

## SEMI-VIRTUAL ASSESSMENT OF CAB SIGNALS IN TRAINS

Tim Heumann<sup>1</sup>, Domenico Tallarico<sup>2</sup>, Fabrice Aubin<sup>3</sup>, Mike Lower<sup>4</sup>

<sup>1</sup>Siemens Mobility, Duisburger Str. 145, Krefeld, Germany

<sup>2</sup>Stadler Rheintal AG, Neudorfstrasse 8, CH-9430 St. Margrethen, Switzerland

<sup>3</sup>SNCF Voyageurs, 4 allée des Gêmeaux 72100 Le Mans, France

<sup>4</sup>ISVR Consulting, University of Southampton, Highfield, Southampton, SO17 1BJ, UK

This paper is a revised version of the authors' submitted manuscript that was peer-reviewed by at least two anonymous reviewers.

### Abstract

Verifying the audibility of warning signals in train driver's cabs currently requires repeated physical measurements. According to TSI LOC&PAS [1], audible signals shall exceed the cab background noise level by 6 dB at maximum operating speed. In practice, regulatory guidance is often interpreted as requiring verification of numerous warning tones, speech messages, and volume-dependent signal variants, resulting in substantial testing effort.

This paper presents a semi-virtual verification procedure based on the measurement of source–receiver impulse responses in a driver's cab and subsequent convolution with digital cab signals. The method captures the complete audio reproduction chain, including loudspeakers and cabin acoustics, and allows signal levels to be evaluated and adjusted in post-processing without repeated physical measurements.

Validation measurements performed on a modern railway vehicle show close agreement between measured and calculated results. Under stable boundary conditions, deviations of the spatially averaged  $L_{p,AFmax}$  values were generally below 0.3 dB. The method proved particularly robust for broadband speech signals, while tonal signals showed greater sensitivity to changing environmental conditions.

This work supports ongoing standardization activities within CEN/TC 256 and demonstrates a practical approach for reducing effort during audibility assessment of cab signals.

**Keywords:** cab signals, semi-virtual verification, impulse response, ISO 3381, acoustics

### 1. INTRODUCTION

Audible driver's cab signals are a fundamental component of railway vehicle safety, operations, and driver ergonomics.

These signals range from high-priority warning tones and alarms to speech announcements intended to convey specific operational states to the train driver. Traditionally, verifying that these signals meet the required sound pressure levels necessitates repetitive play-and-measure testing inside the actual vehicle cab. This conventional approach requires setting up an array of eight microphones and manually playing and recording each individual audio file.

This testing burden is driven by a multi-layered regulatory framework designed to ensure that no critical information is masked by the operational noise of the train. The primary limit value is established by the Technical Specification for Interoperability for Locomotives and Passenger Rolling Stock (TSI LOC&PAS [1]), which mandates that audible signals must be at least 6 dB above the ambient noise level in the cab.

The specific details on how to perform this comparison are outlined in the RFU-RST-090 [2]. This document specifies that the maximum A-weighted sound pressure level of the audible signal with 'F' time weighting, denoted as  $L_{p,AFmax}$  must be compared to the background noise derived from the TSI NOISE [3] cab measurement, at least for the vehicle's maximum operational speed  $v_{max}$ . A physical measurement procedure for capturing both the signal level and the ambient noise is described in the current ISO 3381 [4] standard.

However, a significant challenge arises from what is missing in the current normative framework: a clear, univocal definition of which signals are relevant for this acoustic verification. RFU-RST-090 is unclear on this boundary, and regulatory bodies often interpret it to mean that all audible signals generated by on-board equipment must be measured and verified. Consequently, the number of audio files that must be physically measured during vehicle homologation has grown substantially.

The test burden grows in modern trains equipped with dynamic audio systems. RFU-RST-090 suggests the use of speed-dependent or background-noise-dependent signal levels to prevent signals from being excessively loud at standstill while remaining audible at high speeds. For speed-dependent signals, the number of required measurements rises dramatically, as each signal must be verified across

multiple speed-dependent volume steps. For background-noise-dependent signals, the situation is even more complex, as the regulatory measurement procedure for dynamically adapting systems is currently not fully defined, leading to significant engineering uncertainty and logistical bottlenecks.

To address this inefficiency, a semi-virtual verification procedure has been developed. This paper is based on preliminary work conducted within CEN TC256, which is mandated to write the first European standard for evaluating the audibility of warning signals in the driver's cab. The purpose of this paper is to define, document, and validate a new semi-virtual measurement procedure that can replace iterative direct measurements while adhering to the existing regulatory metrics. The general approach of the procedure is outlined in Fig. 1.

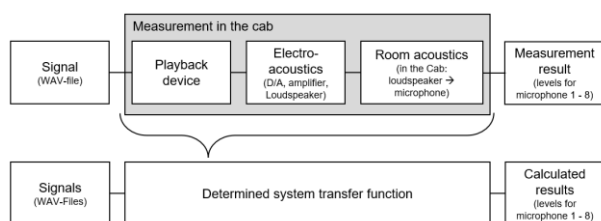


Figure 1. Semi virtual approach for cab signal level determination

## 2. THEORETICAL BACKGROUND AND SYSTEM APPROXIMATION

The fundamental principle underlying the proposed semi-virtual procedure is that the overall acoustic reproduction chain inside a train cab can be approximated as a Linear Time-Invariant (LTI) system. In the context of audible cab signals, this reproduction chain comprises the digital audio file, the digital-to-analogue converter (DAC), the audio amplifier, the physical loudspeakers, and the complex acoustic environment of the cab interior, up to the driver's ear positions.

For a system to be considered an LTI, it must satisfy linearity and time invariance. Linearity implies that the system's output is proportional to its input, and the response to a sum of inputs is the sum of their individual responses. In practical terms, this means that no components in the reproduction chain are driven into non-linear distortion, such as amplifier clipping or loudspeaker cone over-exursion. Time invariance implies that the system's characteristics do not change over time; provided that environmental boundary conditions remain constant.

Within the relevant frequency range for human hearing and typical cab signals (approximately 40 Hz to 20 kHz) and under stable operating conditions, the train cab's audio system behaves as a highly linear and time-invariant system. Therefore, the entire chain can be described mathematically by its impulse response  $h(t)$  in the time domain, or equivalently by its transfer function  $H(f)$  in the frequency domain.

The impulse response represents the acoustic output of the system when excited by a theoretical Dirac delta function. In a physical space like a train cab, the impulse response captures the direct sound travelling from the loudspeaker to the microphone, followed by a complex decay tail consisting of early reflections from the windows, console,

and ceiling, and the late reverberant field dictated by the cab's volume and the absorption coefficients of its interior materials.

If the impulse response  $h(t)$  of the cab is known, the acoustic output  $y(t)$  for any arbitrary input signal  $s(t)$  can be calculated using convolution [5] as:

$$y(t) = h(t) * s(t) = \int_{-\infty}^{+\infty} h(\tau)s(t - \tau)d\tau \quad (1)$$

Specifically, we use an exponentially swept sine as input signal [6]. By moving this process into the digital domain, the physical acoustic space is effectively virtualized, allowing for limitless testing and level adjustment of new input signals in post-processing without requiring access to the physical train.

## 3. METHODOLOGY OF THE SEMI-VIRTUAL PROCEDURE

### 3.1 Physical Measurement Setup

To ensure that the semi-virtual procedure aligns with existing regulatory frameworks and provides a direct alternative to traditional testing, the physical measurement arrangement strictly follows the approach applied for in-cab measurements defined in the current ISO 3381 standard.

The standard dictates that the sound field must be sampled spatially to account for local variations in sound pressure. Therefore, an array of eight measurement microphones is distributed evenly around the theoretical position of the driver's head. These microphones are positioned in a horizontal circle at the height of a seated driver's ears, at a radius of 25 centimetres from the centre point of the head (see Fig 2).

During the initial characterization measurement, the cab must be in a representative operational configuration. All doors and windows must be closed, and the interior geometry must be fully assembled, including seats, sun blinds, and console panels, as these elements significantly influence the acoustic reflections and absorption within the space. Furthermore, environmental boundary conditions, specifically the ambient temperature and relative humidity inside the cab, must be carefully controlled and meticulously documented. As will be discussed in Section 6, these boundary conditions are critical for the consistency of the method.

### 3.2 Excitation and Data Acquisition

In a traditional verification process, the engineer would iteratively play every single warning tone and speech signal at all relevant volumes. In the semi-virtual procedure, the system is excited only once per required loudspeaker routing configuration.

Rather than using a simple impulse (like a starter pistol or balloon pop) or broadband pink noise, we excite the driver's cab with an exponential sine sweep. An exponential sine sweep is a digitally generated signal whose frequency continuously increases in a given time interval, typically covering the entire relevant spectrum from 20 Hz to 20 kHz with a duration of approx. 5 s.



**Figure 2.** Measurement setup acc. to ISO 3381.

The use of an exponential sweep offers significant advantages for deriving acoustic transfer functions in noisy environments like railway vehicles. Since the energy of a sweep is spread over time rather than being concentrated in a short time interval or impulse, it avoids clipping the amplifiers and loudspeakers. The sweep is played through the exact operational playback path that will be used for the actual cab signals. For this purpose, the digital sweep file must be loaded into the train's audio control unit and routed through the onboard digital signal processor (DSP), amplifiers, and specific cab loudspeakers.

Care must be taken that any non-linear signal processing stages in the audio chain, such as dynamic range compression or automatic gain control, are either bypassed or set to a fixed operating point during the sweep measurement, as these would violate the LTI assumption. Routing the sweep through the operational path ensures that all system-specific routing, equalization curves, crossover networks, and limiter settings are inherently captured within the measurement.

The “dry” sweep and the response signal spectrograms are shown in Fig. 3 and 4, respectively. Following deconvolution of the measured response with the reference sweep in the frequency domain, the impulse response can be separated from uncorrelated background noise. This approach provides a high signal-to-noise ratio compared with traditional impulsive excitation methods.

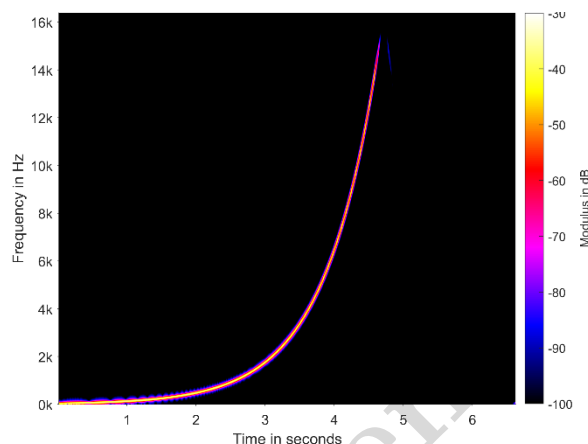
## 4. SIGNAL PROCESSING AND ALGORITHMIC IMPLEMENTATION

### 4.1 Deriving the Impulse Response

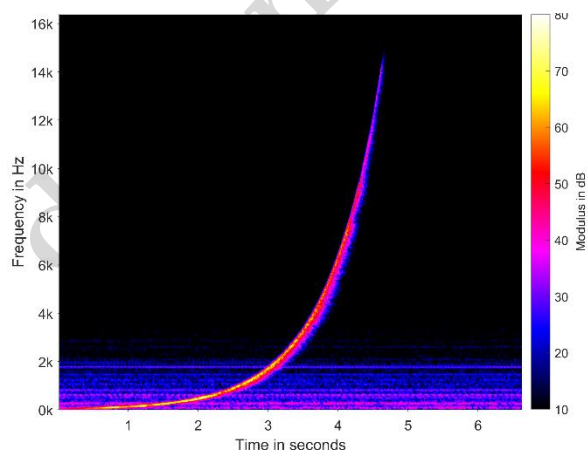
Once the physical sweep response is recorded, the data processing is executed within a dedicated computational script. The first step in the algorithmic workflow is to extract the vehicle-specific impulse response from the recorded sweep data.

For each of the eight microphone channels, the recorded response is cropped to the relevant time segment and precisely time-aligned to the original reference sweep. Fig. 5 compares the recorded sweep responses of two microphone positions. The spatial variation between the individual channels reflects the different source-receiver distances and local reflection patterns within the cab. If the

sample rates of measurement hardware and train's audio files differ, high-quality resampling is applied to ensure a common processing base.



**Figure 3.** Spectrogram of the “dry” sweep test signal



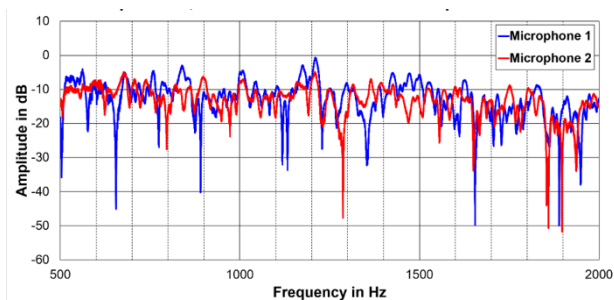
**Figure 4.** Spectrogram of the response signal measured by a single microphone

The transfer function is then calculated in the frequency domain. The complex frequency spectrum of the reference sweep  $X(f)$  and the complex frequency spectrum of the measured acoustic response  $Y(f)$  are computed using a Fast Fourier Transform (FFT). The system's transfer function  $H(f)$  is defined as

$$H(f) = \frac{Y(f)}{X(f)}. \quad (2)$$

The raw, vehicle-specific impulse response  $h(t)$  is subsequently obtained by applying an Inverse Fast Fourier Transform (IFFT) to the transfer function. To isolate the relevant acoustic characteristics and suppress late-arriving background noise or computational artifacts, a time-domain window is applied to the raw impulse response. The procedure utilizes a custom trapezoidal envelope designed to fully contain the initial delay, the direct sound impulse, and the critical early reflections within the cab. For the analysis described in this paper, a window that maintains unity gain for the initial portion of the impulse response, initiates a smooth fade-out at 0.23 seconds, and reaches zero amplitude at 0.33 s was used. This effective window

length of 330 ms has been carefully chosen to capture the complete modal decay and relevant reverberation at the driver's ear position in a typical enclosed train cab. By smoothly truncating the late reverberant tail, the window ensures that time-domain convolutions remain numerically stable and completely unaffected by accumulated low-level measurement noise. This window is carefully designed to retain the initial delay, the direct sound impulse, and the critical early reflections that define the acoustic character at the driver's ear position, while smoothly fading out the late reverberant tail and suppressing computational noise.



**Figure 5.** Frequency spectra of sweep measurements at two microphone positions, normalized to maximum of microphone 1.

## 4.2 Algorithmic Workflow and Post-Processing

After establishing the windowed impulse responses for all eight microphone positions, the physical cab is no longer required. The algorithmic workflow continues by iterating over any arbitrary digital cab signal.

Specifically, the convolution in Eq. (1) is performed for each of the eight microphone channels, thus simulating what each microphone would have recorded had the signal been physically played in the cab.

The workflow carries on evaluating the resulting sound pressure levels according to the regulatory requirements. A-weighting filters are applied to the calculated signals to simulate human hearing sensitivity. The sound pressure level time history is calculated using the time weighting 'F' (0.125 seconds), and the maximum value  $L_{p,AFmax}$  is extracted for each channel. Finally, the arithmetic mean across the eight microphone positions determines the aggregated result, as prescribed by [4] and [2].

Beyond simple evaluation, this algorithm provides a powerful benefit to the railway industry: the ability to adjust levels in post-processing. Indeed, the methodology permits the digital gain adjustments required for the signals to exceed the cab noise by 6 dB. This procedure can also be used to adjust the gain of new audio signals before they are installed, eliminating the trial-and-error level setting process on the actual train. The semi-virtual procedure addresses the signal side of the audibility assessment. The background noise level required for the +6 dB comparison is obtained from the existing TSI NOISE [3] cab measurement at the vehicle's maximum operational speed, as prescribed by RFU-RST-090

## 5. EXPERIMENTAL VALIDATION

To prove the validity, accuracy, and reliability of the semi-virtual procedure, a comprehensive measurement campaign was conducted. The physical cab of a representative modern railway vehicle, an electrical

multiple unit for commuter and regional operation, was equipped with the eight-microphone array according to [4]. The measurement setup can be seen in Fig. 2.

First, the transfer function was measured using the exponential sine sweep to capture the cab response at all eight microphone positions. Subsequently, a series of representative cab signals was physically played through the audio system and recorded. This test suite included five complex, broadband speech announcements in German and one simple, tonal alarm. Each signal was measured in six volume steps. The sweep measurement for the determination of the impulse response was also determined in six volume steps.

The physical measurement campaign spanned an entire day under sunny weather conditions with the vehicle's HVAC system deactivated. Consequently, although exact environmental boundary conditions were not continuously monitored, it is acknowledged that the ambient temperature and relative humidity inside the cab naturally fluctuated over the course of the testing period.

The digital files of the same signals were then processed through the semi-virtual convolution algorithm using the previously derived impulse responses. For each volume level, the corresponding sweep measurement was used. The physically measured sound pressure levels were then compared directly against the semi-virtually calculated levels to determine any differences and assess the method's accuracy.

## 6. RESULTS AND DISCUSSION

### 6.1 Consistency Under Constant Boundary Conditions

The most significant finding of the experimental validation is that the semi-virtual approach yields results that are highly consistent with the direct measurements, provided that the physical boundary conditions remain constant. When the impulse response measurement (via sweep) and the physical validation measurements of the actual signals are conducted in quick succession on the same day, in the same train, and under the same environmental conditions, the calculated  $L_{p,AFmax}$  values match the physically measured  $L_{p,AFmax}$  values with a deviation of less than 0.5 dB.

The results in Table 1 show the absolute values of the difference between the measurement and the convolved (calculated) signals. For brevity, Table 1 presents the results for a representative subset of the tested signals: the speech signal with the lowest deviation, the speech signal with the highest deviation, and the tonal alarm signal. These three cases span the full range of observed deviations and illustrate the method's performance boundaries. The complete dataset for all signals is available from the authors upon request. For better visibility, a colour shading is used to highlight higher deviations.

The use of one impulse response function per volume level improved the quality of the results in comparison to earlier tests, where only one sweep measurement was performed. Using this procedure, a deviation smaller than 0.2 dB for the average level was achieved in most cases. The signal 'Speech WORST' shows consistently elevated deviations of up to 1.4 dB at microphone 6 across multiple volume steps, suggesting a systematic cause such as a minor positional shift of the microphone

between measurements or a local acoustic interaction specific to the spectral content of this signal. The spatially averaged result remains within 0.3 dB, confirming the robustness of the eight-microphone averaging against localised anomalies.

**Table 1.** Difference between measurement and convolved signals with volume specific sweep measurement. Blue shading indicates underestimation by the semi-virtual method, whereas red shading indicates overestimation. Colour intensity is proportional to the absolute deviation.

| signal       |    | Difference (absolute value)<br>measurement – convolved signal<br>mic 1 to 8 + sum(1...8) in dB |     |     |     |     |     |     |     |        |
|--------------|----|------------------------------------------------------------------------------------------------|-----|-----|-----|-----|-----|-----|-----|--------|
|              |    | 1                                                                                              | 2   | 3   | 4   | 5   | 6   | 7   | 8   | (1..8) |
| Speech BEST  | 5  | 0.1                                                                                            | 0.0 | 0.0 | 0.0 | 0.1 | 0.4 | 0.3 | 0.1 | 0.1    |
|              | 6  | 0.0                                                                                            | 0.1 | 0.0 | 0.0 | 0.1 | 0.3 | 0.3 | 0.0 | 0.1    |
|              | 7  | 0.0                                                                                            | 0.1 | 0.0 | 0.0 | 0.0 | 0.2 | 0.2 | 0.0 | 0.1    |
|              | 8  | 0.0                                                                                            | 0.0 | 0.0 | 0.0 | 0.0 | 0.2 | 0.2 | 0.0 | 0.1    |
|              | 9  | 0.0                                                                                            | 0.0 | 0.0 | 0.1 | 0.0 | 0.2 | 0.2 | 0.0 | 0.1    |
| Speech WORST | 10 | 0.1                                                                                            | 0.2 | 0.1 | 0.0 | 0.0 | 0.2 | 0.2 | 0.1 | 0.0    |
|              | 5  | 0.2                                                                                            | 0.2 | 0.2 | 0.3 | 0.1 | 0.2 | 1.4 | 0.2 | 0.4    |
|              | 6  | 0.1                                                                                            | 0.2 | 0.1 | 0.1 | 0.0 | 0.2 | 1.4 | 0.1 | 0.3    |
|              | 7  | 0.0                                                                                            | 0.0 | 0.0 | 0.0 | 0.0 | 0.1 | 1.3 | 0.1 | 0.2    |
|              | 8  | 0.1                                                                                            | 0.1 | 0.1 | 0.1 | 0.1 | 0.0 | 1.2 | 0.2 | 0.1    |
| Tone         | 9  | 0.2                                                                                            | 0.3 | 0.1 | 0.2 | 0.1 | 0.1 | 1.3 | 0.1 | 0.0    |
|              | 10 | 0.4                                                                                            | 0.5 | 0.2 | 0.3 | 0.3 | 0.1 | 1.3 | 0.3 | 0.1    |
|              | 5  | 0.0                                                                                            | 0.0 | 0.1 | 0.0 | 0.1 | 0.1 | 0.2 | 0.0 | 0.1    |
|              | 6  | 0.1                                                                                            | 0.1 | 0.1 | 0.0 | 0.0 | 0.1 | 0.2 | 0.0 | 0.1    |
|              | 7  | 0.2                                                                                            | 0.2 | 0.0 | 0.1 | 0.1 | 0.2 | 0.2 | 0.0 | 0.1    |
|              | 8  | 0.0                                                                                            | 0.0 | 0.0 | 0.0 | 0.0 | 0.0 | 0.1 | 0.0 | 0.0    |
|              | 9  | 0.0                                                                                            | 0.0 | 0.0 | 0.0 | 0.0 | 0.0 | 0.0 | 0.0 | 0.0    |
|              | 10 | 1.7                                                                                            | 1.7 | 1.7 | 1.8 | 1.9 | 1.9 | 1.9 | 1.9 | 1.8    |

Although a dedicated sweep measurement was used for each volume setting, the observed improvement is not necessarily evidence of system non-linearity. A more plausible explanation is that the sweep and signal measurements were performed under more closely matched environmental conditions. The temporal proximity between measurements likely reduced the influence of variations in temperature, humidity, and other boundary conditions. This interpretation is supported by the results obtained using a single sweep measurement, where larger deviations were observed despite the system operating within its nominal linear range.

A notable exception is the tonal signal at volume step 10, for which a systematic deviation of approximately 1.8 dB was observed across all microphone positions. Unlike the localized deviation observed for microphone position 6 in the "Speech WORST" signal, this behaviour affects the complete microphone array and therefore points towards a source-related effect. Preliminary analysis suggests that loudspeaker distortion or saturation effects at this drive level may contribute to the discrepancy. Further investigation is required to establish the underlying cause. The results are shown in Table 2 in a similar manner as Table 1. Here, the deviations for the individual microphones rise significantly, while the overall deviation rises slightly to a maximum of 0.4 dB for the speech signals. The "Tone" signal is more affected, as the deviation of 1.5 dB shows. It should be noted that the deviation pattern attributed to changing boundary conditions is distinctly different from the deviation caused by the known loudspeaker distortion. To ensure the validity of the LTI assumption, it is recommended to verify the system's linearity prior to the sweep measurement, e.g., by comparing two sweep recordings at

different drive levels. A level-proportional response confirms linear operation, while deviations indicate amplifier clipping or loudspeaker over-exursion. A formal measurement uncertainty analysis according to GUM [7] has not been performed within the scope of this study and is subject to future work. Overall, within the linear operating range of the audio system, the LTI assumption holds true, proving that the mathematical convolution accurately represents the acoustic environment of the cab.

**Table 2.** Difference between measurement and convolved signals with a single sweep measurement. Blue shading indicates underestimation by the semi-virtual method, whereas red shading indicates overestimation. Colour intensity is proportional to the absolute deviation.

| signal       | volum | Difference (absolute value)<br>measurement – convolved signal<br>mic 1 to 8 + sum(1...8) in dB |     |     |     |     |     |     |     |        |
|--------------|-------|------------------------------------------------------------------------------------------------|-----|-----|-----|-----|-----|-----|-----|--------|
|              |       | 1                                                                                              | 2   | 3   | 4   | 5   | 6   | 7   | 8   | (1..8) |
| Speech BEST  | 5     | 0.1                                                                                            | 0.0 | 0.0 | 0.0 | 0.1 | 0.4 | 0.3 | 0.1 | 0.1    |
|              | 6     | 0.0                                                                                            | 0.1 | 0.0 | 0.0 | 0.1 | 0.3 | 0.3 | 0.0 | 0.1    |
|              | 7     | 0.0                                                                                            | 0.1 | 0.0 | 0.0 | 0.0 | 0.2 | 0.2 | 0.0 | 0.1    |
|              | 8     | 0.6                                                                                            | 0.4 | 0.6 | 0.0 | 0.6 | 0.0 | 0.4 | 0.3 | 0.2    |
|              | 9     | 0.5                                                                                            | 0.3 | 0.5 | 0.2 | 0.6 | 0.0 | 0.3 | 0.2 | 0.1    |
| Speech WORST | 10    | 0.1                                                                                            | 0.1 | 0.4 | 0.1 | 0.5 | 0.1 | 0.4 | 0.3 | 0.0    |
|              | 5     | 0.2                                                                                            | 0.2 | 0.2 | 0.3 | 0.1 | 0.2 | 1.4 | 0.2 | 0.4    |
|              | 6     | 0.1                                                                                            | 0.2 | 0.1 | 0.1 | 0.0 | 0.2 | 1.4 | 0.1 | 0.3    |
|              | 7     | 0.0                                                                                            | 0.0 | 0.0 | 0.0 | 0.0 | 0.1 | 1.3 | 0.1 | 0.2    |
|              | 8     | 0.3                                                                                            | 0.1 | 1.0 | 0.0 | 0.3 | 0.4 | 1.7 | 0.5 | 0.4    |
| Tone         | 9     | 0.1                                                                                            | 0.3 | 0.8 | 0.1 | 0.4 | 0.3 | 1.6 | 0.4 | 0.3    |
|              | 10    | 0.1                                                                                            | 0.8 | 0.7 | 0.5 | 0.5 | 0.1 | 1.5 | 0.1 | 0.1    |
|              | 5     | 0.0                                                                                            | 0.0 | 0.1 | 0.0 | 0.1 | 0.1 | 0.2 | 0.0 | 0.1    |
|              | 6     | 0.1                                                                                            | 0.1 | 0.1 | 0.0 | 0.0 | 0.1 | 0.2 | 0.0 | 0.1    |
|              | 7     | 0.2                                                                                            | 0.2 | 0.0 | 0.1 | 0.1 | 0.2 | 0.2 | 0.0 | 0.1    |
|              | 8     | 3.9                                                                                            | 2.5 | 3.4 | 1.1 | 0.3 | 1.9 | 1.9 | 3.6 | 1.5    |
|              | 9     | 3.5                                                                                            | 2.0 | 3.5 | 1.1 | 0.4 | 2.2 | 1.9 | 2.8 | 1.2    |
|              | 10    | 1.8                                                                                            | 0.1 | 0.3 | 2.7 | 2.5 | 4.2 | 0.5 | 0.9 | 0.8    |

## 6.2 Robustness of Speech Signals

When analysing the data further, a distinct difference in robustness emerges between different types of audio signals. The semi-virtual approach is highly robust for broadband signals, particularly human speech. For dynamic speech announcements, the deviation between the physical measurement and the semi-virtual calculation remains negligible even if minor environmental changes occur.

Because speech is a broadband signal characterized by rapidly changing frequencies and transient amplitudes, spatial variations and standing waves within the cab average out across the broader frequency spectrum. The complex frequency content of speech interacts with the cab's geometry in a broadband manner, leading to highly stable sound pressure levels across the eight microphone positions. Consequently, the convolution algorithm captures and predicts speech levels with excellent reliability.

## 6.3 Sensitivity of Tonal Signals to Boundary Condition Changes

Conversely, for signals containing prominent, continuous tones or multi-tonal chords—such as traditional warning beeps and system alarms—the procedure exhibits a higher sensitivity to changes in boundary conditions. While the method shows excellent agreement with measurements when boundary conditions are maintained, applying a transfer function measured under one set of conditions to a physical validation measurement conducted under different conditions reveals notable variances.

This sensitivity is not a flaw in the semi-virtual calculation, but rather a physical reality of acoustics. Tonal signals are inherently more sensitive to changes in the physical testing environment than speech signals. The sensitivity is driven by several factors:

First, boundary conditions such as temperature and humidity drastically affect the sound field. In a relatively small, enclosed space like a driver's cab, pure continuous tones can cause standing waves and resonances. The spatial distribution of acoustic nodes (areas of minimum pressure) and antinodes (areas of maximum pressure) often matches the wavelength of a tone. Since the speed of sound depends directly on air temperature and humidity, a tiny shift in the cab's climate alters the speed of sound thus shifts the spatial locations of the nodes and antinodes within the cab compared to the fixed microphone positions.

Second, measurement setup variations introduce significant variance. Even small changes in the positioning of the eight microphones between different measurement campaigns will alter the recorded sound pressure of high-frequency tones, as the microphones may move into or out of local acoustic nodes.

Third, geometric changes due to manufacturing tolerances play a role. The high-frequency acoustic response varies slightly from vehicle to vehicle within the same fleet due to minor differences in the stiffness of interior panelling or the exact torque applied to speaker mountings.

Therefore, if a physical play-and-measure test is repeated on a different day with a slightly different ambient temperature, or on a different train of the same type, the physical sound field inside the cab can change. The semi-virtual method captures the acoustic state of the cab at the time of the sweep measurement. The observed variances simply highlight that tonal signals are highly sensitive to environmental shifts.

Furthermore, preliminary data suggests that there is a further frequency-dependent sensitivity regarding these boundary condition changes. Higher frequency tones appear to exhibit different spatial variance patterns than lower frequency tones when temperature or geometry changes. The exact nature of this frequency-dependent sensitivity, and how it interacts with the ISO 3381 microphone array dimensions, requires further dedicated investigation in future studies to determine constraints of the boundary conditions for the sweep measurement. The sensitivity of the semi-virtual procedure to environmental changes suggests that future normative methods should define controlled boundary conditions for the measurement of the impulse response to minimize sensitivity.

## 7. CONCLUSION

The proposed semi-virtual verification procedure offers an innovative, scientifically sound, and highly efficient alternative to traditional physical measurements for audible driver's cab signals. Driven by the +6 dB requirements of TSI LOC&PAS and the expansive scope of RFU-RST-090, the railway industry requires new testing methods to manage the continuing increase in required acoustic verification, particularly for speed-dependent and background-noise-dependent audio systems.

By utilizing a computational transfer-function method based on a single sweep measurement, acousticians can accurately calculate in-cab sound pressure levels for an unlimited number of audio files. This allows for the predictive adjustment of new signals in post-processing before physical implementation, eliminating costly trial-and-error track time.

The method yields consistent results with direct measurements, especially when boundary conditions are constant. While broadband speech signals prove to be highly robust against environmental shifts, tonal signals exhibit a higher sensitivity to changes in temperature, humidity, microphone placement, and manufacturing tolerances. This sensitivity highlights the physical realities of standing waves in enclosed spaces rather than a deficit of the calculation method. Standardizing this semi-virtual approach, as currently explored within CEN TC256, will significantly streamline measurement efforts during initial vehicle homologation and lifecycle software updates, providing a more consistent and efficient path to acoustic compliance.

## 8. Acknowledgments

The authors would like to express their sincere gratitude to Jonas Tumbrägel for his foundational work in defining the semi-virtual verification procedure presented in this paper. Furthermore, special thanks are extended to Nicolai Stenzel, David Glomsda, and Christian Dreier for their invaluable contributions, technical insights, and continuous support in further developing and refining the methodology and its algorithmic implementation.

## REFERENCES

- [1] European Commission, "Commission Regulation (EU) No 1302/2014 concerning a technical specification for interoperability relating to the 'rolling stock — locomotives and passenger rolling stock' subsystem of the rail system in the European Union (TSI LOC&PAS).
- [2] NB-Rail Coordination Group, "RFU-RST-090 Issue 04: Audible information in cab," 2025.
- [3] European Commission, "Commission Regulation (EU) No 1304/2014 on the technical specification for interoperability relating to the subsystem 'rolling stock — noise' (TSI NOISE)," *Official Journal of the European Union*, 2014. (as amended)
- [4] CEN ISO 3381:2023, "Railway applications - Acoustics - Measurement of noise inside railbound vehicles," International Organization for Standardization, 2023.
- [5] Hespanha, J.P. *Linear System Theory*. Princeton university press, 2009.
- [6] Farina, A. "Simultaneous measurement of impulse response and distortion with a swept-sine technique", AES Convention, 2000.
- [7] JCGM 100:2008, "Evaluation of measurement data Guide to the expression of uncertainty in measurement (GUM)," Joint Committee for Guides in Metrology, 2008.