



Potential increase in classroom voice support from focusing retroreflectors

Chang Hyok Cha (1), Densil Cabrera (1) and Shuai Lu (1,2)

(1) School of Architecture, Design and Planning, University of Sydney, Sydney, Australia

(2) Shenzhen International Graduate School, Tsinghua University, Shenzhen, China

ABSTRACT

Focusing retroreflectors are capable of providing a high level of voice support (ST_V) for a nearby talker in the high frequency range, particularly at 2 kHz and above. The purpose of this study is to evaluate the potential use of focusing retroreflectors to enhance voice support for teachers in classrooms, thereby improving acoustic conditions for speaking with the potential to reduce vocal strain. Pelegrín-García's prediction model for the average ST_V in a room was extended to include the contribution of special reflectors to enumerate the potential of a focusing retroreflector installation to improve voice support for a talker. A prediction model with retroreflection shows that ST_V values substantially increase across different room volumes with various reverberation times. Based on the prediction model, two types of actual rooms (small and large university classrooms) were selected. Oral-binaural room impulse responses (OBIRs) were measured, to which retroreflective energy obtained from a focussing retroreflector was added through signal processing in order to examine the voice support achievable in the rooms under the assumption that retroreflectors are installed. The results show that the addition of retroreflective energy increased ST_V values in both rooms, with a particularly greater enhancement in the level of voice support observed in the larger and more absorptive room. The study indicates that focusing retroreflectors could serve as an effective architectural treatment to improve voice support for teachers in classrooms in which reverberant support is insufficient.

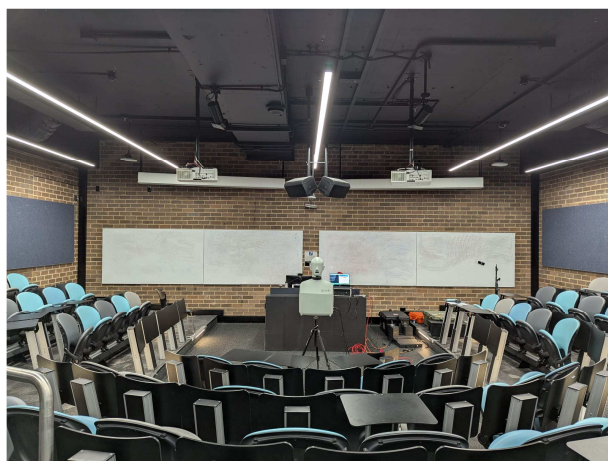
1 INTRODUCTION

Acoustics is important in classroom environments. Poor acoustics can hinder effective communication (Berg et al., 1996; Gheller et al., 2019), leading talkers to raise their voice levels against background noise (Lombard, 1911), causing vocal strain (Sierra-Polanco et al. 2021) and negative educational performance (Mealings, 2021). Modern classroom acoustical design addresses these issues by deploying extensive sound absorption to reduce background noise and reverberation time (Choi, 2014). A very short reverberation time is considered by some to be ideal (between 0.1 and 0.3 s), while 0.4 s is generally recommended as a practical design (Bistafa and Bradley, 2000). However, this approach may result in insufficient voice support for talkers due to weak reflections. Classroom acoustics considerably impacts teachers' voices, particularly talkers with voice problems appear more sensitive to acoustic conditions (Pelegrín-García et al., 2010). For this reason, a reverberation time between 0.45 and 0.6 s in a fully occupied classroom (or between 0.6 and 0.7 s in an unoccupied but furnished condition) was recommended to support both vocal comfort for teachers and speech intelligibility for students in flexible classroom acoustic settings (Pelegrín-García et al., 2014). However, longer reverberation time may increase background noise and reduce intelligibility, raising the question of whether voice support could be provided in another way. This study considers the possibility of providing voice support through focusing retroreflectors. Prior studies on retroreflector arrays in architectural environments have mainly reported substantial sound concentration on the source in octave bands above 2 kHz (Cabrera et al., 2018, 2021, 2022; Rapp, 2022). It was also observed that this high frequency bias became potentially stronger with focusing (Cabrera et al., 2023; Lu et al., 2026). Cansu et al. (2025) reported that reflective ceilings provided greater comfort to talkers compared with baseline and absorptive treatments. In particular, the results indicate that a decrease of 1.3 dB was observed in mean vocal intensity when participants performed the task in the diffuser configuration. Therefore, this study aims to explore the potential of applying focusing retroreflectors installed above the talker's head in classrooms to enhance vocal comfort. In addition, it extends Pelegrín-García's voice support (ST_V) prediction model by incorporating retroreflection and evaluates its performance against analysed and predicted values.

2 METHODS

2.1 University classrooms

Two classrooms at The University of Sydney were selected for this study. Room A (Fig 1a) is a small lecture room with a seating capacity of 100 and a volume of approximately 383 m³. It has a reverberation time of 0.8 s at mid frequencies (500 Hz and 1 kHz mean). Room B (Fig 1b) is a medium-sized open-plan classroom with a volume of approximately 630 m³ and a mid-frequency reverberation time of 0.45 s. A Brüel & Kjær 5128C head and torso simulator (HATS) was used to measure oral-binaural room impulse responses (OBRIR). The HATS was positioned with ear-height of 1.5 m at three typical didactic locations in each room. To minimise the influence of early reflections, it was placed at least 1.5 m away from the walls and other obstacles.



(a)

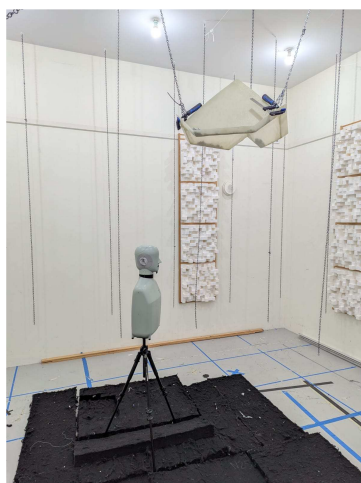


(b)

Figure 1: University classrooms with HATS (a) Room A; (b) Room B

2.2 Retroreflection

The retroreflected sound energy from a reflector prototype was obtained from the HATS in the reverberant chamber (130 m³) of the Acoustics Laboratory at the University of Sydney (Fig 2), although the reverberant part of the OBRIRs was discarded. The HATS was at the same height as it was in the classrooms. Floor reflections were attenuated by applying 100 mm of polyester wool sound-absorption treatment only to the floor area beneath the HATS. Four configurations of the focusing retroreflector were applied to examine the effect of their placement—two installation angles (60 and 90° elevation on the median plane) and two heights (2.0 and 2.5 m from the floor). The focusing retroreflector used in this study was a single parabolic trihedron fabricated from fibreglass, with an edge length of 0.5 m. The geometry and acoustic performance of a parabolic trihedron are described by Cabrera et al. (2023).



(a)



(b)

Figure 2: Configurations of HATS and the focusing retroreflector in the reverberation chamber: (a) placed 2.5 m above the floor and at 60° elevation; (b) placed 2.0 m above the floor and at 90° elevation

2.3 Signal synthesis

A synthesised OBRIR with retroreflections, as if the focusing retroreflector had been installed in the classroom, was generated by employing a signal substitution method combining the OBRIRs in the original classrooms and the retroreflection from the laboratory measurements. The waveform arriving within 3 to 13 ms was regarded as retroreflection, based on the distance between the mouth and the focusing retroreflector. Figure 3a shows the original OBRIR of Room B before signal synthesis. Figure 3b presents the synthesised signal, in which the retroreflection measured with the focusing retroreflector positioned 2.5 m above the floor and at 60° elevation was substituted into the corresponding time window (3 to 13 ms). In this example, two discrete reflections of Room B were overwritten by the retroreflector's impulse response.

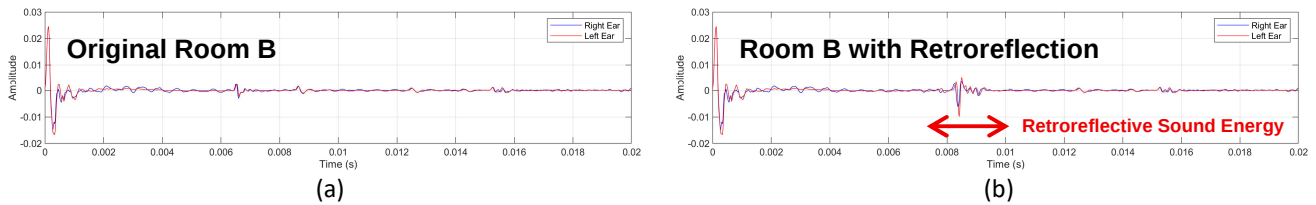


Figure 3: (a) Oral-binaural room impulse response (OBRIR) of Room B, and (b) the OBRIR after the laboratory-measured retroreflection was inserted.

Pelegrín-García (2011) simplified the measurement method of voice support initially proposed by Brunskog et al. (2009). The direct sound of the original room without the focusing retroreflector $h_{o,d}(t)$ is obtained by applying a window $w(t)$ to the measured OBRIR $h(t)$.

$$h_{o,d}(t) = h(t) \times w(t), \quad (1)$$

where

$$w(t) = \begin{cases} 1 & t < 4.5 \text{ ms} \\ 0.5 + 0.5 \cos [2\pi(t - t_0)/T_W] & 4.5 \text{ ms} < t < 5.5 \text{ ms} \\ 0 & t > 5.5 \text{ ms} \end{cases} \quad (2)$$

with $t_0 = 4.5$ ms as a time-shift parameter and $T_W = 2$ ms as the period of the cosine function. The reflected sound of the original room without the focusing retroreflector $h_{o,r}(t)$ is the complementary signal.

$$h_{o,r}(t) = h(t) \times (1 - w(t)) = h(t) - h_{o,d}(t). \quad (3)$$

The energy levels corresponding to the direct sound $L_{E,o,d}$ and the reflected sound $L_{E,o,r}$ in the original rooms are calculated as

$$L_{E,o,d} = 10 \log \frac{\int_0^\infty h_{o,d}^2(t) dt}{E_0} \text{ (dB)}, \quad (4)$$

$$L_{E,o,r} = 10 \log \frac{\int_0^\infty h_{o,r}^2(t) dt}{E_0} \text{ (dB)}. \quad (5)$$

The voice support of the original room $ST_{V,o}$ is defined as the difference between the reflected sound and the direct sound energy levels (i.e., the energy ratio expressed in decibels).

$$ST_{V,o} = L_{E,o,r} - L_{E,o,d} \text{ (dB)}. \quad (6)$$

However, in the case of the voice support from the OBRIRs with synthesised retroreflection, the retroreflected sound energy arrives sometimes earlier than 5 ms. The conventional method of calculating ST_V incorporates the initial retroreflected energy within the direct sound. Consequently, this approach was not appropriate for ST_V with early-arriving retroreflection. Therefore, this study adopted the method of determining the room gain (G_{RG}) as defined by Brunskog et al. (2008) and subsequently converting this value into voice support.

Room gain requires the measurement of two OBRIRs, one at the room of interest $h(t)$ and another in a reflection-free environment (such as an anechoic chamber, $h_{ach}(t)$). The direct sound in the original room without the focusing retroreflector was assumed to be equivalent to the anechoic condition, and this value was used as the

OBRIR of anechoic reference. In the following, L_E is the total energy level in the room including retroreflection. $L_{E,ach}$ is the energy level in anechoic-equivalent conditions.

$$L_E = 10 \log \frac{\int_0^\infty h^2(t) dt}{E_0} \text{ (dB)}, \quad (7)$$

$$L_{E,ach} = 10 \log \frac{\int_0^\infty h_{ach}^2(t) dt}{E_0} \text{ (dB)} \approx \text{Eq. (4)}, \quad (8)$$

$$G_{RG} = L_E - L_{E,ach} \text{ (dB)}. \quad (9)$$

ST_V with focusing retroreflection could be obtained from G_{RG} under the assumption that the total energy is approximately the sum of the energies corresponding to the direct and the reflected sound after windowing.

$$ST_{V,retro} \approx 10 \log \left(10^{\frac{G_{RG}}{10}} - 1 \right) \text{ (dB)}. \quad (10)$$

2.4 Proposed ST_V prediction model

Pelegrín-García (2012) proposed a prediction model for voice support (Eq. (11)) that accounts for the relationship between direct and reflected sound at the ears when the sound is emitted from the mouth. The prediction model estimates the position-averaged voice support in a room. The voice support (ST_V) in each octave band (oct) is

$$ST_{V,oct} = 10 \log \left[\left(\frac{cT_{oct}}{6 \ln(10)V} - \frac{4}{S} + \frac{Q_{oct}^*}{4\pi(2d)^2} \right) \cdot S_{ref} \right] + \Delta L_{HRTF,oct} - K_{oct} \text{ (dB)}. \quad (11)$$

In this equation, c denotes the speed of sound in air (≈ 343 m/s); S is the total surface area of the room (m^2); T is the reverberation time (s); V is the room volume (m^3); Q^* is the directivity of the source in the downward direction; d is the distance from the mouth to the floor ($= 1.5$ m); S_{ref} is the reference area (≈ 1 m^2). ΔL_{HRTF} is the diffuse-field head-related transfer function (dB). To incorporate the effect of retroreflection into the original prediction model, the retroreflected sound level relative to the direct sound level is calculated as follows:

$$L_{E,retro} - L_{E,ach} = 10 \log \frac{E_{retro}}{E_{ach}} = 10 \log \frac{\int_0^{13 \text{ ms}} h^2(t) dt}{\int_0^\infty h_{ach}^2(t) dt}. \quad (12)$$

To account for the contribution of retroreflection, the retroreflected energy relative to the direct sound is not inserted into the reflection-related term of the original prediction model (Eq. (11)). Instead, an extended prediction model ($ST_{V,retro,oct}$) is obtained by directly summing the energy of the theoretical ST_V without retroreflection (from Eq. (11)) and the measured retroreflected energy relative to the direct sound (Eq. (12)). Hence, Eq. (13) presents a hybrid prediction of octave band voice support incorporating theoretical support from the general room acoustics and the measured support from a distinct reflector.

$$ST_{V,retro,oct} = 10 \log \left\{ 10^{\left(\frac{ST_{V,oct}}{10}\right)} + 10^{\left(\frac{L_{E,retro,oct} - L_{E,ach,oct}}{10}\right)} \right\} \text{ (dB)}. \quad (13)$$

The energy level values in rows 8 to 11 of Table 1 are for the second term of Equation 13 based on the described measurement results. Table 1 provides all of the values required to evaluate Eq. (13).

Table 1: Relevant frequency-dependent quantities used in the prediction model of voice support

Category	Value					
Centre frequency (Hz)	125	250	500	1000	2000	4000
(1) Typical speech SPL on-axis at 1 m $L_{D,ref,1m,oct}$ (dB)	44.9	57.3	61.8	58.2	53.7	48.9
(2) Difference from SPL at eardrum (measured) $L_{D,oct} - L_{D,1m,oct}$ (dB)	13.1	11.8	11.7	13.5	15.3	14.1
(3) Typical speech levels at the eardrum = (1) + (2) $L_{D,ref,oct}$ (dB)	58.0	69.1	73.5	71.7	69.0	63.0
(4) Difference between on-axis SPL at 1 m and L_W $L_{D,1m,oct} - L_{W,oct}$ (dB)	-9.5	-8.1	-9.2	-9.5	-7.0	-6.0
(5) Constant K for model Eq. = (2) + (4) k_{oct} (dB)	3.6	3.7	2.5	4.0	8.3	8.1

Category	Value					
(6) Directivity of human speech on downward direction Q_{oct}^*	0.95	0.78	0.79	0.60	0.21	0.25
(7) Diffuse field HRTF $\Delta L_{\text{HRTF},\text{oct}}$ (dB)	0	0	2	4	11	13
(8) Retroreflection Energy related to direct sound energy - 2.5 m (Height), 90° (Angle) $L_{E,\text{retro},\text{RR}2.5,90^\circ} - L_{E,d}$ (dB)	-13.9	-19.3	-18.1	-13.8	-6.6	-2.9
(9) Retroreflection Energy related to direct sound energy - 2.5 m (Height), 60° (Angle) $L_{E,\text{retro},\text{RR}2.5,60^\circ} - L_{E,d}$ (dB)	-15.0	-20.9	-16.9	-11.4	-7.8	-2.2
(10) Retroreflection Energy related to direct sound energy - 2.0 m (Height), 90° (Angle) $L_{E,\text{retro},\text{RR}2.0,90^\circ} - L_{E,d}$ (dB)	-11.0	-23.5	-25.6	-20.4	-14.2	-7.8
(11) Retroreflection Energy related to direct sound energy - 2.0 m (Height), 60° (Angle) $L_{E,\text{retro},\text{RR}2.0,60^\circ} - L_{E,d}$ (dB)	-10.1	-22.0	-25.3	-17.3	-14.1	-13.1

3 RESULTS AND DISCUSSION

3.1 Analysis voice support from signal synthesis

Figure 4 shows the ST_V values in octave bands for Room A and B. The black line indicates the actual measured values in the original rooms without the focusing retroreflector, while the coloured lines represent the results analysed through signal synthesis under each combination of reflector position and angle. Overall, an increase in ST_V values was observed across all frequency bands in both rooms following the substitution of retroreflection signals in almost all cases. The results show that the ST_V was strengthened when the retroreflector was positioned closer to the HATS, except at 250 Hz. Voice support increases from the focusing retroreflector were sometimes strongest at higher frequencies (2 and 4 kHz). This may be attributed to the fact that the focusing retroreflector used in the previous study exhibited stronger high-frequency bias due to the spectral characteristics of diffraction and focusing. This is potentially beneficial, as discussed by Rapp et al. (2021), who found that voice regulation in conversations is most influenced by high-frequency voice support. Therefore, the enhancement of voice support at high frequencies through the application of a focusing retroreflector has the potential to be effective in classroom environments for teachers.

The spatial characteristics of the two classrooms also affected the results. Room B, which has a larger volume and shorter reverberation time, reveals a greater increase in overall speech-weighted ST_V from retroreflection signal synthesis than that of Room A. Because of its originally lower ST_V , Room B had more to gain from the retroreflective treatment. While Room A showed an overall speech-weighted ST_V enhancement of 1 to 2.1 dB with the retroreflection, Room B exhibited a larger gain of 3.5 to 6.3 dB. Although the ST_V of Room B without a retroreflector was 3.2 dB lower than that of Room A, the difference was reduced to less than 0.4 dB under all synthesised retroreflection conditions.

The overall speech-weighted ST_V has a strong frequency emphasis on 500 Hz when calculated from the values of each octave band (Category 3 in Table 1). However, since the focusing retroreflector provided greater voice support enhancement at the high frequencies (2 and 4 kHz), the overall speech-weighted ST_V had a more modest increase.

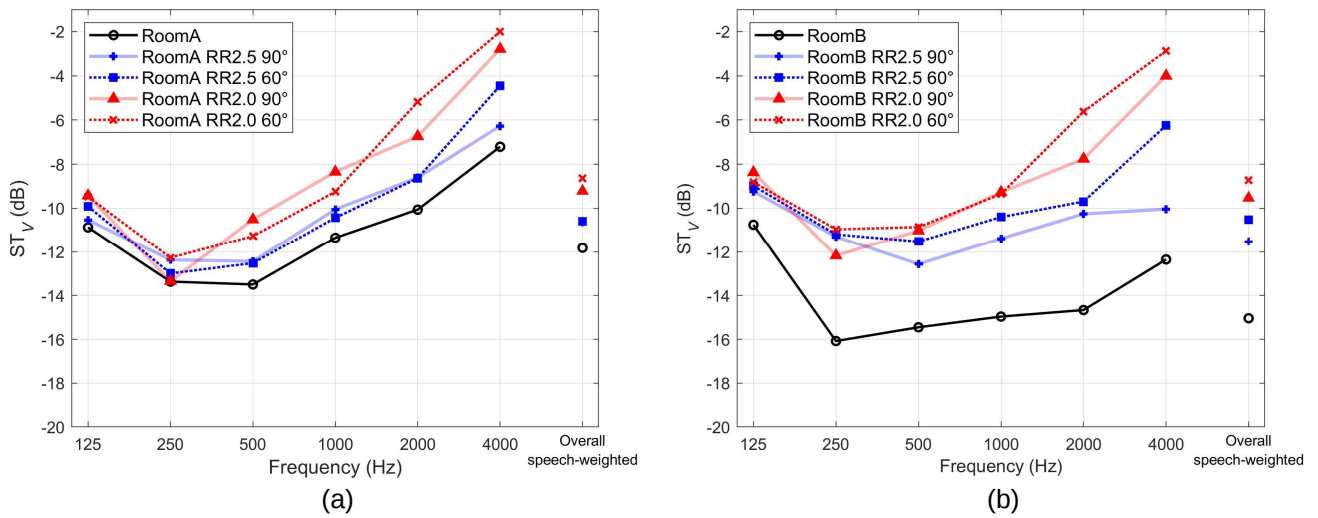


Figure 4: octave-band ST_V values obtained in (a) Room A and those in (b) Room B, including measurements from the original rooms (Black line) and from synthesized signals under each retroreflector condition. The overall speech-weighted ST_V is also dotted for each case.

The angle of the retroreflector positioned along the median vertical axis influences the results, and this is likely due to the directivity of human speech. Figure 5 shows the directional distribution of human speech in median vertical plane per octave band, based on Chu and Warnock (2022). The sound energy emitted from the mouth exhibits strong directivity in the downward direction. However, apart from this angle, considerable energy is also radiated around 60° at 1, 2, and 4 kHz. This is particularly evident in the 2 kHz band. This directivity gain could be a contributing factor to the observed increase in voice support at high frequencies.

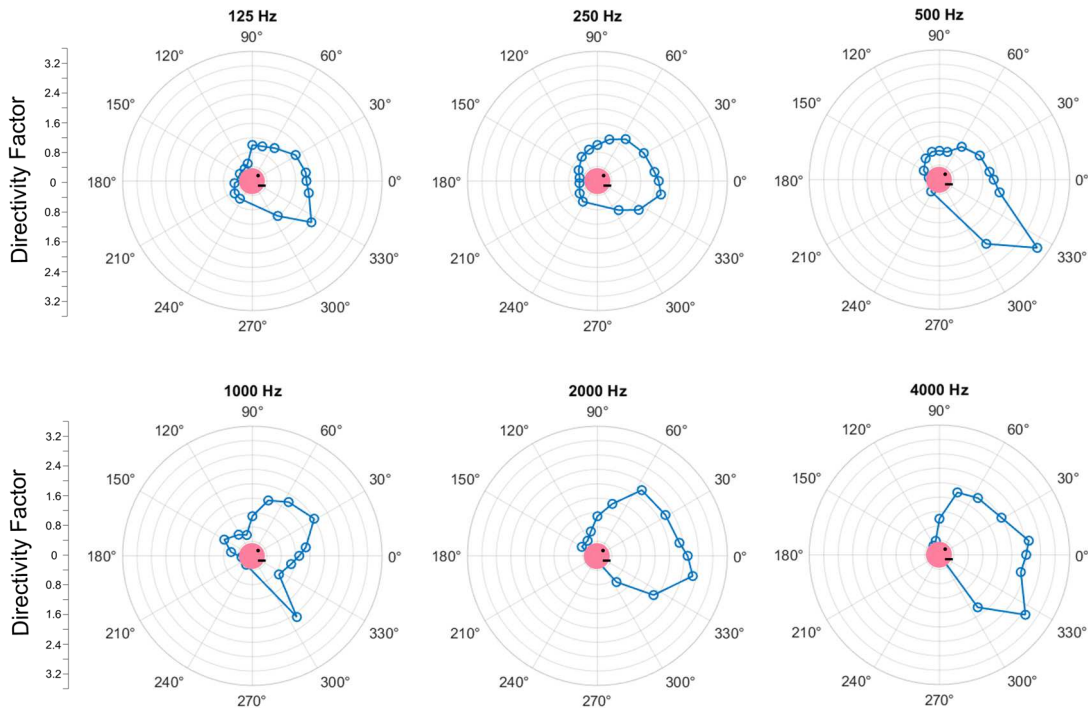


Figure 5: Long-term directivity distribution of human speech in the median vertical plane calculated from Chu and Warnock (2022). Values are directivity factor (Q), i.e., the energy ratio of sound radiated in the specified direction to that radiated to all directions.

3.2 Evaluation of a prediction model

Figure 6 presents a comparison between the ST_V values analysed through signal synthesis and those from the extended prediction model introduced above. The detailed room configurations for each case are presented in Table 2. cases 2 to 5 and 7 to 10 were calculated using the extended prediction model, whereas case 1 and 6 could be regarded as predicted values from the original model proposed by Pelegrín-García. To analyse the correlation between two values, regression lines, the coefficient of determination (R^2), bias (σ_T), residual standard deviation (σ_ϵ), and p -value (p) were used.

For 125, 250, 500 and 1000 Hz bands, the analysed ST_V values exceeded the predicted values (except cases 2 and 3 at 125 Hz), indicating poor or weak predictions. The corresponding bias values were -2.39, -3.52, -2.37 and -2.49 dB, and the residual standard deviations were 2.59, 1.7, 1.66, and 1.31, respectively. Particularly, the prediction for 250 Hz yielded the lowest R^2 value of 0.07. Conversely, at 2 and 4 kHz, the prediction model closely matched the analysed values. The slopes of the regression lines at 2 and 4 kHz were 1.1 and 1.2, respectively, and 2 kHz showed the strongest correlation ($R^2 = 0.93$). Bias values of less than 1 dB were observed in both bands. In particular, the residual deviation at 2 kHz was the lowest ($\sigma_\epsilon = 0.87$), indicating a reasonable degree of confidence in prediction.

The comparison predicted versus analysed overall speech-weighted values of voice support is shown in Figure 7. Despite poor agreement at low frequency and strong correction factor at specific frequency (500 Hz) in the speech-weighted calculation, fair predictive values were obtained, with the bias of -1.04 dB and residual deviation of 1.24 dB.

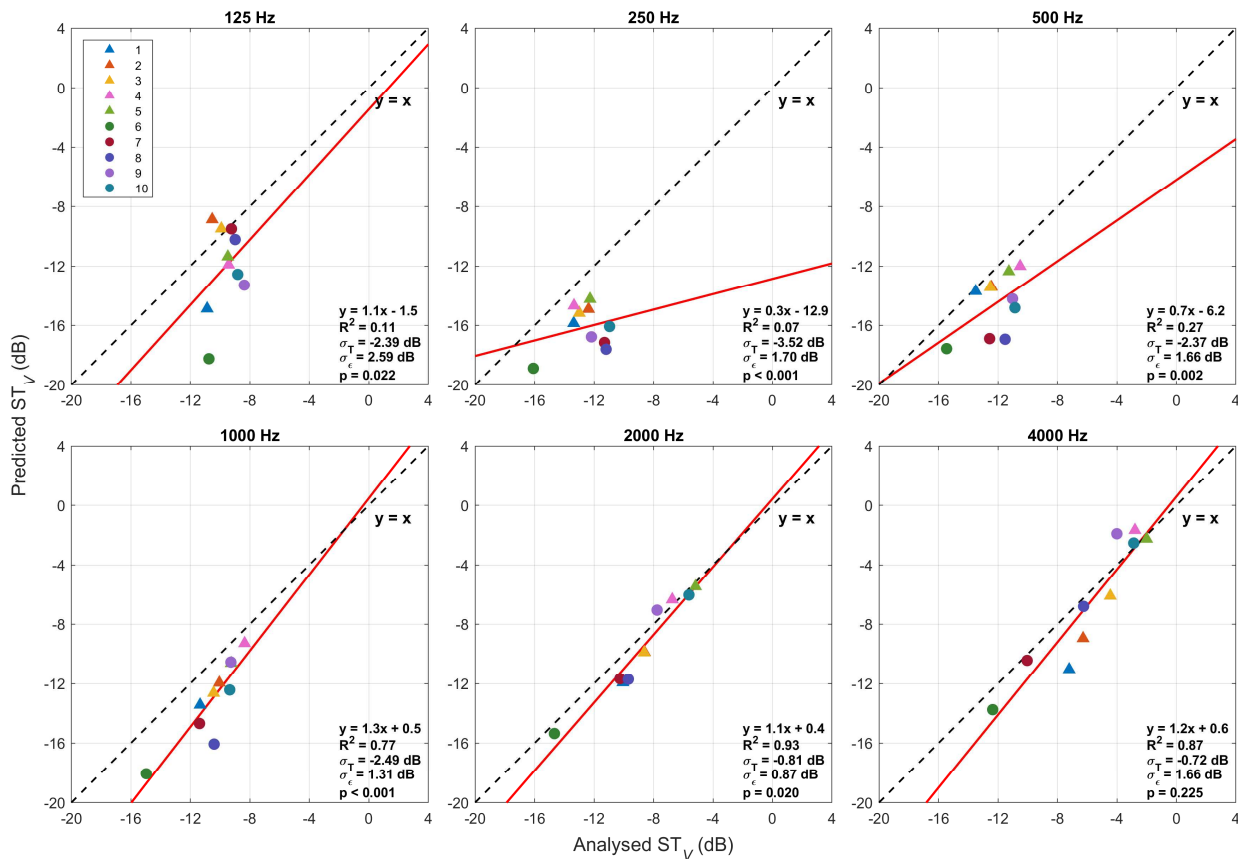


Figure 6: Scatter plot between predicted and analysed ST_V values from signal synthesis in frequency bands. The dashed lines represent perfect agreement between the two values. The solid red lines indicate the linear regression line for prediction. The values are represented by triangles for Room A and by circles for Room B.

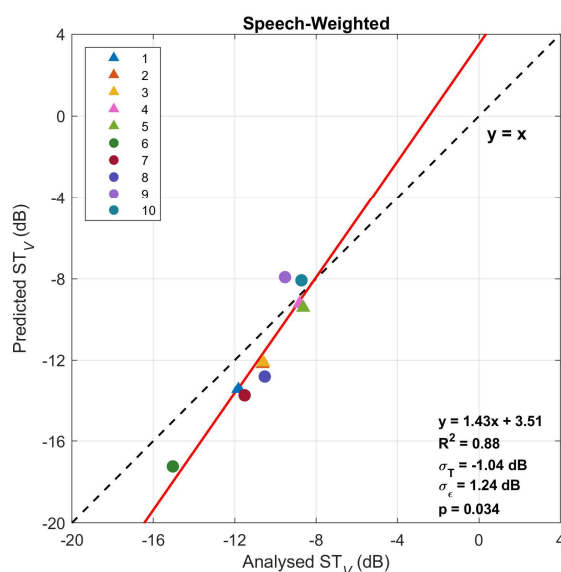


Figure 7: Scatter plot between predicted and analysed overall speech-weighted ST_V values from signal synthesis. The dashed line represents perfect agreement between the two values. The solid red line indicates the linear regression line for prediction. The values are represented by triangles for Room A and by circles for Room B.

Table 2: Room configuration details by type: retroreflector height and median vertical angle

Number	Room	Height (m)	Angle (°)
1	A	-	-
2	A	2.5	90
3	A	2.5	60
4	A	2.0	90
5	A	2.0	60
6	B	-	-
7	B	2.5	90
8	B	2.5	60
9	B	2.0	90
10	B	2.0	60

Figure 8a shows overall speech-weighted ST_V predicted values by Pelegrín-García's an original model, which was plotted against room volume and reverberation time. Figure 8b shows those predicted by the extended model that takes focusing retroreflection into account. When the measured ST_V of Room A ($T = 0.8$ s) and B ($T = 0.45$ s) are compared with the predicted values from Pelegrín-García's original model (Fig 7a), it is observed that the value of Room A is similar to the original prediction model, whereas the prediction for Room B is likely to be slightly underestimated. On the other hand, the analysed values of both Room A and B are higher than those of the extended prediction model. This is likely due to deviations in the predicted values at 125, 250, and 500 Hz, as discussed above.

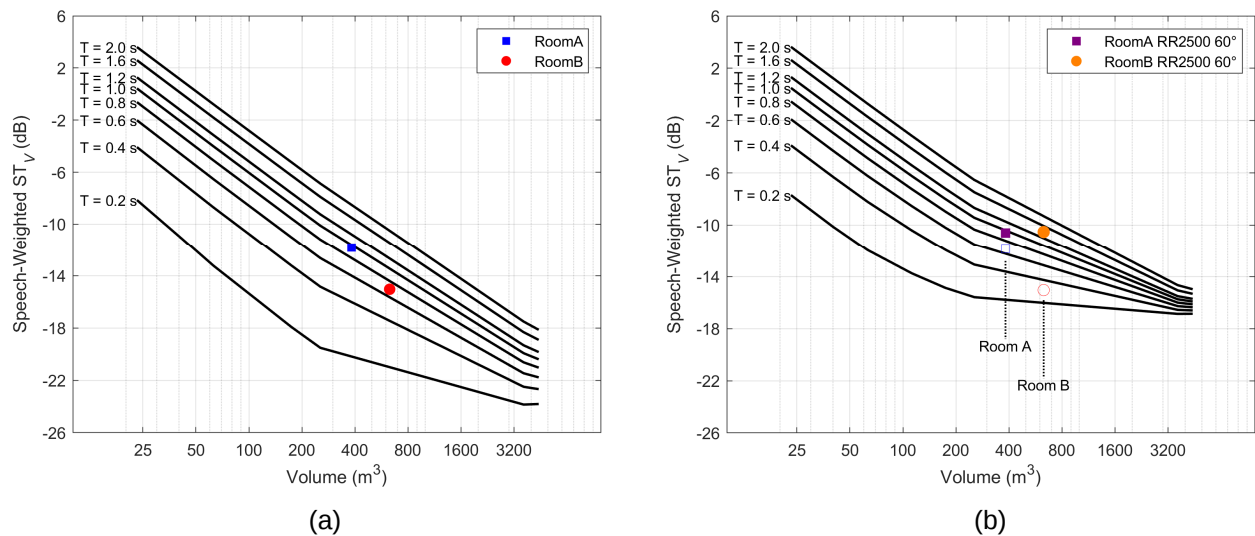


Figure 8: Overall speech-weighted ST_V values across room volumes and reverberation times predicted by Pelegrín-García's original model (a), showing plots for Rooms A and B without focusing retroreflection. Values calculated using the extended prediction model (b) are shown for the same rooms with a focusing retroreflector (Height 2.5 m and 60° off).

4 CONCLUSIONS

Two university classrooms were used to explore the potential of improving didactic voice support through the introduction of focusing retroreflectors that reflect high-frequency energy towards the source position. The results of the study were analysed using synthesised oral–binaural room impulse responses (OBRIRs), obtained by combining the OBRIRs measured in original classrooms without retroreflection and the OBRIRs measured in a separate room with the focusing retroreflector, as if the retroreflectors had been installed in the classroom. The focusing retroreflectors were installed at heights of 2.0 m and 2.5 m above the floor, in combinations of angles of 90° and 60° on the median vertical plane from the mouth of the head and torso simulator (HATS). Since retroreflection occurred from 3 to 13 ms due to the distance between the focusing retroreflectors and the ears of the HATS, the signal synthesis was carried out by substituting the corresponding time window in the OBRIR of the original room. Accordingly, ST_V was analysed from the synthesised signal, which included retroreflection within 3 to 13 ms. In addition, an extended version of Pelegrín-García's ST_V prediction model was proposed by incorporating retroreflected energy, and the predicted results were compared with those derived from the signal synthesis.

The results of the experiment are as follows:

1. Both rooms with synthesised retroreflection showed notable improvements at higher frequencies (2 and 4 kHz).
2. The ST_V values were greater when the focusing retroreflector was positioned closer to the HATS and at 60° along the median vertical axis.
3. Rooms with larger volumes and shorter reverberation times are more affected by retroreflection with respect to voice support enhancement, exhibiting greater increases in overall speech-weighted ST_V .
4. The extended prediction model showed close agreement with the ST_V values obtained from signal synthesis at 2 and 4 kHz.

Through signal synthesis, this study verified the potential of focusing retroreflectors to enhance voice support for talkers in the high-frequency range. Nonetheless, the research was limited by the lack of field experiments, and it did not investigate the extent to which speakers regulate their voice level in response to retroreflection.

Future research should aim to validate the model by comparing it with more empirical results, such as classrooms with physically installed focusing retroreflectors, wave-based acoustic simulations, and investigation into the potential to reduce vocal strain by preventing talkers from unnecessarily raising their voice levels. In addition, the effect of combining retroreflection with spaces where speakers use audio systems could also be investigated. It would also be possible to conduct experiments using focusing retroreflectors of different sizes, designs, and quantities.

ACKNOWLEDGEMENTS

The authors thank The University of Sydney School of Architecture, Design and Planning for providing facilities. This research was supported by the Australian Government through the Australian Research Council's Discovery Projects funding scheme (project DP230101357).

REFERENCES

- Berg, Frederick S, James C Blair, and Peggy V Benson. 1996. *Classroom acoustics: The problem, impact, and solution*. Language, Speech, and Hearing Services in Schools 27 (1): 16-20.
- Bistafa, Sylvio R, and John S Bradley. 2000. *Reverberation time and maximum background-noise level for classrooms from a comparative study of speech intelligibility metrics*. The Journal of the Acoustical Society of America 107 (2): 861-875.
- Brunskog, Jonas, Anders Christian Gade, Gaspar Payá Bellester, and Lilian Reig Calbo. 2009. "Increase in voice level and speaker comfort in lecture rooms." *The Journal of the Acoustical Society of America* 125 (4): 2072-2082.
- Cabrera, Densil, Jonathan Holmes, Hugo Caldwell, Manuj Yadav, and Kai Gao. 2018. *An unusual instance of acoustic retroreflection in architecture—Ports 1961 Shanghai flagship store façade*. Applied Acoustics 138: 133-146.
- Cabrera, Densil, Manuj Yadav, Jonathan Holmes, Osborn Fong, and Hugo Caldwell. 2020. *Incidental acoustic retroreflection from building façades: Three instances in Berkeley, Sydney and Hong Kong*. Building and Environment 172: 106733.
- Cabrera, Densil, Jonathan Holmes, Shuai Lu, Mary Rapp, Manuj Yadav, and Oliver Hutchison. 2021. *Voice support from acoustically retroreflective surfaces*. Proceedings of the Euronoise.
- Cabrera, Densil, Shuai Lu, Jonathan Holmes, and Manuj Yadav. 2023. *The potential of focusing acoustic retroreflectors for architectural surface treatment*. Applied Sciences 13 (3): 1547.
- Cabrera, Densil, Jonathan Holmes, and Shuai Lu. 2024. *Concentration of reflected sound in a room treated with cube corner retroreflectors*. The Journal of the Acoustical Society of America 155 (3): 1747-1758.
- Choi, Young-Ji. 2014. *An optimum combination of absorptive and diffusing treatments for classroom acoustic design*. Building Acoustics 21 (2): 175-179.
- Chu, Wing Tin, and ACC Warnock. 2002. "Detailed directivity of sound fields around human talkers." *Report B3144* 6.
- Gheller, Flavia, Elisa Lovo, Athena Arsie, and Roberto Bovo. 2020. *Classroom acoustics: Listening problems in children*. Building Acoustics 27 (1): 47-59.
- Lombard, Etienne. 1911. *Le signe delevation de la voix*. Annu. maladies oreille larynx nez pharynx 27: 101-119.
- Lu, Shuai, Densil Cabrera, Jonathan Holmes, and Ross Ferraro. 2023. *Cube-corner retroreflectors for acoustic support outdoors: A comparison of simple and optimised designs*. Building and Environment 236: 110268.
- Lu, Shuai, Densil Cabrera, and Jonathan Holmes. 2026. *Geometric design and acoustic performance of dihedral focusing retroreflectors based on confocal conics*. Applied Acoustics 241: 111031.
- Mealings, Kiri. 2022. *The effect of classroom acoustic conditions on literacy outcomes for children in primary school: A review*. Building Acoustics 29 (1): 135-156.
- Pelegri-García, David, Viveka Lyberg-Åhlander, Roland Rydell, Jonas Brunskog, and Anders Lofqvist. 2010. *Influence of classroom acoustics on the voice levels of teachers with and without voice problems: A field study*. Proceedings of Meetings on Acoustics.
- Pelegri-García, David, Jonas Brunskog, Viveka Lyberg-Åhlander, and Anders Löfqvist. 2012. *Measurement and prediction of voice support and room gain in school classrooms*. The Journal of the Acoustical Society of America 131 (1): 194-204.
- Rapp, Mary, Densil Cabrera, and Manuj Yadav. 2021. *Effect of voice support level and spectrum on conversational speech*. The Journal of the Acoustical Society of America 150 (4): 2635-2646.
- Sierra-Polanco, Tomás, Lady Catherine Cantor-Cutiva, Eric J Hunter, and Pasquale Bottalico. 2021. *Changes of voice production in artificial acoustic environments*. Frontiers in Built Environment 7: 666152.