

presentation script
In Which to Trust?
An Artistic Way of Presenting
Aurally Diverse Listening Experiences
by Jay Afrisando

for Aural Diversity Workshop 4
September 16, 2022

SLIDE 1

Hello everyone. I am Jay Afrisando, a multimedia artist, composer, and researcher. Today, I will present my paper, “*In Which to Trust?: An Artistic Way of Presenting Aurally Diverse Listening Experiences*”

SLIDE 2

There are 5 parts in my presentation. First, I will discuss the importance of comparative listening experiences to understand soundscape through aurally diverse perspectives. Second, I will discuss how *In Which to Trust?* came into being, and why. Third, I will discuss the making process of *In Which to Trust?*. Fourth, I will discuss the result of the work. Finally, my final thoughts will conclude this presentation.

SLIDE 3

Perception and soundscape have been research subjects, but how people perceive soundscapes differently is still a big mystery. Solving this mystery will be possible if we acknowledge that everybody hears differently. Here, I won't go into details since this topic has been covered in the previous Aural Diversity Workshops.

Aurally diverse listeners have documented their listening experiences through essays, films, workshops, and forums. Still, such accounts are limited and not easy to find.

In addition, documenting listening experiences is challenging. As Helmholtz said, observing the sensations of the object's objective nature around us is not something we are trained in. We might know what we feel, but manifesting our bodily sensations towards stimuli using tangible media which other people can learn is challenging.

Further, there is a need to comparatively comprehend to what extent aurally diverse listeners perceive the same stimuli. Through comparisons, we will know exactly how people with diverse hearings make sense of their aural experiences.

So, is there any alternative way to facilitate comparative listening experiences and share them with the public?

SLIDE 4

To answer that question, I created an arts-based research project called *In Which to Trust?* that presents comparative listening experiences of aurally diverse listeners.

This project manifests in a 5-channel sound-captioned video and stereo audio installation and uses sound captions—captioned by five aurally diverse listeners—, audio, and visuals as artistic methods to gather and present aurally diverse people’s perceived soundscapes.

The differences between the audiovisual stimuli and the various captioners’ interpretations leave us a choice of what we should refer to when talking about sound: the objects triggering the sensations, our hearing apparatuses, or both? In which will you trust?

For those who don’t know, sound captions are captions that describe sound. Typically, captions are used to convey speech, but nowadays, captions are also used to convey signs and, although slowly but surely, to elaborate sound, which mostly occurs in audiovisual arts like film.

This arts-based research shares comparative listening experiences with a broader range of people through alternative media.

SLIDE 5

Why arts-based research? Before going into the details, I should mention briefly that arts-based research is research that uses arts as a main methodology in any or all steps of the process.

Artistic methodologies allow flexibility in formulating research questions and diversifying the research process and the outcome.

Also, arts have the potential to provide ways of seeing, extending comprehension instead of concretizing meaning.

Further, artistic methodologies allow collaboration to form among the people who participate and can make research questions, outcomes, and dissemination more accessible and engaging to participants and readers/audiences.

I want to highlight the latter because I seek to invite more and more people to be aware of aural diversity and strive to make this research more accessible to non-specialists.

SLIDE 6

Next, I want to show you the making process of *In Which to Trust?*.

First, I invited 5 aurally diverse listeners to collaborate as sound-captioners.

They are Bill Davies whose current hearing condition is a typical age-related loss at high frequencies, a noise-induced loss from loud rock music, and some processing differences associated with being autistic;

Josephine Dickinson whose current hearing condition is profoundly deaf since childhood, totally deaf since 2012, and receiving cochlear implant;

Ed Garland whose current hearing condition is experiencing permanent tinnitus and noise-induced hearing loss above 3kiloHertz;

Alan Jacques whose current hearing condition is having Ménière's disease for 15 years, severe bilateral hearing impairment, able to converse one-to-one with hearing aids, having almost complete cochlear amusia for pitch, and using the psychological "inner ear" to continue playing piano;

and Terry Perdanawati whose current hearing condition is hearing.

All the collaborators have been aware of the notions of aural diversity, although only Perdanawati hasn't had any experience documenting her own listening experience. Nevertheless, all the collaborators have writing skills, which was a good start for this initial effort.

SLIDE 7

Second, the collaborators were provided with 32 short audiovisual scenes with various qualities and continua derived from everyday environments.

All the qualities include loud to soft, pitchy to percussive, cacophonous to monophonic, high-pitched to low-pitched, dry to wet, human to animal to non-living non-human, and musical instruments to outdoor environments.

Some examples of the scenes include a train coming into a train station, machine sound, bird sound, balloon sound, slide flute sound, and traffic sounds.

SLIDE 8

Third, the five collaborators wrote sound captions on the 32 short audiovisual scenes.

I provided them with captioning guidelines to allow them to express their aural listening experiences based on what they truly hear and feel, not the objects that generate the sounds or what is supposed to be the consensus of sonic understandings dictated by the 'normal' hearing.

They were invited to consider (but not limited to) these conditions/situations that will help the writing, including loudness (barely anything, excessively sensitive, etc.), frequency (certain pitches, no pitches, or multiple pitches heard, etc.), sound quality (sharp, blunt, dense, sparse, bright, dark, etc.), body (painful, dizzy, partially different on two ears, etc.), memory (association with past memories but not related to the actual object that generates the sound in the video), timing (may denote timing of the captions), the number of lines of captions, listening mode if needed, and captioning style which allows any style of captioning.

SLIDE 9

Fourth, the final video is presented in black and white to display cohesive colors across the various footage.

This is also helpful in dealing with uncontrolled colors in the environment (such as walls and other objects within the building) where this project is exhibited.

SLIDE 10

Fifth, finally, the project is presented using 5-channel video. An optional approach will be 1-channel video if multichannel monitors are not possible.

This project was premiered at Sound Scene 2022 at the Smithsonian Hirshhorn National Museum of Modern Art in Washington, DC, the United States, on June 4-5.

The viewers were estimated at 500++ people per day who visited the installation, making it 1,000++ viewers in total.

The visitors were diverse in terms of age and ability.

At this exhibition, the project was installed using 5 digital signage monitors which take each clip via a collage mode and a pair of stereo loudspeakers which take stereo audio input.

SLIDE 11

The making process results in a project which is comparative, quirky, iridescent, and accessible.

The sound captions gathered are not only diverse as expected but also showcase various and detailed listening experiences induced by the various audiovisual scenes.

The sound captions show comparative perceived soundscapes toward the particular audiovisual scenes, which help viewers know how different listeners make sense of aural experiences towards the same sources.

The sound captions display subjective, quirky ways of captioning in terms of the words and marks used to express the experiences.

The sound captions exhibit the spectra of experiences, providing specific contexts via brief descriptions corresponding to the multifarious audiovisual scenes.

Furthermore, the form of the research project makes it more accessible to the general public to learn about how aurally diverse people understand soundscape and the existence of aural diversity.

SLIDE 12

In summary, from my perspective, each sound-captioner reveals contrasting perceived soundscape with their unique captioning style.

Davies reveals his soundscape which is mostly sequential and showcases details that other sound-captioners didn't—and possibly other 'normal' listeners do not—perceive. His soundscape exhibits complex and extremely detailed experiences.

Dickinson reveals her soundscape which attempts to verbatim capture her aural experience through onomatopoeia, making it performative. She also uses dynamics used in Western notation (like fortissimo and mezzo piano), imagery, dots, and spacing to denote event sequences.

Garland reveals his soundscape which exhibits architectural aural experience from outer and inner sources of sound, despite the brevity. His soundscape shows that inner ear ringing is a significant factor that makes his soundscape.

Jacques reveals his soundscape which exhibits painful and complex listening experiences, possibly not easy to imagine for ‘normal’ listeners. His sound-captions are denotative, which includes possibilities of his own listening and listening modes through his hearing aids.

Perdanawati reveals her soundscape which is mostly descriptive. Occasionally, she uses analogies of her memories to show the similarity of the experiences. In addition, some of her captions are onomatopoeic, and most of her captions show sequential soundscape events.

It is worth noting that Dickinson’s and Jacques captions reveal that the visuals of the sources are relevant to their soundscapes. Dickinson includes the expression “Julie Mehretu!” to indicate her imagined sound when noticing the word “Julie Mehretu” in the street traffic scene. On the other hand, Jacques includes descriptive explanations that the visuals of the sources helped him make sense of his listening experience. Jacques also told me that he listened to the audiovisual scenes without the visuals at first, and then with the visuals next.

SLIDE 13

Next, I want to show you the trailer of *In Which to Trust?*. It lasts just a bit more than 30 seconds.

[play the video]

SLIDE 14

Finally, I want to share with you my final thoughts.

In Which to Trust? exhibits aurally diverse people’s comparative, various listening experiences on the same stimuli and an alternative approach of sharing perceived soundscapes through arts.

I want to stress that alternative doesn’t mean secondary. Instead, it means the possibility of creating more opportunities for anything.

Instead of limiting, the captions’ brevity—coupled with the audiovisual scenes—reveals the potential of alternative ways for people to learn how we, aurally diverse people, listen.

This project has been a reflection for me that arts-based research offers a prospective approach to creating an alternative form of knowledge through collaboration and making the dissemination of knowledge more friendly and welcoming.

Future projects may:

- invite more aurally diverse people to collaborate in different geographical locations in the future,
- end up with massive video channels, funding permits,
- lead to meaningful workshops on documenting listening experiences for those not having prior skills, highlighting the importance of documenting listening experiences and the significance of experience as building blocks of knowledge, and
- include tactile as an additional or main stimulus.

SLIDE 15

That's it. Thank you!

Much gratitude to the collaborators Bill Davies, Josephine Dickinson, Ed Garland, Alan Jacques, and Terry Perdanawati, and for the Jerome Foundation that supported this project through the Jerome Hill Artist Fellowship 2021-22.

Please let me know if you have any questions.