

PREDICTION OF OBJECT GEOMETRY FROM ACOUSTIC SCATTERING USING CONVOLUTIONAL NEURAL NETWORKS

Ziqi Fan¹, Vibhav Vineet², Chenshen Lu¹, T.W. Wu³, Kyla McMullen¹

¹University of Florida, Gainesville, USA

²Microsoft Research, Redmond, USA

³University of Kentucky, Lexington, USA

ABSTRACT

Acoustic scattering is strongly influenced by boundary geometry of objects over which sound scatters. The present work proposes a method to infer object geometry from scattering features by training convolutional neural networks. The training data is generated from a fast numerical solver developed on CUDA. The complete set of simulations is sampled to generate multiple datasets containing different amounts of channels and diverse image resolutions. The robustness of our approach in response to data degradation is evaluated by comparing the performance of networks trained using the datasets with varying levels of data degradation. The present work has found that the predictions made from our models match ground truth with high accuracy. In addition, accuracy does not degrade when fewer data channels or lower resolutions are used.

Index Terms— Object geometry prediction, acoustic scattering, convolutional neural network, fast acoustic simulation

1. INTRODUCTION

Acoustic reflection models have been extensively researched to predict room geometry [1–20]. A widely-used approach is based on the assumption that walls can be described by parametric lines in a 2-D plane [1, 3–5, 9, 14, 16, 19]. In this approach, ellipses with source-receiver pairs as their focal points are constructed. After a sonic event, acoustic measurements such as time of arrival (TOA) [1, 3, 5, 9, 14, 16, 19] and direction of arrival (DOA) [4, 14] are extracted from the room reflections and used to determine shapes of the ellipses. The parametric lines forming the walls are determined from common tangential lines of the ellipses. In another popular approach, the virtual image sources mirrored from real sources are used to determine the orientations of the walls [6, 7, 10–13, 15, 17, 20]. In this approach, to determine the wall orientation, a series of circles centered at different receiver locations are constructed. The radii of these circles are determined by the TOAs of echos from the wall. The image source is at the intersection of the circles and the wall is perpendicular to the line segment connecting the image source and the real source.

Recently, Lindell et al. [21] proposed predicting the shape of an object hidden behind a corner using measurements of reflections. In their work, acoustic waves experience 3 reflections before finally arriving at a microphone array. A computational model is derived for the mapping from received acoustic waves to surface geometry. The object surface is reconstructed from the microphone-recorded signals. Their work shows that reflection models are able to predict object geometry in more complicated environments.

The previously mentioned studies only use reflection, despite the many other main effects of acoustic scattering. Critical features

such as diffraction and occlusion are neglected. Recently, Fan et al. [22] proposed to predict acoustic scattering from objects using a convolutional neural network (CNN). The CNN was trained using images representing object geometry and acoustic fields. Reflections, diffractions, and occlusions were accurately constructed from the CNN, according to their test cases. Inspired by [22], it is believed that the inverse problem, prediction of object geometry from acoustic scattering using a CNN, is also feasible.

The present work presents the first study that predicts object geometry from acoustic scattering using CNNs. To achieve this goal, we developed a fast acoustic solver on an Nvidia GPU to generate a large amount of training pairs. An example of these pairs is shown in Fig. 1. In reality, collecting data shown in Fig. 1 is expensive, so we aim to reduce the number of channels and image resolution for a more practical solution. Through selecting multiple source locations, differentiating low and high frequency ranges, and varying density of virtual sensors, we created 24 datasets from the original data and a separate CNN was trained for each of them. The 24 CNN models were carefully studied to evaluate the influence of data degradation on prediction accuracy.

Our work shows that synthetic data can be used with CNN models to accurately predict object geometry. Further, using degraded acoustic images does not significantly decrease accuracy. The results indicate that predicting object geometry from acoustic scattering using CNNs is a robust and promising approach. Our training and test data and the numerical solver are shared at the following address: <https://faculty.eng.ufl.edu/soundpad-lab/research/aml/>.

2. METHODOLOGY

2.1. Problem formulation

Let Ω denote a space of interest and $\mathcal{R}$ be a centered subspace of Ω . There exist objects in $\mathcal{R}$. In this initial study, we restrict the number of objects to 1. Let $\mathcal{V}$ and $\partial\mathcal{V}$ denote the volume and boundary of this object. Region $\mathcal{R}$ is inaccessible and information of $\partial\mathcal{V}$ is unknown. Let $\mathcal{S} = \Omega \setminus \mathcal{R}$ be an open region free of objects. A sound source can be placed at any location $\mathbf{x}_s \in \mathcal{S}$. Our task is to predict $\partial\mathcal{V} \subseteq \mathcal{R}$ using acoustic information in $\mathcal{S}$ induced by sources and $\partial\mathcal{V}$.

Pressure field p induced by a point source at $\mathbf{x}_s$ at an angular frequency ω is described by the Helmholtz equation:

$$\nabla^2 p(\omega, \mathbf{x}; \mathbf{x}_s) + k^2 p(\omega, \mathbf{x}; \mathbf{x}_s) = -\delta(\mathbf{x} - \mathbf{x}_s), \quad \mathbf{x} \in \Omega - \mathcal{V}, \quad (1)$$

where $k = \omega/c$ denotes the wave number and δ denotes the Dirac delta. A unique solution of Eq. (1) is ensured by the Sommerfeld radiation condition at infinity [23] and a frequency-dependent bound-

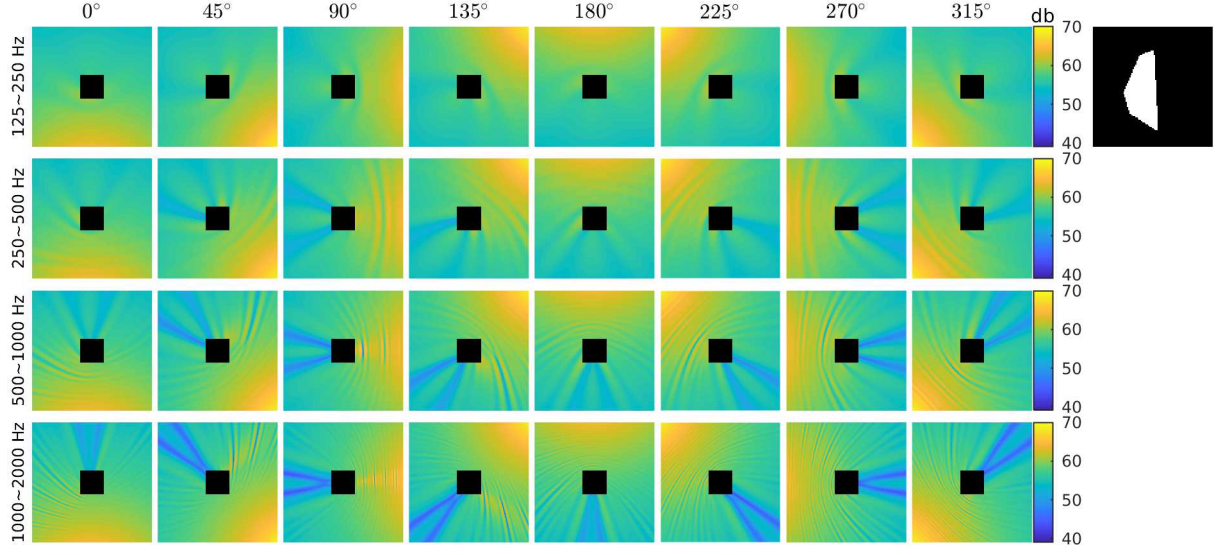


Fig. 1: An input-output pair to the neural network of the largest possible number of channels and the highest resolution. The input is a 32-channel image, 4 channels for octave bands and 8 channels for source locations. Bright yellow pixels far from the dark region denote locations near sound sources. Bright yellow curves radiating from the dark region denote reflection from the object. Dark blue streaks denote occlusion of sound propagation. Green lobes between streaks of occlusion denote diffraction around the object. The output is a binary image indicating the geometry of an object in the inaccessible region. White and black pixels denote object occupancy and air respectively.

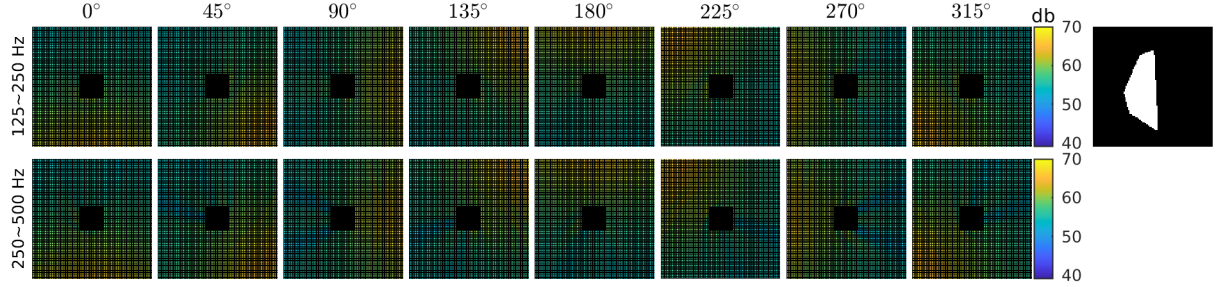


Fig. 2: An input-output pair using low-frequency octave bands only and a spatial sampling factor of 2. Black pixels denote unknown loudness. A much smaller number of measurements is needed for this input, as compared to that in Fig. 1. The number of sources can be reduced to generate images of fewer source channels. Density of virtual probes can be further decreased to lower image resolution.

any condition defined on $\partial\mathcal{V}$:

$$\alpha(\omega, \mathbf{x}) p(\omega, \mathbf{x}) + \beta(\omega, \mathbf{x}) q(\omega, \mathbf{x}) = \gamma(\omega, \mathbf{x}), \quad \mathbf{x} \in \partial\mathcal{V}, \quad (2)$$

where $q = \partial p / \partial n(\mathbf{x})$ denotes the normal velocity on $\partial\mathcal{V}$. Finally, we define loudness at a point $\mathbf{x}$ as spectral energy concentration:

$$L_{i,j}(\mathbf{x}) = 10 \log_{10} \frac{1}{\omega_{i+1} - \omega_i} \int_{\omega_i}^{\omega_{i+1}} |p(\omega, \mathbf{x}; \mathbf{x}_{s,j})|^2 d\omega, \quad \mathbf{x} \in \mathcal{S}, \quad (3)$$

where $\omega_i = [\omega_i, \omega_{i+1})$ denotes the i -th frequency range and $\mathbf{x}_{s,j}$ denotes the j -th source location.

Loudness fields L and boundary $\partial\mathcal{V}$ are related by Eq. (2). Two boundary properties, geometry $\mathcal{G}$ and acoustic impedance $\mathcal{Z}$ are associated with $\partial\mathcal{V}$. In this study, we assume $\partial\mathcal{V}$ to be acoustically rigid, corresponding to the case $\alpha = \gamma = 0, \beta \neq 0$ in Eq. (2). Thus, the only variable related to $\partial\mathcal{V}$ is its geometry $\mathcal{G}$. A total number of N_s source locations $\mathbf{x}_{s,j}$ uniformly distributed around an object and N_f frequency ranges ω_i in the form of octave bands are used. For each $\mathcal{G}$, $N_s \cdot N_f$ loudness fields are defined on $\mathcal{S}$, each loudness field corresponding to a unique combination of ω_i and $\mathbf{x}_{s,j}$. The mapping

$\mathcal{M} : \{L\} \rightarrow \{\mathcal{G}\}$ can be constructed with a large number of data pairs $(L, \mathcal{G})$. Once $\mathcal{M}$ is constructed, prediction of an unknown $\mathcal{G}$ can be conveniently made if L measurements are provided.

2.2. Data representation

In this study, data pairs $(L, \mathcal{G})$ are provided by simulations. Fan et al. proposed a parallel acoustic solver developed in CUDA [24]. This solver is based on GPU implementation of the boundary element method [25] and provides a fast estimation of pressure distributions on $\partial\mathcal{V}$ through solving Eq. (1). The solver was further developed to extrapolate p on $\partial\mathcal{V}$ into $\Omega \setminus \mathcal{V}$ and to calculate L of Eq. (3) in $\mathcal{S}$.

Both L and $\mathcal{G}$ are represented using voxel grids. Similar to [26], space Ω is divided into many small cubes and these cubes form a voxel grid $\mathbf{G}(\Omega) = \mathbf{G}(\mathcal{R}) \cup \mathbf{G}(\mathcal{S})$. Loudness samples $L_{i,j}(\mathbf{x})$ are calculated at every center point of a voxel in $\mathbf{G}(\mathcal{S})$. Consequently, each voxel in $\mathbf{G}(\mathcal{S})$ is associated with a numerical value in decibels. Voxels in $\mathbf{G}(\mathcal{R})$ are associated with a special value denoting unknown, as $\mathcal{R}$ is inaccessible. All the loudness grids form an image $\mathcal{I}(L)$ of $N_f \cdot N_s$ channels, each channel denoting one combination

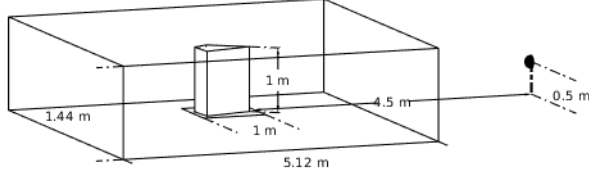


Fig. 3: Illustration of our simulation region. Only 1 out of 8 potential source locations around the region is shown.

of a source location and a frequency range. The geometric distribution of the object in Ω is also defined on $\mathbf{G}(\Omega)$. Each voxel of $\mathbf{G}(\Omega)$ is associated with a binary value, 1 denoting object occupancy and 0 denoting air. This association leads to a binary image $\mathcal{I}(\mathcal{G})$ representing the geometry of the object. As objects are originally represented using mesh models composed of triangular elements, object occupancy is determined from triangle-cube intersection tests.

As a result, both L and $\mathcal{G}$ of an object are represented using 3-D images $\mathcal{I}(L)$ and $\mathcal{I}(\mathcal{G})$. The mapping $\mathcal{M} : \{L\} \rightarrow \{\mathcal{G}\}$ is then turned into a new mapping $\mathcal{M}' : \{\mathcal{I}(L)\} \rightarrow \{\mathcal{I}(\mathcal{G})\}$. CNNs are a good candidate for constructing $\mathcal{M}'$ in the form of images. Our formulation is image segmentation in nature, where each pixel in an output image is of two categories: occupied by an object or not.

3. EXPERIMENTAL EVALUATION

3.1. Simulation setup

The simulation region shown in Fig. 3 is $5.12 \text{ m} \times 5.12 \text{ m} \times 1.44 \text{ m}$. This region is turned into a voxel grid using unit cubes of volume 1 cm^3 . There is a $1 \text{ m} \times 1 \text{ m}$ square region at the center of the simulation region. An object whose geometry is to be determined resides within the square region. The distance from a sound source to the center of the region is fixed at $r_s = 5 \text{ m}$. The sound source can be positioned at any of the following $N_s = 8$ directions:

$$\theta_j = j \cdot \pi/4, j = 0, 1, 2, \dots, 7. \quad (4)$$

We also use $N_f = 4$ frequency ranges in the form of octave bands:

$$\omega_i = [\omega_i, \omega_{i+1}), \omega_i = 2\pi \cdot 125 \cdot 2^{i+1}, i = 0, 1, 2, 3. \quad (5)$$

3.2. Generation of objects and input-output pairs

We make the assumption that objects to be predicted are horizontally convex. We designed the pipeline shown in Fig. 4 for object generation. The pipeline is adapted from [22]. Different from [22], a step of translation is inserted before extrusion to ensure the diversity of object positions within the square region. We generated five categories of prisms, namely, triangular, quadrilateral, pentagonal, hexagonal and heptagonal to simulate training data. There are $N_{tc} = 3150$ instances in each category, summing up to a total number of $N_{tr} = 15750$. Examples of these objects are shown in Fig. 5. We also generated a smaller dataset for generalization test. Instead of collocating vertices on the circle inscribed to the square, we collocated vertices on circles of 4 different radii: 0.3 m, 0.35 m, 0.4 m and 0.45 m. The step scaling was not used in generating objects of the test dataset. For each radius, we generated the same 5 categories of prisms, each of 35 instances, summing up to $N_{tst} = 700$. We compared our test objects to training objects to ensure no overlap.

We simulated loudness fields L for all objects in the training set and the test set. For each object, combinations from $N_s = 8$

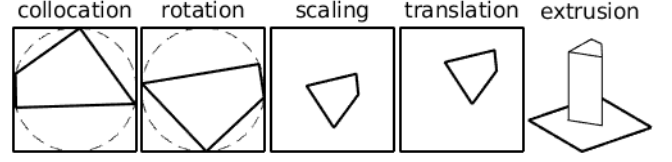


Fig. 4: Our pipeline to generate convex objects of diverse geometry, adapted from [22]. A translation step is added to ensure diverse locations. Objects from this pipeline only vary horizontally.

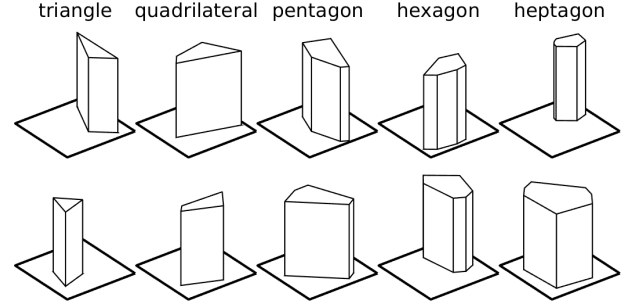


Fig. 5: Examples of objects for our training dataset. The objects cover 5 categories, each one extruded from a type of polygons.

source locations and $N_f = 4$ octave bands lead to 32 independent simulations and they sum up to a 3-D image of 32 channels. We only took the center slice from each channel, as there is no variation for objects in the vertical direction. Thus, the input data of each object is a 32-channel 2-D image. The output data is a 1-channel binary image indicating object occupancy in the horizontal plane. The simulation for 15 750 training objects took 20 days on a Tesla P100 GPU.

An example of our input-output pairs is shown in Fig. 1. On the left-hand side are 32 image channels representing simulations of loudness fields for 4 frequency ranges and 8 source directions. The black region at the center of an image channel denotes region $\mathcal{R}$, which contains an object. Pixels for $\mathcal{S}$ are associated with numerical values in decibels. On the right-hand side is a binary image representing object occupancy in $\mathcal{R}$. Black pixels of value 0 denote air occupancy and white pixels of value 1 denote object occupancy.

3.3. Resolution degradation and channel reduction

Acquisition of input images in Fig. 1 is expensive, considering the spatial density of probes and the numbers of source locations and energy bands. Thus, we aim to study the robustness of our approach to the degradation of image resolution and the reduction of image channels. Our frequency ranges are divided into a low group and a high group, each of 2 octave bands. Using either one of them or both leads to 3 situations. Further, using half of the source locations or using all of them leads to 2 situations. Finally, we define 4 spatial sampling factors: 1, 2, 4 and 8, which correspond to situations of no sampling, sampling every 1 in 2 pixels, every 1 in 4 pixels and every 1 in 8 pixels in 2 dimensions. All the above situations lead to 24 combinations of image channels and resolutions. An example of the dataset using all 8 source locations, low octave bands, and a spatial sampling factor of 2 is shown in Fig. 2.

3.4. Convolutional neural network

We used the full resolution residual network (FRRN) [28] since it has been shown to successfully learn the mapping from geometry to

Table 1: IMED [27] between predictions and ground truth, normalized between 0 and 1. SSF is short for spatial sampling factor.

SSF	4 Sources			8 Sources		
	Frequency Band			Frequency Band		
	Low	High	Full	Low	High	Full
8	0.1453	0.1819	0.1844	0.1420	0.1993	0.1680
4	0.1491	0.2136	0.1753	0.1524	0.1881	0.1610
2	0.1603	0.2036	0.1677	0.1499	0.1880	0.1732
1	0.1497	0.2004	0.1600	0.1480	0.1893	0.1679

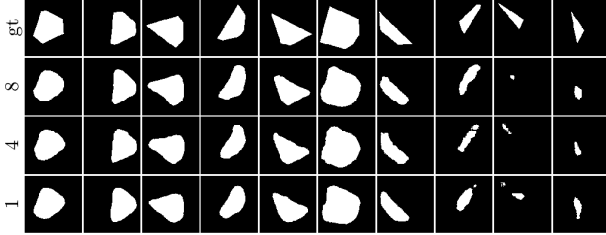


Fig. 6: Comparison of predictions using different spatial subsampling factors. Notation “gt” is short for “ground truth”. Predictions using spatial sampling factors 8, 4, 1 are shown in separate rows. Source numbers are uniformly 8 and only low octave bands are used.

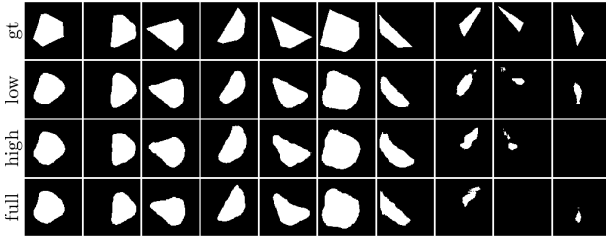


Fig. 7: Comparison of predictions using different combinations of octave bands. Predictions using lower octave bands, higher octave bands and both are illustrated in different rows. Numbers of sources are uniformly 8 and spatial sampling factors are uniformly 1.

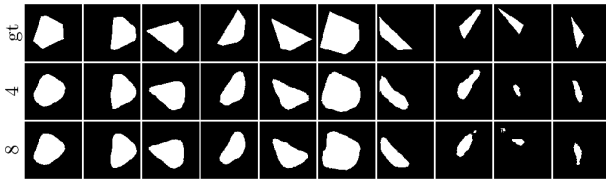


Fig. 8: Comparison of predictions using different numbers of sources. Predictions using 4 and 8 source channels are illustrated in different rows. Spatial sampling factors are uniformly 1 and only low octave bands are used.

loudness [22], adopting the implementation from Shah [29]. We kept the architecture of the FRRN unchanged from Pohlen’s work [28] and modified the numbers of input and output channels to our needs. Cross-entropy was used as the loss metric. Targets were selected as the binary images of size 100×100 denoting the inaccessible region. We also adopted mini-batches of size 2, selected stochastic gradient descent as our optimizer and set the learning rate, momentum, and weight decay as $1e-5$, 0.99 and $1.0e-5$ respectively. Combinations of spatial sampling factors, octave bands, and source locations were

differentiated in the data loader. We trained 24 networks corresponding to the 24 datasets on a GPU cluster. Each network was trained for at least 4 epochs. Training took a total of 4 days.

3.5. Generalization test

Our test data was processed in the same manner as the training data. As a result, we obtained 24 test datasets for the 24 CNN models. Test input images were fed into their corresponding models and each model output 700 predictions for geometry of test objects. Examples of the geometric predictions are shown in Fig. 6, Fig. 7 and Fig. 8.

For the comparison in Fig. 6, choices of octave bands and number of sources are kept the same and the only variable is the spatial sampling factor. In general, severe degradation is not observed when a coarser spatial resolution is used, as shown by a comparison between row 2, 3 and 4. A direct comparison between the predictions and the ground truth reveals that sharp corners in ground-truth objects are rounded in predictions. This discrepancy is induced by the smoothing effect of CNNs. It should also be noted that the predictions of smaller objects are not as accurate as larger ones, as shown by a comparison between columns 8, 9, and 10 and all previous columns. Specifically, only small portions of the triangles in columns 9 and 10 are predicted and the predicted object boundaries diverge from the ground truth. Similar tendencies can also be found in Fig. 7 and Fig. 8, where the variables for comparison are respectively frequency ranges and source distributions.

The Image Euclidean Distance (IMED) [27] is used to quantify prediction accuracy of the 24 models. Surprisingly, models using low octave bands are superior to models using high octave bands or both, as shown in Table 1. Diffraction is more prominent at lower frequencies, and acoustic waves tend to “wrap” around objects as diffraction takes place. Thus, one potential explanation for the superiority of using low octave bands is, the “wrapping” effect of diffraction helps probe a whole surface, while only partial local features are revealed by refraction and occlusion, which are more prominent at higher frequencies. Further, using coarser images or fewer source locations does not significantly degrade prediction accuracy. This finding is consistent with direct observation from Fig. 6 and Fig. 8. The only model showing decreased accuracy with spatial sampling is the one trained with 4 source locations and 2 frequency bands. Thus, our CNN models are robust to reduction of both source channels and spatial resolution.

4. CONCLUSION

In this study, inference of object geometry is formulated as image segmentation. Training and test data are generated from a fast acoustic solver. CNN models are shown to predict object geometry with high accuracy. Further, prediction robustness to reduction of acoustic image channels and of spatial resolution is validated. Specifically, models trained using low-frequency images are superior to models trained using high-frequency images or both, indicating the significance of diffraction in the inference.

There is still much to improve in our work. Only one object is considered in our inference and our objects are convex horizontally and do not vary vertically. These simplifications restrict the generalizability of our work. In the future, we aim to infer the geometry of multiple prisms, lifting the restriction on convexity. Further, no background noise is assumed. In a future work, reflections from the ground should be included in datasets. Finally, our models are trained and tested using only simulations, and it would be interesting to examine the robustness of our approach to real measurements.

5. REFERENCES

- [1] F. Antonacci, A. Sarti, and S. Tubaro, "Geometric reconstruction of the environment from its response to multiple acoustic emissions," in *2010 IEEE International Conference on Acoustics, Speech and Signal Processing*, March 2010, pp. 2822–2825.
- [2] D. Ba, F. Ribeiro, C. Zhang, and D. Florêncio, "L1 regularized room modeling with compact microphone arrays," in *2010 IEEE International Conference on Acoustics, Speech and Signal Processing*, March 2010, pp. 157–160.
- [3] Y. E. Baba, A. Walther, and E. A. P. Habets, "Reflector localization based on multiple reflection points," in *2016 24th European Signal Processing Conference (EUSIPCO)*, Aug 2016, pp. 1458–1462.
- [4] A. Canclini, P. Annibale, F. Antonacci, A. Sarti, R. Rabenstein, and S. Tubaro, "From direction of arrival estimates to localization of planar reflectors in a two dimensional geometry," in *2011 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, May 2011, pp. 2620–2623.
- [5] A. Canclini, F. Antonacci, M. R. P. Thomas, J. Filos, A. Sarti, P. A. Naylor, and S. Tubaro, "Exact localization of acoustic reflectors from quadratic constraints," in *2011 IEEE Workshop on Applications of Signal Processing to Audio and Acoustics (WASPAA)*, Oct 2011, pp. 17–20.
- [6] I. Dokmanić, R. Parhizkar, A. Walther, Y. M. Lu, and M. Vetterli, "Acoustic echoes reveal room shape," *Proceedings of the National Academy of Sciences*, vol. 110, no. 30, pp. 12 186–12 191, 2013.
- [7] I. Dokmanić, Y. M. Lu, and M. Vetterli, "Can one hear the shape of a room: The 2-d polygonal case," in *2011 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, May 2011, pp. 321–324.
- [8] J. Filos, A. Canclini, F. Antonacci, A. Sarti, and P. A. Naylor, "Localization of planar acoustic reflectors from the combination of linear estimates," in *2012 Proceedings of the 20th European Signal Processing Conference (EUSIPCO)*, Aug 2012, pp. 1019–1023.
- [9] J. Filos, E. A. Habets, and P. A. Naylor, "A two-step approach to blindly infer room geometries," in *Proc. Int. Workshop Acoust. Echo Noise Control (IWAENC)*. Citeseer, 2010.
- [10] M. Lovedee-Turner and D. Murphy, "Three-dimensional reflector localisation and room geometry estimation using a spherical microphone array," *The Journal of the Acoustical Society of America*, vol. 146, no. 5, pp. 3339–3352, 2019.
- [11] A. H. Moore, M. Brookes, and P. A. Naylor, "Room geometry estimation from a single channel acoustic impulse response," in *21st European Signal Processing Conference (EUSIPCO 2013)*, Sep. 2013, pp. 1–5.
- [12] H. Naseri and V. Koivunen, "Cooperative simultaneous localization and mapping by exploiting multipath propagation," *IEEE Transactions on Signal Processing*, vol. 65, no. 1, pp. 200–211, Jan 2017.
- [13] E. Nastasia, F. Antonacci, A. Sarti, and S. Tubaro, "Localization of planar acoustic reflectors through emission of controlled stimuli," in *2011 19th European Signal Processing Conference*, Aug 2011, pp. 156–160.
- [14] L. Remaggi, P. J. B. Jackson, P. Coleman, and W. Wang, "Room boundary estimation from acoustic room impulse responses," in *2014 Sensor Signal Processing for Defence (SSPD)*, Sep. 2014, pp. 1–5.
- [15] L. Remaggi, P. J. B. Jackson, P. Coleman, and W. Wang, "Acoustic reflector localization: Novel image source reversion and direct localization methods," *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 25, no. 2, pp. 296–309, Feb 2017.
- [16] L. Remaggi, P. J. B. Jackson, W. Wang, and J. A. Chambers, "A 3d model for room boundary estimation," in *2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, April 2015, pp. 514–518.
- [17] L. Remaggi, P. J. Jackson, and P. Coleman, "Source, sensor and reflector position estimation from acoustical room impulse responses," *Proc. Int. Congr. Sound Vibration*, pp. 1–8, 2015.
- [18] S. Tervo and T. Tossavainen, "3D room geometry estimation from measured impulse responses," in *2012 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, March 2012, pp. 513–516.
- [19] F. Antonacci, J. Filos, M. R. P. Thomas, E. A. P. Habets, A. Sarti, P. A. Naylor, and S. Tubaro, "Inference of room geometry from acoustic impulse responses," *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 20, no. 10, pp. 2683–2695, Dec 2012.
- [20] I. Dokmanić, L. Daudet, and M. Vetterli, "From acoustic room reconstruction to slam," in *2016 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, March 2016, pp. 6345–6349.
- [21] D. B. Lindell, G. Wetzstein, and V. Koltun, "Acoustic non-line-of-sight imaging," in *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2019.
- [22] Z. Fan, V. Vineet, H. Gamper, and N. Raghuvanshi, "Fast acoustic scattering using convolutional neural networks," in *2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2020, pp. 171–175.
- [23] E. G. Williams, *Fourier acoustics: sound radiation and nearfield acoustical holography*. Elsevier, 1999.
- [24] Z. Fan, T. Arce, C. Lu, K. Zhang, T. W. Wu, and K. McMullen, "Computation of head-related transfer functions using graphics processing units and a perceptual validation of the computed hrtfs against measured HRTFs," in *Audio Engineering Society Conference: 2019 AES International Conference on Headphone Technology*, Aug 2019.
- [25] T. Wu, "Boundary element acoustics fundamentals and computer codes," 2002.
- [26] N. Raghuvanshi and J. Snyder, "Parametric directional coding for precomputed sound propagation," *ACM Trans. Graph.*, vol. 37, no. 4, pp. 108:1–108:14, July 2018.
- [27] L. Wang, Y. Zhang, and J. Feng, "On the euclidean distance of images," *IEEE transactions on pattern analysis and machine intelligence*, vol. 27, no. 8, pp. 1334–1339, 2005.
- [28] T. Pohlen, A. Hermans, M. Mathias, and B. Leibe, "Full-resolution residual networks for semantic segmentation in street scenes," in *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, July 2017.
- [29] M. P. Shah, "Semantic segmentation architectures implemented in pytorch," <https://github.com/meetshah1995/pytorch-semseg/>, 2017, accessed: June, 2019.