog. The values are normalized to the open-ear KEMAR result. This

means that the value of -12.5 dB observed for MaX indicates that the ratio between transient noise and speech is 12.5 dB lower than in the unaided measurement. The reduction of the transient-to-speech ratio (or increase in the speech-to-transient ratio) is an indication of how much better the speech is preserved in the presence of transient sounds.

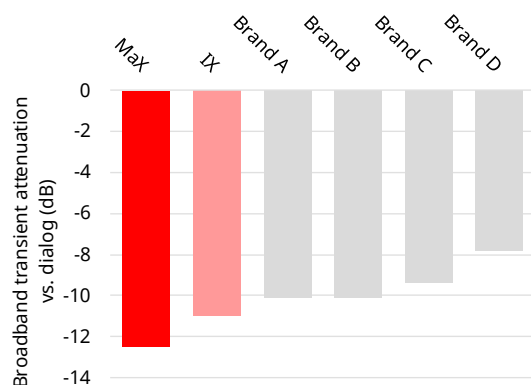


Figure 6. Broadband attenuation of the transient-to-speech ratio, averaged across different types of transients. The attenuation is indicated relative to the transient-to-speech ratio measured in the open KEMAR ear. Lower (more negative) values indicate more attenuation.

The results in Figure 6 show that Signia MaX offers the highest relative attenuation of transient sounds. Signia MaX provides a 12.5 dB attenuation of the transient-to-speech ratio compared to the open ear KEMAR, and SoundSmoothing 2.0 brings a significant 1.5 dB benefit compared to Signia IX. Compared to the four competitor brands, the benefit is even higher, ranging from 2.4 dB (Brand A) to 4.7 dB (Brand D). A further 2.4 dB of extra attenuation of transients relative to speech corresponds to a 42% reduction of the intensity ratio between transients and speech, while a 4.7 dB reduction corresponds to a 66% reduction of the intensity ratio. Thus, compared to the competitors, MaX offers superior preservation of speech in the presence of transient sounds.

eWindScreen 2.0

The MaX platform also provides a significant update when it comes to protecting the wearer against wind noise without sacrificing speech perception.

The new eWindScreen 2.0 introduces detection and attenuation of wind noise in each of the first 16 processing channels, which cover the frequency range affected by wind noise. This marks substantial progress compared to the

previous version's broadband detection.

Furthermore, the eWindScreen 2.0 performance is improved by operating on a cleaner input signal, since it adaptively fades to the microphone with the least wind noise (out of the two microphones available in each hearing aid). This is particularly beneficial when the wind comes from one direction, e.g., during walking, running or riding a bike.

In combination, the updates enable a stronger, more accurate, and more stable attenuation of wind noise. This allows for better preservation of the quality of the target sounds and – if speech is present – the speech cues that are important to speech understanding. This will have a positive effect in all types of windy situations, but in particular during conversations in wind.

In an objective wind tunnel test, we investigated the ability of eWindScreen 2.0 to handle wind while preserving speech. We used Signia IX and the previous version of eWindScreen as reference in these measurements.

In the test, the hearing aids were fitted on the ears of a KEMAR manikin, which was on axis to a wind tunnel producing a constant wind speed of 5 m/s, corresponding to a moderate breeze. At the same time, speech from a male speaker was presented from a loudspeaker on the left side of the KEMAR (at 315° azimuth) at a level of 70 dB SPL. Recordings of the processed combined sound (speech and wind noise) were made in the left (better) ear of the KEMAR.

The Hagerman phase-inversion procedure was applied to estimate the isolated wind noise signals on the output side of the hearing aids. This was done by doing two recordings, one with and one without phase inversion of the speech signal. In this way, the speech could be eliminated by adding the two recordings, allowing calculation of the long-term spectrum and level of the processed wind noise. Since wind noise is not acoustically reproducible, a similar Hagerman phase-inversion recording approach could not be used to estimate the speech signal. Instead, 256 independent recordings were made, and by averaging across all the recordings (and assuming independence between the random wind noise samples), the wind noise could be averaged out to an extent that allowed a precise estimate of the processed speech signal.

For the measurements, both test hearing aids were programmed to a flat hearing loss of 50 dB HL using the proprietary fitting rationales. The measurements with both pairs of hearing aids were done in two settings, with eWindScreen turned on (in the default medium setting) and off, respectively. For both hearing aids, this allowed estimation of the improvement in the output SNR and estimation of the magnitude of the wind noise reduction. Both outcomes were measured as a function of frequency, and a Speech Intelligibility Index (SII) weighting (American National Standards Institute, 2024) was applied to obtain one outcome value representing the overall performance.

The SII-weighted SNR improvements and noise reductions are plotted in Figure 7. The length of the red and gray bars indicates the SNR improvement and the amount of noise reduction, respectively. Thus, more negative values of the gray noise reduction bars indicate more noise reduction.

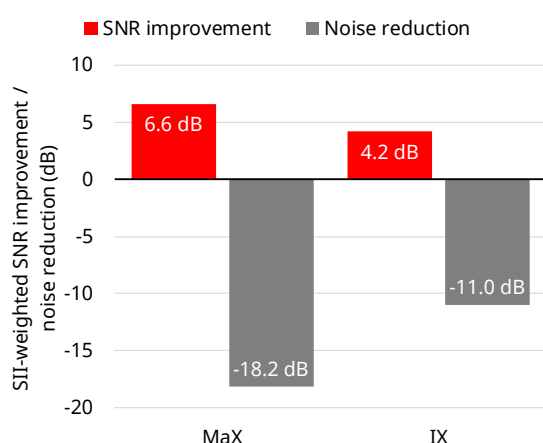


Figure 7. The SII-weighted SNR improvement (red bars) and noise reduction (gray bars) for Signia MaX (eWindScreen 2.0) and Signia IX (eWindScreen). Note that higher is better (more SNR improvement) for the red bars, while lower is better (more noise reduction) for the gray bars.

Figure 7 shows the SNR improvement provided by MaX was 6.6 dB, compared to 4.2 dB for IX, in both cases compared to turning off the feature. Thus, MaX provides a 2.4 dB improvement in speech-to-wind-noise ratio compared to IX. This corresponds to a 74% improvement in the intensity ratio between speech and wind noise.

The figure also shows that MaX reduced the wind noise by 18.2 dB, while the reduction was 11.0 dB for IX, resulting in a 7.2 dB advantage provided by MaX compared to IX.

The results of this study show that MaX offers a considerable amount of comfort by decreasing the level of wind noise by 18.2 dB, but at the same time, it offers a substantial improvement in SNR. Thus, MaX provides greater protection against wind noise while better preserving speech – both compared to turning the feature off and compared to the previous version implemented in IX. This is an important contribution to MaX's overall aim of improving the handling of speech in noise, supporting conversations everywhere – also outdoors in windy conditions.

MaX-Fit rationale

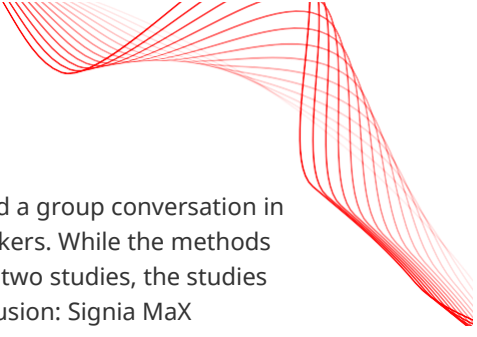
When turning on their Signia MaX hearing aids, the wearer will not only experience the impact of Acoustic Intelligence and the variety of AI-controlled processing technologies described in Part I of this paper. The foundation upon which all signal processing operates is built by the new fitting rationale, MaX-Fit.

Developed specifically for the MaX platform, MaX-Fit combines optimized target curves to deliver a crystal-clear, natural, and uncompromised sound experience. For example, the fact that the OVP^{AI} is now automatically available to all wearers eliminates the traditional fitting trade-off, allowing MaX-Fit to prescribe optimal gain for external speech (and other sounds) while maintaining a comfortable level of the sound of the wearer's own voice. This makes an important contribution to MaX's handling of speech, which is designed to offer clear speech perception in both quiet and noisy environments.

In addition, MaX-Fit also includes updated simulation methods for improved match-to-target performance. By providing a better alignment between prescribed and actual real-ear insertion gain, MaX not only ensures that wearers receive the intended amplification and sound experience but also simplifies in-situ fittings.

Discussion

In this paper, we have presented the Signia MaX platform with Acoustic Intelligence, which introduces AI-driven system-level intelligence in hearing aids. The introduction of the world's first hearing aids with four DNNs having different purposes and being trained in different ways, but at the same time being interconnected and working as the intelligent core of the system,



marks a completely new way of applying AI in hearing aids.

Acoustic Intelligence allows Signia MaX to increase the focus on what always has been a strong Signia focus point: Processing of speech – in all types of environments. In particular, Acoustic Intelligence and its control of the four key processing features provided by MaX offers improved performance in group conversations in noise. This was clearly demonstrated in the technical measurements of the output SNR in a noisy group conversation simulated in the lab. The results indicated superior SNR performance of MaX, both compared to the previous Signia IX platform and compared to the latest RIC hearing aids from four main competitors. The same overall pattern of results was observed at two different input SNRs, 0 dB and +5 dB. These SNRs cover a range that many wearers will encounter in their everyday lives (Smeds et al., 2015), and where they often will struggle to follow a conversation and therefore need the best possible support from their hearing aids.

The fact that the Signia IX performance also was shown to be superior to the competitors once again demonstrates the solid foundation laid out by the split processing and multi-stream processing that is already provided by IX. Signia MaX adds to this by offering the effects of Acoustic Intelligence, which enhances the processing features and steers them for the moment.

While there often is a link between improved SNR and improved speech understanding, the observed improvement of the measured output SNR cannot in itself be used to conclude that MaX provides better speech understanding, since speech understanding is affected by other factors than just SNR. Thus, human performance studies are needed to conclude on speech understanding.

To be able to draw such conclusions, two human performance studies were carried out to investigate speech understanding in noisy group conversations simulated in the lab. One study was done at ORCA Labs in Chicago, USA, and the other was done at Hörzentrum Oldenburg in Germany. Both studies compared Signia MaX to Signia IX and two of the four competitors included in the SNR measurements presented in this paper (including the best performing competitor in the SNR study). In both studies, a speech test was

carried out that simulated a group conversation in noise with alternating talkers. While the methods used varied between the two studies, the studies provided the same conclusion: Signia MaX provided significantly better speech understanding in the simulated group conversation in noise than both Signia IX and the competitors included in the studies. Thus, the studies showed that the SNR improvements presented in this paper indeed have a positive effect on speech understanding in a group conversation scenario. More details about the two studies are provided by Korhonen et al. (2026) and Jensen et al. (2026).

In combination, the results from the technical SNR measurements and the human performance studies demonstrate the advantages of MaX in group conversations in noise. The observed SNR and speech understanding improvements, compared to both IX and competitors, can be directly linked to the effects of Acoustic Intelligence and the AI-steered processing provided by MaX.

Besides the effects of Acoustic Intelligence, this paper has also presented results driven by new MaX processing features that do not rely directly on Acoustic Intelligence. The technical studies on SoundSmoothing 2.0 and eWindScreen 2.0 have shown the potential of these features to provide increased comfort to the wearer by efficiently detecting and reducing transient sounds and wind noise. But just as importantly, the results for both features also demonstrate how the aim of MaX is not just to reduce noise, but also to preserve speech, as demonstrated by the reduction of the transient-to-speech ratio in the SoundSmoothing 2.0 measurements and the improvement of the SNR in the eWindScreen 2.0 measurements. These effects contribute strongly to the overall aim of Signia MaX to offer the wearer a perfect balance between speech and noise in all types of sound environments.

Summary

With the introduction of system-level intelligence, Signia MaX with Acoustic Intelligence offers a completely new and holistic approach to AI-driven signal processing in hearing aids, allowing optimization of the sound for the full moment, not just the noise. Controlled by four different but

highly interconnected DNNs, Acoustic Intelligence and the four key technologies provide improved support to the MaX wearer in all types of listening and communication situations, including challenging group conversations in noise.

Compared to Signia IX and four main competitors, a technical study showed that Signia MaX outperforms both IX and all the competitors when it comes to improving the SNR in group conversations in noise. The SNR advantage vs. the best performing competitor reached 3.1 dB at an input SNR of +5 dB. The performance of Signia IX was also better than all four competitors, showing the power of multi-stream processing.

Technical studies of the new SoundSmoothing 2.0 and eWindScreen 2.0 features both demonstrated the ability of MaX to protect the wearer against unwanted sounds without sacrificing speech. The results showed that SoundSmoothing 2.0 reduced the transient-to-speech ratio by up to 4.7 dB more than competitors, corresponding to a 66% reduction of the transient/speech intensity ratio, while eWindScreen improved the signal-to-wind-noise ratio by 2.4 dB compared to Signia IX, corresponding to a 74% improvement of the intensity ratio between speech and wind noise.

In combination, the technical measurements presented in this paper demonstrate the superiority of Signia MaX when it comes to handling of speech in noise.

References

American National Standards Institute (2024). Methods for calculation of the speech intelligibility index (ANSI/ASA S3.5-1997 (R2024)). Acoustical Society of America.

Andersen, A. H., Santurette, S., Pedersen, M. S., Alickovic, E., Fiedler, L., Jensen, J., & Behrens, T. (2021). Creating Clarity in Noisy Environments by Using Deep Learning in Hearing Aids. *Seminars in Hearing*, 42(3), 260-281. doi: 10.1055/s-0041-1735134.

Ashkanichenarlogh, V., Folkeard, P., Scollie, S., Kühnel, V., & Parsa, V. (2025). Objective Evaluation of a Deep Learning-Based Noise Reduction Algorithm for Hearing Aids Under Diverse Fitting and Listening Conditions. *Trends in Hearing*, 29,

23312165251396644. doi: 10.1177/23312165251396644.

Bisgaard, N., Vlaming, M. S., & Dahlquist, M. (2010). Standard audiograms for the IEC 60118-15 measurement procedure. *Trends in Amplification*, 14(2), 113-120. doi: 10.1177/1084713810379609.

Fitzgerald, M. B., Athreya, V. M., Srour, M., Rejimon, J. P., Venkitakrishnan, S., Bhowmik, A. K., Jackler, R. K., Steenerson, K. K., & Fabry, D. A. (2025). Effectiveness of deep neural networks in hearing aids for improving signal-to-noise ratio, speech recognition, and listener preference in background noise. *Frontiers in Audiology and Otology*, 3, 1677482. doi: 10.3389/fauot.2025.1677482.

Folkeard, P., Jensen, N. S., Kamkar Parsi, H., Bilert, S., & Scollie, S. (2024). Hearing at the Mall: Multibeam Processing Technology Improves Hearing Group Conversations in a Real-World Food Court Environment. *American Journal of Audiology*, 33, 782-792. doi: 10.1044/2024_AJA-24-00027.

Hagerman, B., & Olofsson, Å. (2004). A method to measure the effect of noise reduction algorithms using simultaneous speech and noise. *Acta Acustica United with Acustica*, 90(2), 356-361.

Hernandez, A. R., Chalupper, J., & Powers, T. (2006). An Assessment of Everyday Noises and Their Annoyance. *Hearing Review*, 13(7), 16.

Holube, I., Fredelake, S., Vlaming, M., & Kollmeier, B. (2010). Development and analysis of an international speech test signal (ISTS). *International Journal of Audiology*, 49(12), 891-903. doi: 10.3109/14992027.2010.506889.

Jensen, N. S., Samra, B., Best, S., Wilson, C., & Taylor, B. (2025). Improving Speech Understanding in Noisy Group Conversation: 86% of Participants Performed Better with Signia Integrated Xperience Versus Key Competitor. *AudiologyOnline*, Article 29273. Retrieved from www.audiologyonline.com.

Jensen, N. S., Samra, B., Kamkar Parsi, H., Bilert, S., & Taylor, B. (2023). Power the conversation with Signia Integrated Xperience and RealTime Conversation Enhancement. *Signia White Paper*. Retrieved from www.signia-library.com.

Jensen, N. S., Samra, B., Taghvaei, N., & Taylor, B. (2024a). Improving the Real-World Conversation

Experience With a Multi-Stream Architecture.
Hearing Review, 31(9), 16-20.

Jensen, N. S., Wilson, C., Kamkar Parsi, H., Ramirez, I., & Vormann, M. (2026). Signia Multi-adaptive Xperience outperforms key competitors in group conversation in noise. Signia White Paper. Retrieved from www.signia-library.com.

Jensen, N. S., Wilson, C., Kamkar Parsi, H., Samra, B., Hain, J., Best, S., & Taylor, B. (2024b). Signia IX delivers more than twice the speech enhancement benefit in a noisy group conversation than the closest competitors. Signia White Paper. Retrieved from www.signia-library.com.

Korhonen, P., Slugocki, C., Kuk, F., & Peeters, H. (2026). Enhancing group conversations in noise with Signia's new RealTime Conversation Enhancement^{AI} powered by Acoustic IntelligenceTM. Signia White Paper. Retrieved from www.signia-library.com.

LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436-444. doi: 10.1038/nature14539.

Moore, B. C. J. (2026). Hearing Aids: What Works Well and What Can Be Improved. *Journal of the Association for Research in Otolaryngology*. doi: 10.1007/s10162-026-01031-5.

Picou, E. M. (2025). Are Hearing Aid Benefit and Satisfaction Rates Higher for Traditional or Over-the-Counter Hearing Aids? *Seminars in Hearing*, 46(3), 196-208. doi: 10.1055/s-0045-1812039.

Powers, T., Froehlich, M., Branda, E., & Weber, J. (2018). Clinical Study Shows Significant Benefit of Own Voice Processing. *Hearing Review*, 25(2), 30-34.

Samra, B., Jensen, N. S., & Lotter, T. (2025). Own Voice Processing 2.0 linked to 23% fewer returns. Signia Data Insights. Retrieved from www.signia-library.com.

Smeds, K., Wolters, F., & Rung, M. (2015). Estimation of signal-to-noise ratios in realistic sound scenarios. *Journal of the American Academy of Audiology*, 26(2), 183-196. doi: 10.3766/jaaa.26.2.7.

**Be
Brilliant™**