

UNITED STATES PATENT AND TRADEMARK OFFICE

BEFORE THE PATENT TRIAL AND APPEAL BOARD

GOOGLE LLC
Petitioner

v.

JAWBONE INNOVATIONS, LLC
Patent Owner

Case No. IPR2023-01131
U.S. Patent No. 8,326,611

**PETITION FOR *INTER PARTES* REVIEW
OF U.S. PATENT NO. 8,326,611**

TABLE OF CONTENTS

I.	REQUIREMENTS FOR IPR.....	1
A.	Grounds for Standing.....	1
B.	Challenge and Relief Requested	1
C.	Priority Date.....	2
1.	Dynamic Drinkware Analysis	3
II.	BACKGROUND	4
A.	The '611 Patent.....	4
B.	Prosecution History.....	6
C.	Level of Ordinary Skill	6
D.	Claim Construction.....	6
III.	GROUND 1: AVENDANO AND VISSER (CLAIMS 1-7, 25-28)	7
A.	Avendano Overview	7
B.	Visser Overview	12
C.	Combination of Avendano and Visser.....	15
D.	Claim 1.....	22
E.	Claim 2.....	35
F.	Claims 3 and 4	36
G.	Claims 5 and 6	37
H.	Claim 7.....	38
I.	Claim 25.....	40

J.	Claim 26.....	41
K.	Claims 27 and 28	41
IV.	GROUND 2: Avendano, Visser, And Bisgaard (Claims 8-16, 23, 24).....	43
A.	Bisgaard Overview	43
B.	Combination of Avendano, Visser, and Bisgaard.....	44
C.	Claim 8.....	47
D.	Claim 9.....	49
E.	Claim 10.....	50
F.	Claim 11.....	51
G.	Claim 12.....	52
H.	Claim 13.....	52
I.	Claim 14.....	53
J.	Claim 15.....	53
K.	Claim 16.....	54
L.	Claim 23.....	54
M.	Claim 24.....	57
V.	GROUND 3: Avendano, Visser, Bisgaard, And Hou (Claims 17- 19).....	58
A.	Hou Overview	58
B.	Combination of Avendano, Visser, Bisgaard, and Hou	61
C.	Claim 17.....	66
D.	Claims 18 and 19	69

VI.	GROUND 4: AVENDANO, VISSER, BISGAARD, HOU, AND FREQUENCY ART (BYRNE, BURNETT, AND/OR BERGLUND) (CLAIMS 20-22)	71
A.	Byrne Overview.....	71
B.	Burnett Overview.....	73
C.	Berglund Overview.....	73
D.	Combination of Avendano, Visser, Bisgaard, Hou, and Frequency Art (Byrne, Burnett, and/or Berglund).....	75
E.	Claim 20.....	77
F.	Claims 21 and 22	81
VII.	INSTITUTION IS APPROPRIATE HERE	91
A.	The <i>Fintiv</i> Factors Support Institution	91
1.	Factor 1: Potential Stay.....	91
2.	Factor 2: Proximity of Trial to FWD.....	92
3.	Factor 3: Investment in Parallel Proceeding.....	92
4.	Factor 4: Overlapping Issues	93
5.	Factor 5: Parties in Parallel Proceedings	93
6.	Factor 6: Other Circumstances	93
B.	Google’s Previously Filed IPR Petition Does Not Warrant Discretionary Denial (<i>General Plastic</i>)	94
1.	<i>General Plastic</i> Factors 1-3	94
2.	<i>General Plastic</i> Factors 4 and 5.....	97
3.	<i>General Plastic</i> Factors 6 and 7.....	98
C.	Discretionary Denial Under § 325(d) is Also Not Appropriate.....	99

VIII. MANDATORY NOTICES UNDER 37 C.F.R § 42.8(a)(1).....	100
A. Real Party-In-Interest Under 37 C.F.R. § 42.8(b)(1)	100
B. Related Matters Under 37 C.F.R. § 42.8(b)(2)	100
C. Lead And Back-Up Counsel Under 37 C.F.R. § 42.8(b)(3)	103
D. Service Information	104
E. Conclusion.....	104

TABLE OF AUTHORITIES

	Page(s)
Cases	
<i>Apple Inc. v. Fintiv, Inc.</i> , IPR2020-00019, Paper 11 (P.T.A.B. Mar. 20, 2020)	91, 94
<i>Apple Inc. v. Jawbone Innovations, LLC</i> , IPR2022-01085, Paper 2 (P.T.A.B. Jun. 3, 2022)	95
<i>Apple Inc. v. Jawbone Innovations, LLC</i> , IPR2022-01085, Paper 14 (P.T.A.B. Jan. 9, 2023)	96
<i>Apple Inc. v. Uniloc 2017 LLC</i> , IPR2020-00854, Paper 9 (P.T.A.B. Oct. 28, 2020)	94
<i>Cisco Sys., Inc. v. Centripetal Networks, Inc.</i> , IPR2022-01151, Paper 12 (P.T.A.B. Jan. 4, 2023)	95
<i>ClearValue, Inc. v. Pearl River Polymers, Inc.</i> , 668 F.3d 1340 (Fed. Cir. 2012)	81
<i>Code200, UAB v. Bright Data Ltd.</i> , IPR2022-00861, Paper 18 (P.T.A.B. Aug. 23, 2022)	94, 95, 99
<i>E.I. DuPont de Nemours & Co. v. Synvina C.V.</i> , 904 F.3d 996 (Fed. Cir. 2018)	81
<i>Global Tel*Link Corp. v. HLFIP Holding, Inc.</i> , IPR2021-00444, Paper 14 (P.T.A.B. Jul. 22, 2021)	92
<i>Google LLC v. Jawbone Innovations, LLC</i> , IPR2022-00604, Paper 8 (P.T.A.B. Jul. 11, 2022)	95
<i>Google LLC v. Jawbone Innovations, LLC</i> , IPR2022-00604, Paper 12 (P.T.A.B. Oct. 6, 2022)	95, 98
<i>Google LLC v. Jawbone Innovations, LLC</i> , IPR2022-00630, Paper 10 (P.T.A.B. Sept. 13, 2022)	93

<i>Samsung Elecs. Am. Inc. v. Snik LLC</i> , IPR2020-01428, Paper 10 (P.T.A.B. Mar. 9, 2021)	92
<i>Sand Revolution II, LLC v. Cont'l Intermodal Grp.-Trucking LLC</i> , IPR2019-01393, Paper 24 (P.T.A.B. Jun. 16, 2020)	93
<i>Well-man, Inc. v. Eastman Chem. Co.</i> , 642 F.3d 1355 (Fed. Cir. 2011)	6
Statutes	
35 U.S.C. § 315(c)	96
35 U.S.C. § 316(a)(11).....	99
35 U.S.C. § 325(d)	99

LIST OF EXHIBITS

Exhibit	Description
EX. 1001	U.S. Patent No. 8,326,611 to Petit <i>et al.</i> (“the ’611 patent”)
EX. 1002	Excerpts from the Prosecution History of the ’611 patent (“the Prosecution History”)
EX. 1003	Declaration of Dr. Thomas Kenny
EX. 1004	<i>Curriculum Vitae</i> of Dr. Thomas Kenny
EX. 1005	U.S. Patent No. 8,194,880 B2 (“Avendano ’880”)
EX. 1006	U.S. Patent No. 7,464,029 B2 (“Visser”)
EX. 1007	RESERVED
EX. 1008	U.S. Patent No. 7,155,019 B2 (“Hou”)
EX. 1009	Byrne, D, et al, “An international comparison of long-term average speech spectra,” 1994 Oct; J. Acoust. Soc. Am.; 96(4): 2108-2120 (“Byrne”).
EX. 1010	U.S. Publication No. US 2011/0103626 A1 (“Bisgaard”)
EX. 1011	U.S. Provisional App. No. 60/816,244 (“the Bisgaard Provisional”)
EX. 1012	U.S. Publication No. US 2002/0198705 A1 (“Burnett”)
EX. 1013	Berglund, B, et al, “Sources and effects of low-frequency noise,” 1996 May; J. Acoust. Soc. Am; 99(5): 2985-3002 (“Berglund”).
EX. 1014	Declaration of June Ann Munford
EX. 1015	Declaration of June Ann Munford – Appendix
EX. 1016	<i>Curriculum Vitae</i> of June Ann Munford
EX. 1017	Declaration of Dr. David V. Anderson

Exhibit	Description
EX. 1018	<i>Curriculum Vitae</i> of Dr. David V. Anderson
EX. 1019	<i>Amazon.com, Inc. v. Jawbone Innovations, LLC</i> , IPR2023-00286, Decision on Institution, Paper 10 (Jun. 7, 2023)
EX. 1020	<i>Jawbone Innovations, LLC v. Google LLC</i> , No. 6:21-cv-00985, ECF 105, Order Transferring Case (W.D. Tex. Feb. 1, 2023)
EX. 1021	<i>Jawbone Innovations, LLC v. Google LLC</i> , 3:23-cv-00466, ECF 137, Order Granting Stay Pending <i>Inter Partes</i> Review (N.D. Cal. Apr. 27, 2023)

LISTING OF CHALLENGED CLAIMS

Claim 1	
[1pre]	A method comprising:
[1a]	forming a first virtual microphone by combining a first signal of a first physical microphone and a second signal of a second physical microphone;
[1b]	forming a filter that describes a relationship for speech between the first physical microphone and the second physical microphone;
[1c]	forming a second virtual microphone by applying the filter to the first signal to generate a first intermediate signal, and summing the first intermediate signal and the second signal;
[1d]	generating an energy ratio of energies of the first virtual microphone and the second virtual microphone; and
[1e]	detecting acoustic voice activity of a speaker when the energy ratio is greater than a threshold value.
Claim 2	
[2]	The method of claim 1, wherein the first virtual microphone and the second virtual microphone are distinct virtual directional microphones.
Claim 3	
[3]	The method of claim 2, wherein the first virtual microphone and the second virtual microphone have approximately similar responses to noise.
Claim 4	
[4]	The method of claim 3, wherein the first virtual microphone and the second virtual microphone have approximately dissimilar responses to speech.
Claim 5	

[5]	The method of claim 1, comprising applying a calibration to at least one of the first signal and the second signal.
Claim 6	
[6]	The method of claim 5, wherein the calibration compensates a second response of the second physical microphone so that the second response is equivalent to a first response of the first physical microphone.
Claim 7	
[7]	The method of claim 5, comprising applying a delay to the first intermediate signal.
Claim 8	
[8]	The method of claim 7, wherein the delay is proportional to a time difference between arrival of the speech at the second physical microphone and arrival of the speech at the first physical microphone.
Claim 9	
[9]	The method of claim 8, wherein the forming of the first virtual microphone comprises applying the filter to the second signal.
Claim 10	
[10]	The method of claim 9, wherein the forming of the first virtual microphone comprises applying the calibration to the second signal.
Claim 11	
[11]	The method of claim 10, wherein the forming of the first virtual microphone comprises applying the delay to the first signal.
Claim 12	

[12]	The method of claim 11, wherein the forming of the first virtual microphone by the combining comprises subtracting the second signal from the first signal.
Claim 13	
[13]	The method of claim 12, wherein the filter is an adaptive filter.
Claim 14	
[14]	The method of claim 13, comprising adapting the filter to minimize a second virtual microphone output when only speech is being received by the first physical microphone and the second physical microphone.
Claim 15	
[15]	The method of claim 13, wherein the adapting comprises applying a least-mean squares process.
Claim 16	
[16]	The method of claim 13, comprising generating coefficients of the filter during a period when only speech is being received by the first physical microphone and the second physical microphone.
Claim 17	
[17]	The method of claim 13, wherein the forming of the filter comprises: generating a first quantity by applying a calibration to the second signal; generating a second quantity by applying the delay to the first signal; forming the filter as a ratio of the first quantity to the second quantity.
Claim 18	
[18]	The method of claim 17, wherein the generating of the energy ratio comprises generating the energy ratio for a frequency band.
Claim 19	

[19]	The method of claim 17, wherein the generating of the energy ratio comprises generating the energy ratio for a frequency subband.
Claim 20	
[20]	The method of claim 19, wherein the frequency subband includes frequencies higher than approximately 200 Hertz (Hz).
Claim 21	
[21]	The method of claim 19, wherein the frequency subband includes frequencies in a range from approximately 250 Hz to 1250 Hz.
Claim 22	
[22]	The method of claim 19, wherein the frequency subband includes frequencies in a range from approximately 200 Hz to 3000 Hz.
Claim 23	
[23]	The method of claim 12, wherein the filter is a static filter.
Claim 24	
[24]	The method of claim 23, wherein the forming of the filter comprises: determining a first distance as distance between the first physical microphone and a mouth of the speaker; determining a second distance as distance between the second physical microphone and the mouth; and forming a ratio of the first distance to the second distance.
Claim 25	
[25]	The method of claim 1, comprising generating a vector of the energy ratio versus time.
Claim 26	
[26]	The method of claim 1, wherein the first and second physical microphones are omnidirectional microphones.
Claim 27	

[27]	The method of claim 1, comprising positioning the first physical microphone and the second physical microphone along an axis and
	separating the first physical microphone and the second physical microphone by a first distance.
Claim 28	
[28]	The method of claim 27, wherein a midpoint of the axis is a second distance from a mouth of the speaker, wherein the mouth is located in a direction defined by an angle relative to the midpoint.

Petitioner Google LLC (“Petitioner”) requests an *inter partes* review (“IPR”) of claims 1-28 (the “Challenged Claims”) of U.S. Patent No. 8,326,611 (“the ’611 Patent”). This petition is substantively the same as the petition in IPR2023-00286 (which is instituted) and is concurrently filed with a motion requesting joinder with that proceeding.

I. REQUIREMENTS FOR IPR

A. Grounds for Standing

Petitioner certifies that the ’611 patent is available for IPR. This petition is accompanied by a motion for joinder. Pursuant to 37 C.F.R. § 42.122(b), Petitioner is not barred or estopped from requesting this review.

B. Challenge and Relief Requested

Petitioner requests IPR on the following grounds.

Ground	Claims	§103 Basis
1	1-7, 25-28	Avendano, Visser
2	8-16, 23, 24	Avendano, Visser, Bisgaard
3	17-19	Avendano, Visser, Bisgaard, Hou
4	20-22	Avendano, Visser, Bisgaard, Hou, and Frequency Art (Byrne, Burnett, and/or Berglund)

C. Priority Date

The '611 patent was filed 10/26/2009 as a continuation-in-part of applications filed 05/25/2007 and 06/13/2008, and claims priority to a provisional application filed 10/24/2008.

The Challenged Claims are not entitled to the 05/25/2007 and 06/13/2008 dates because neither CIP application discloses: “generating an energy ratio of energies of the first virtual microphone and the second virtual microphone” and “detecting acoustic voice activity of a speaker when the energy ratio is greater than a threshold value.” Thus, the earliest possible priority date is 10/24/2008 (“Critical Date”).

Each reference qualifies as prior art:

Reference	Date	Section
Avendano	01/29/2007 (filed)	§102(e)
Visser	07/22/2005 (filed)	§102(e)
Bisgaard	06/25/2007 (filed) 06/23/2006 (filed, provisional application)	§102(e)
Hou	03/14/2001 (filed)	§102(e)
Byrne	October 1994 (published) ¹	§102(b)
Burnett	12/26/2002 (published)	§102(b)
Berglund	May 1996 (published) ²	§102(b)

¹ Ex. 1013, ¶¶6-8.

² *Id.*, ¶¶9-11.

Bisgaard qualifies as prior art because its filing date (06/25/2007) and the filing date of its provisional application (06/23/2006) predate the Critical Date.

1. Dynamic Drinkware Analysis

Bisgaard claims priority to U.S. 60/816,244 (“Bisgaard Provisional”). Ex. 1011, Cover. The Bisgaard Provisional is incorporated in its entirety in Bisgaard. Ex. 1010, [0001]. Bisgaard and the Bisgaard Provisional share a similar specification and similar claims. Ex. 1010, [0002]-[0006], [0008]-[0019], [0022]-[0027], [0033]-[0083], [0088]-[0090], claim 1; Ex. 1011, 1:3-3:18, 3:26-12:18, claim 1. Bisgaard is entitled to the 06/23/2006 filing date because the Bisgaard Provisional includes the relevant prior art disclosure and supports at least one of Bisgaard’s claims (claim 1), as shown below.

a. A hearing instrument, comprising:

Ex. 1011, claim 1, 1:3-5, 2:1-9, 2:15, 4:1-4, FIG. 1; Ex. 1003, ¶81.

b. at least two microphones for reception of sound and conversion of the received sound into corresponding electrical sound signals that are input to the signal processor;

Ex. 1011, claim 1, 1:13-15, 2:1-9, 4:1-9, 5:1-20, FIGS. 1-2; Ex. 1003, ¶81.

c. wherein the signal processor is configured to process the electrical sound signals into a combined signal with

a directivity pattern with at least one adaptive null direction θ ; and

Ex. 1011, claim 1, 2:1-9, 3:3-13, 5:21-6:6, 11:30-12:18, FIGS. 1-2; Ex. 1003,

¶81.

- d. wherein the signal processor is further configured to prevent the at least one null direction θ from entering a prohibited range of directions, wherein the prohibited range is a function of a parameter of the electrical sound signals.**

Ex. 1011, claim 1, 2:6-9, 3:3-13, 6:18-19; Ex. 1003, ¶81.

II. BACKGROUND

A. The '611 Patent

The '611 patent “relates to noise suppression systems, devices, and methods for use in acoustic applications.” Ex. 1001, 1:16-18. A first virtual microphone (V_1) is generated by (i) applying a delay filter ($z^{-\gamma}$) to a signal from a first physical microphone (O_1), (ii) applying a calibration filter ($\alpha(z)$) and an adaptive filter ($\beta(z)$) to a signal from a second physical microphone (O_2), and (iii) combining the filtered signals. *Id.*, 5:20-6:19, FIG. 4.

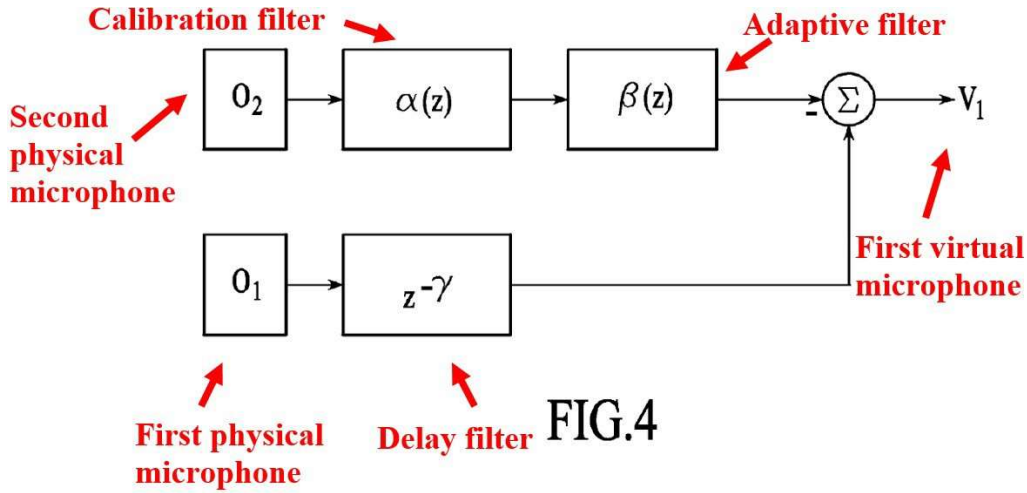


FIG.4

Ex. 1001, FIG. 4³

Further, a second virtual microphone (V_2) is generated by (i) applying an adaptive filter ($\beta(z)$) and a delay filter ($z^{-\gamma}$) to the signal from a first physical microphone (O_1), (ii) applying a calibration filter ($\alpha(z)$) to the signal from a second physical microphone (O_2), and (iii) combining the filtered signals. *Id.*, 5:20-6:19, FIG. 3.

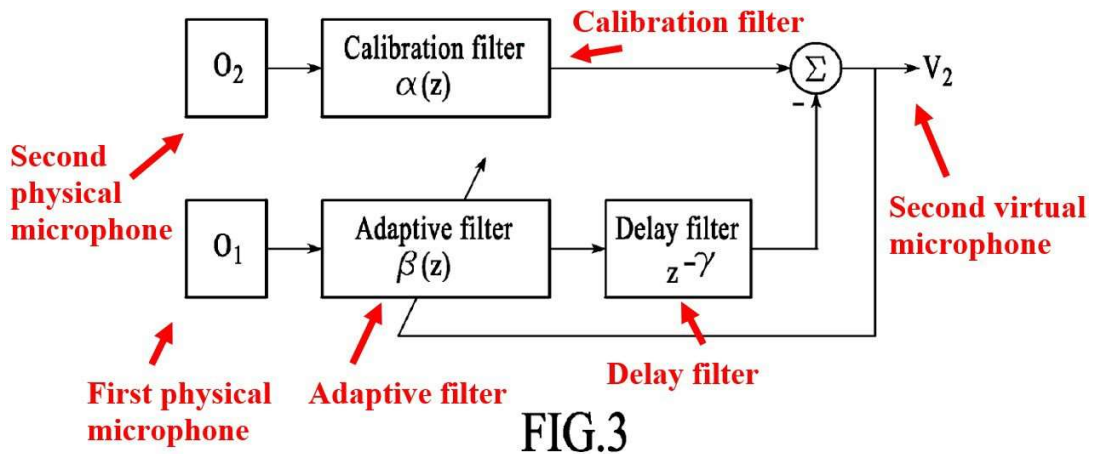


FIG.3

Ex. 1001, FIG. 3

³ Red annotations added throughout.

The ratio of energies of the first and second virtual microphones is used “to determine when speech is occurring.” *Id.*, 6:20-10:8, FIGS. 5-11. A ratio that is greater than a threshold value is indicative of acoustic voice activity, whereas a ratio that is less than the threshold value is indicative of an absence of acoustic voice activity. *Id.*, 6:47-51, 7:5-7, FIGS. 5-11; Ex. 1003, ¶¶42-49.

B. Prosecution History

The claims were allowed after the filing of a terminal disclaimer over U.S. 12/606,146. Ex. 1002, 195-196, 218-219, 227-233.

C. Level of Ordinary Skill

A person of ordinary skill in the art (“POSITA”) would have at least a bachelor of science in electrical engineering, computer engineering, computer science, mechanical engineering, or a related discipline, with at least two years of relevant experience in a field related to acoustics, speech recognition, speech detection, or signal processing. Ex. 1003, ¶¶22-23. Additional education or industry experience may compensate for a deficit in the other. *Id.*

D. Claim Construction

No formal claim constructions are necessary because “claim terms need only be construed to the extent necessary to resolve the controversy.” *Well-man, Inc. v.*

Eastman Chem. Co., 642 F.3d 1355, 1361 (Fed. Cir. 2011).⁴

III. GROUND 1: AVENDANO AND VISSER (CLAIMS 1-7, 25-28)

A. Avendano Overview

Avendano determines “inter-microphone level differences (ILD) ... based on energy level differences of a pair of omni-directional microphones,” and uses ILD “to attenuate noise and enhance speech.” Ex. 1005, 2:5-9.

Avendano discloses “audio device 104” having “primary microphone 106” and “secondary microphone 108,” which may be “omni-directional microphone[s].” *Id.*, 3:27-35; FIGS. 1a-1b.

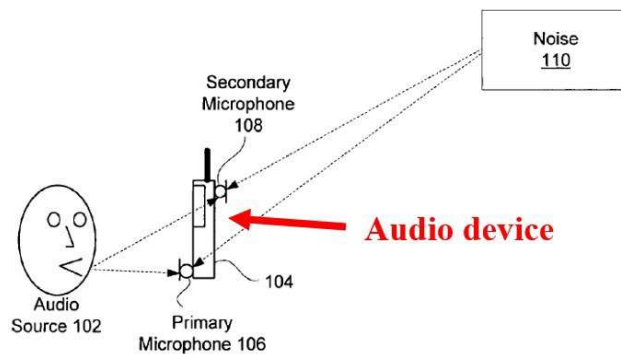
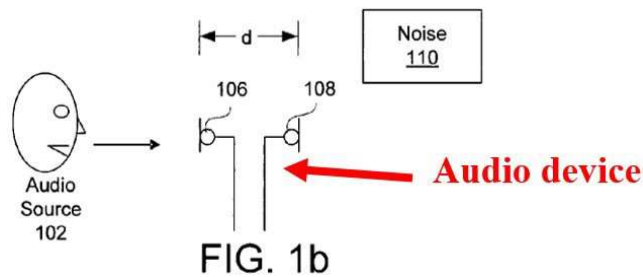


FIG. 1a

Ex. 1005, FIG. 1a

⁴ Petitioner is neither conceding that each claim satisfies all statutory requirements, such as §§101 and 112, nor waiving any arguments concerning claim scope or grounds that can only be raised in district court.



Ex. 1005, FIG. 1b

Avendano’s “primary microphone 106 is much closer to [an] audio source 102 than the secondary microphone 108,” and thus “the intensity level is higher for the primary microphone 106 resulting in a larger energy level during a speech/voice segment.” *Id.*, 3:45-55, FIGS. 1a-1b.

Avendano uses this “level difference ... to discriminate speech and noise in the time-frequency domain.” *Id.*, 3:55-57. For example, Avendano receives signals from the two microphones (signals x_1 and x_2), and processes the signals using “differential microphone array (DMA) module 302” to “create two different directional patterns around the audio device 104.” *Id.*, 4:20-41. As Avendano explains, “[e]ach directional pattern is a region about the audio device 104 in which sounds generated by an audio source 102 within the region may be received by the microphones 106 and 108 with little attenuation,” and “[s]ounds generated by audio sources 102 outside of the directional pattern may be attenuated.” *Id.*, 4:41-46.

Avendano’s DMA module 302 generates (i) a first processed signal having a directional pattern for receiving sounds “within a front cardioid region around the

audio device 104” (*i.e.*, “cardioid primary signal (C_f)”), and (ii) a second processed signal having a directional pattern for receiving sounds “within a back cardioid region around the audio device 104” (*i.e.*, “cardioid secondary signal (C_b)”). *Id.*, 4:47-52, 5:25-35, 9:29-42, Figure 6 (below).

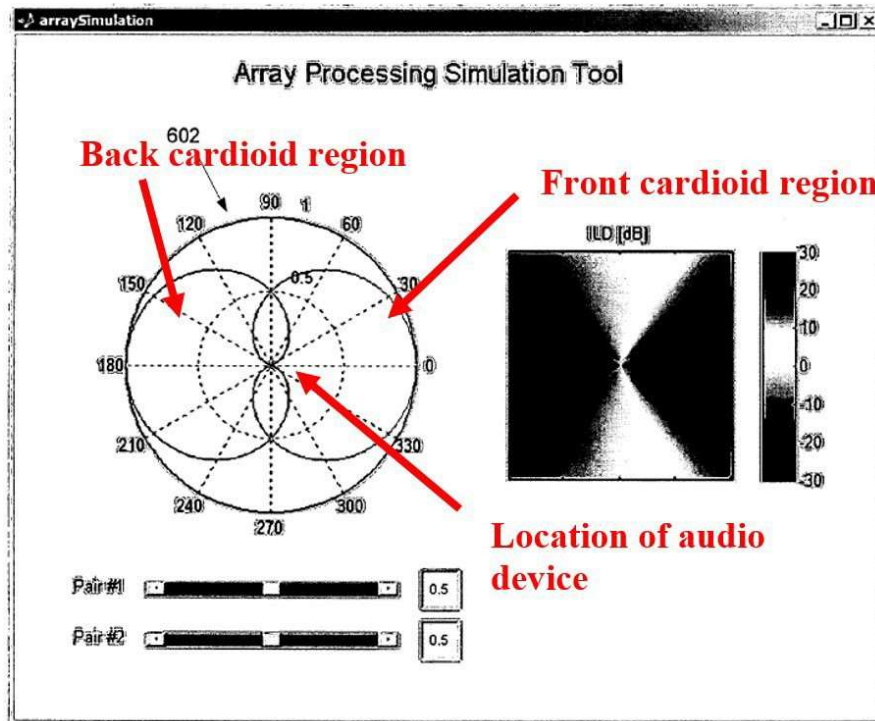


FIG. 6

Ex. 1005, FIG. 6

Avendano’s “cardioid primary signal (C_f)” is generated by combining (i) signal x_1 from primary microphone 106, and (ii) signal x_2 from secondary microphone 108 (signal x_2 having been filtered by “delay node 404” and “gain module 406”). *Id.*, 5:15-35, FIG. 4a. Avendano’s “cardioid secondary signal (C_b)” is generated by combining (i) signal x_2 from secondary microphone 108, and (ii)

signal x_1 from primary microphone 106 (signal x_1 having been filtered by “delay node 402”). *Id.* The “delay nodes” are implemented using filters (e.g., “allpass filters”). *Id.*, 8:47-51.

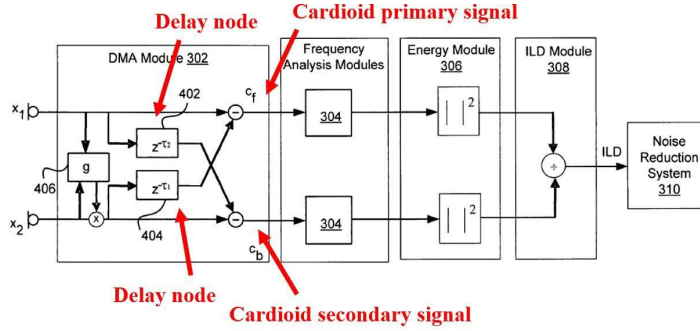


FIG. 4a

Ex. 1005, FIG. 4a

Further, Avendano detects speech based on the ratio between (i) the energy of “cardioid primary signal (C_f)” and (ii) the energy of “cardioid secondary signal (C_b).” *Id.*, 5:49-6:34. Specifically, an “energy level” (E_f) associated with “cardioid primary signal (C_f)” is calculated:

$$E_f(t, \omega) = \int_{frame} |C_f(t', \omega)|^2 dt',$$

Ex. 1005, 5:60

Further, an “energy level” (E_b) associated with “cardioid secondary signal (C_b)” is calculated:

$$E_b(t, \omega) = \int_{frame} |C_b(t', \omega)|^2 dt'.$$

Ex. 1005, 6:5

The ratio between these two energy levels (ILD) is determined:

$$ILD(t, \omega) = \frac{\int |C_f(t', \omega)|^2 dt'}{\int_{frame} |C_b(t', \omega)|^2 dt'} \cdot$$

 **Energy level of C_f**

 **Energy level of C_b**

Ex. 1005, 6:16

Avendano compares the ratio (ILD) to a “threshold” to determine the presence or absence of speech:

$$\lambda_I(t, \omega) = \begin{cases} \approx 0 & \text{if } ILD(t, \omega) < \text{threshold} \\ \approx 1 & \text{if } ILD(t, \omega) > \text{threshold} \end{cases}$$

Ex. 1005, 6:61

If the ratio (ILD) is “smaller than a threshold value (e.g., threshold=0.5) above which speech is expected to be,” a value λ_1 is set to zero (e.g., indicating an absence of speech). *Id.*, 6:58-7:3. However, if the ratio (ILD) “starts to rise (e.g., because speech is present within the large ILD region), λ_1 increases” (e.g., is set to one, indicating a presence of speech). *Id.*

Avendano’s ratio (ILD) is used to process audio signals “through a noise reduction system 310” to “enhance the speech of the primary acoustic signal.” *Id.*, Abstract, 6:35--8:23, 10:18-50, FIGS. 7-8; Ex. 1003, ¶¶51-67.

B. Visser Overview

Visser “improv[es] the quality of a speech signal extracted from noisy acoustic environment” using a “voice activity detector.” Ex. 1006, 6:57-60. Visser’s “speech separation process 100” separates speech from “sound signals from microphones ... 102 and 104.” *Id.*, 8:4-8, FIG. 1. Specifically, “voice activity detector (VAD) 106 ... receives two input signals 105, with one of the signals defined to hold a stronger speech signal,” and generates “control signal 107 ... to activate the signal separation process only when speech is occurring.” *Id.*, 8:33-40.⁵

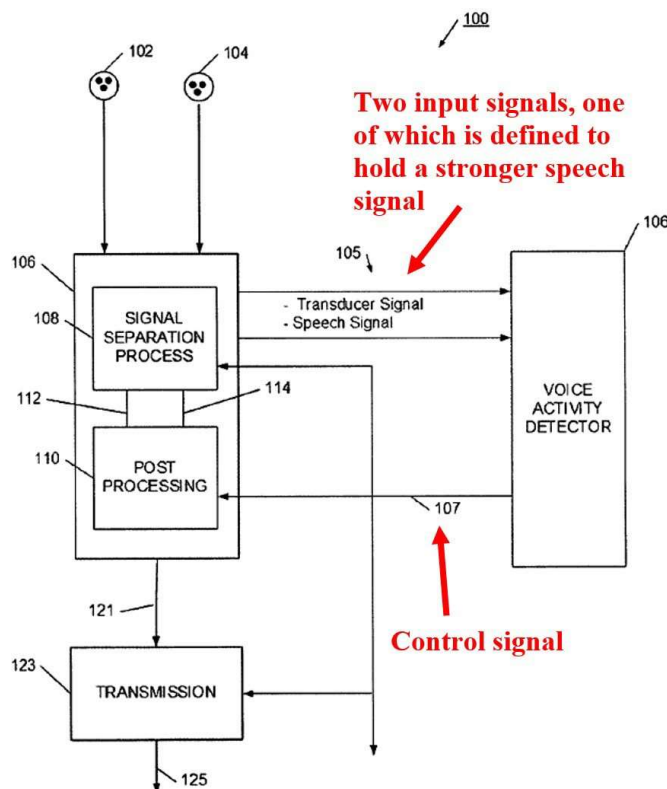


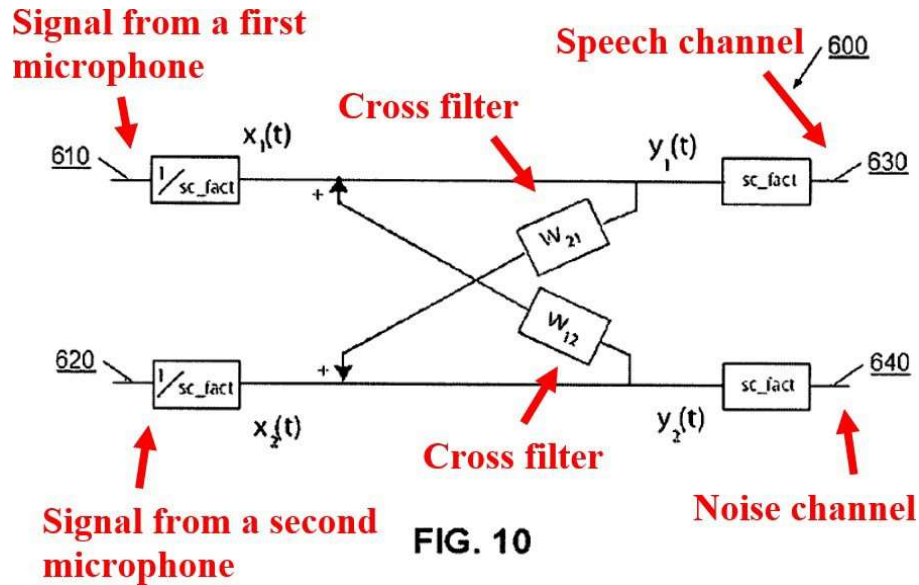
FIG. 1

⁵ All emphasis added.

Ex. 1006, FIG. 1

Visser's "signal separation process" is performed based on signals generated by an "ICA [independent component analysis] or BSS [blind signal source] processing function." *Id.*, 8:16-18, 16:3-28, 17:29-30, FIG. 10. Visser's "ICA or BSS processing function" receives "signals X_1 and X_2 ... from channels 610 and 620," each of the signals "typically ... com[ing] from at least one microphone." *Id.*, 17: 36-39, FIG. 10. Further, Visser's "ICA or BSS processing function" generates (i) "channel 630 of separated signals U_1 " (*i.e.*, "speech channel") that "contains predominantly desired signals," namely speech, and (ii) "channel 540 of separated signals U_2 " (*i.e.*, "noise channel") that "contains predominantly noise signals." *Id.*, 17:39-44.

Visser's system generates the "speech channel" by combining "input signal X_1 " and "input signal X_2 " ("input signal X_2 " having been filtered by "cross filter w_{12} "). *Id.*, FIG. 10. Further, the system generates the "noise channel" by combining "input signal X_2 " and "input signal X_1 " ("input signal X_1 " having been filtered by "cross filter w_{21} "). *Id.*



Ex. 1006, FIG. 10

These channels are provided as inputs to “VAD 106” (e.g., “input signals 105”) to determine “when speech is present.” *Id.*, 8:33-36.

As Visser describes, “cross filters W_{21} and W_{12} can have sparsely distributed coefficients over time to capture a long period of time delays.” *Id.*, 17:56-64. Further, the “speech separation process ... may be adaptive and learn according to the specific acoustic environment,” and to “adapt to particular microphone placement, the acoustic environment, or a particular user's speech.” *Id.*, 9:8-12. In Visser's “ICA process,” each filter “ha[s] an adaptable and adjustable filter coefficient.” *Id.*, 16:63-65. Specifically, “the coefficients are adjusted to improve separation performance ... and the new coefficients are applied. This continual adaptation of the filter coefficients enables the [speech

separation] process ... to provide a sufficient level of separation, even in a changing acoustic environment.” *Id.*, 16:65-17:4, FIG. 9; Ex. 1003, ¶¶68-77.

C. Combination of Avendano and Visser

A POSITA would have found it obvious to combine Avendano and Visser. Ex. 1003, ¶¶101-125. Both references come from the same field of endeavor of enhancing speech and attenuating noise based on voice activity detection. Ex. 1005, Abstract, 1:24-26, 3:13-26, 3:42-60; Ex. 1006, Abstract, 1:19-23, 6:57-7:25, 8:4-47. Further, both references describe identifying voice activity by analyzing (i) a first processed signal representative of speech detected by two physical microphones, and (ii) a second processed signal representative of noise detected by the two physical microphones. Ex. 1005, Abstract, 1:24-26, 3:13-7:9; Ex. 1006, Abstract, 1:19-23, 6:57-7:25, 8:4-47, 17:29-50. Further still, modifying Avendano’s system in view of Visser’s disclosure would have improved Avendano’s system by enabling the system to “adapt” and “learn” to separate speech “according to the specific acoustic environment,” such that the system can accurately separate speech “even in a changing acoustic environment.” Ex. 1006, 16:65-17:4, FIG. 9; Ex. 1003, ¶104.

For example, Avendano determines the presence of speech based on a comparison between (i) a first processed signal (“cardioid primary signal (C_f)” directed to the front) and (ii) a second processed signal (“cardioid secondary signal (C_b)” directed to the back). Section III(A); Ex. 1005, 4:47-7:9. This comparison is

particularly suitable for determining presence of speech, as the source of speech (“audio source 102”) is positioned on a front side of “audio device 104,” and thus “cardioid primary signal (C_f)” would have a greater response to speech than “cardioid secondary signal (C_b).” Ex. 1003, ¶¶105-106.

Specifically, as taught by Avendano, the differences in the “levels” between the two processed signals is represented by a ratio (ILD) between energy levels of these two processed signals. Ex. 1005, 5:49-6:34, 6:58-7:3. If the ratio is sufficiently high (*e.g.*, when the energy level of front-facing “cardioid primary signal C_f ” is sufficiently higher than that of back-facing “cardioid secondary signal C_b ”), this is indicative of a presence of speech activity. *Id.*, 6:58-7:3.

Further, Avendano’s “cardioid primary signal (C_f)” and “cardioid secondary signal (C_b)” are each generated by applying a delay to a signal received from one of the physical microphones, and combining the delayed signal with a signal received from the other physical microphone. Section III(A); Ex. 1005, 5:15- 35.

Visser determines the presence of speech using principles similar to those taught by Avendano. For example, Visser’s “VAD 106” determines the presence of speech based on two processed signals, “with one of the signals defined to hold a stronger speech signal.” Section III(B); Ex. 1006, 8:33-35. Specifically, Visser’s processed signals include (i) a “speech channel” that “contains predominantly

desired signals,” and (ii) a “noise channel” that “contains predominantly noise signals” (*e.g.*, generated by “ICA or BSS processing”). Ex. 1006, 17:36-44.

Like Avendano, Visser’s processed signals are generated by applying a delay to a signal received from one of the physical microphones (*e.g.*, using a cross filter having “a delay gain factor for the time delay between the output signal and the feedback input signal”), and combining the delayed signal with a signal received from the other physical microphone. *Id.*, 17:39-62, FIG. 10.

Visser additionally describes that the separation between the “speech channel” and the “noise channel” is further enhanced by generating and adapting the filters according to an “ICA process.” *Id.*, 9:20-62, 17:29-25:9, FIGS. 10-13. Such a process enables a system to adapt to “the specific acoustic environment,” including “[a] particular microphone placement, the acoustic environment, or a particular user’s speech,” in order “to provide a sufficient level of separation, even in a changing acoustic environment.” *Id.*, 9:8-12, 16:65-17:4.

Accordingly, a POSITA would have been motivated to incorporate Visser’s “ICA process” into Avendano’s system to further enhance the differences in speech content in Avendano’s two processed signals, such that speech is detected with a greater degree of accuracy, “even in a changing acoustic environment.” Ex. 1003, ¶112.

For example, a POSITA would have recognized that, like Avendano’s “delay node 404,” Visser’s “cross filter w_{12} ” is configured to filter one microphone signal, in order to generate a first processed signal representative of speech (e.g., by combining the filtered signal with another microphone signal). Ex. 1005, 5:15-35, FIG. 4a; Ex. 1006, 17:36-44, FIG. 10; Ex. 1003, ¶¶113-115.

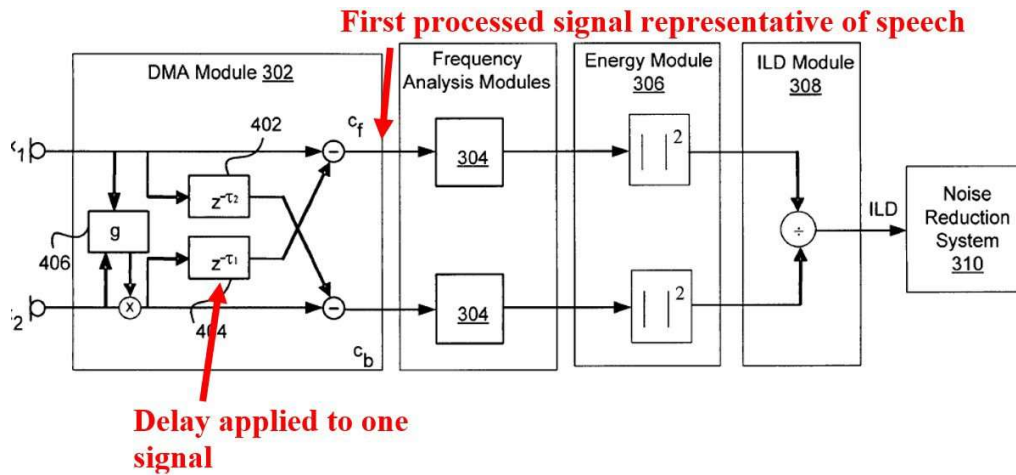


FIG. 4a

Ex. 1005, FIG. 4a

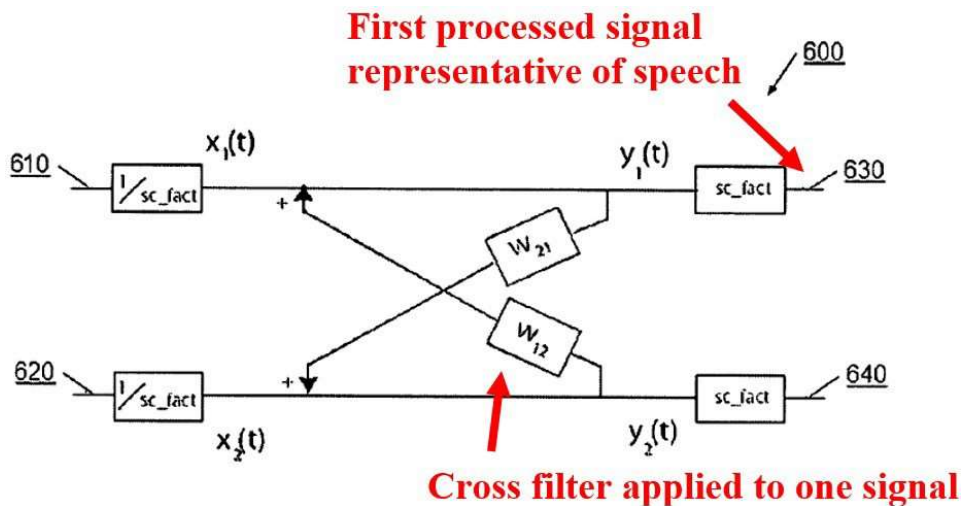


FIG. 10

Ex. 1006, FIG. 10

Similarly, a POSITA would have recognized that, like Avendano's "delay node 402," Visser's "cross filter w_{21} " is configured to filter another microphone signal, in order to generate a second processed signal representative of noise (*e.g.*, by combining the filtered signal with the other microphone signal). Ex. 1005, 5:15-35, FIG. 4a; Ex. 1006, 17:36-44, FIG. 10; Ex. 1003, ¶¶116-118.

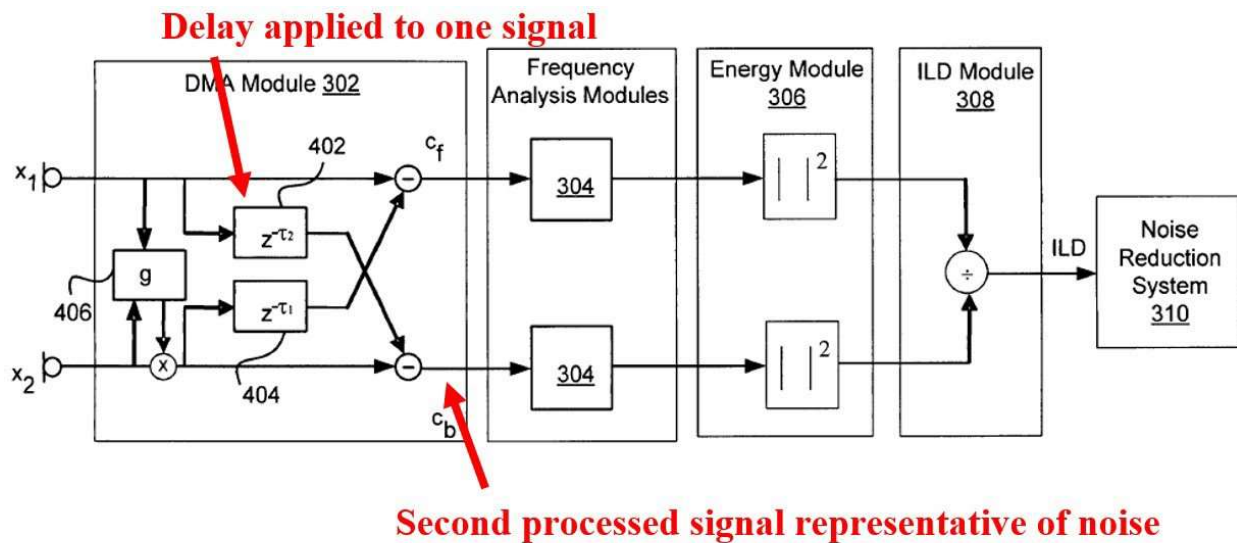
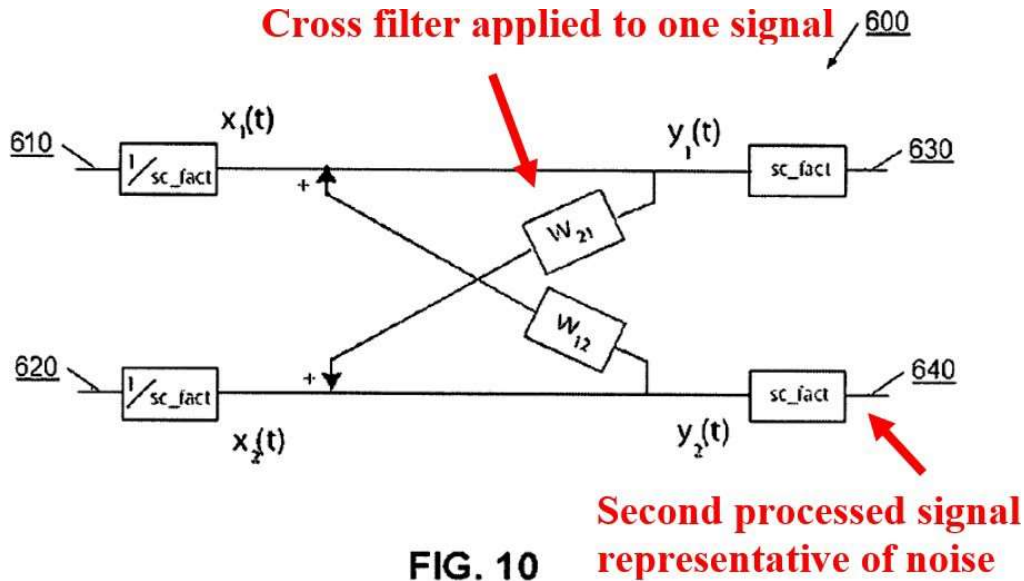


FIG. 4a

Ex. 1005, FIG. 4a



Ex. 1006, FIG. 10

From this disclosure, and as an example of applying additional filtering to Avendano's signals, a POSITA would have found it obvious to incorporate Visser's "cross filters w_{21} and w_{12} " into Avendano's system. Ex. 1003, ¶119. Such a modification would merely involve using a known technique (*e.g.*, using Visser's "cross filters") to improve a similar device in a similar way (*e.g.*, to further enhance the separation of speech and noise, such that voice activity is detected more accurately). *Id.*

In this example, Visser's cross filter w_{12} would have been applied to Avendano's second signal x_2 (e.g., to further increase speech content in the first virtual microphone C_f). *Id.*, ¶120. Further, Visser's cross filter w_{21} would have been applied to Avendano's first signal x_1 (e.g., to further decrease speech content in the second virtual microphone C_b). *Id.* An example of this modification is shown below:

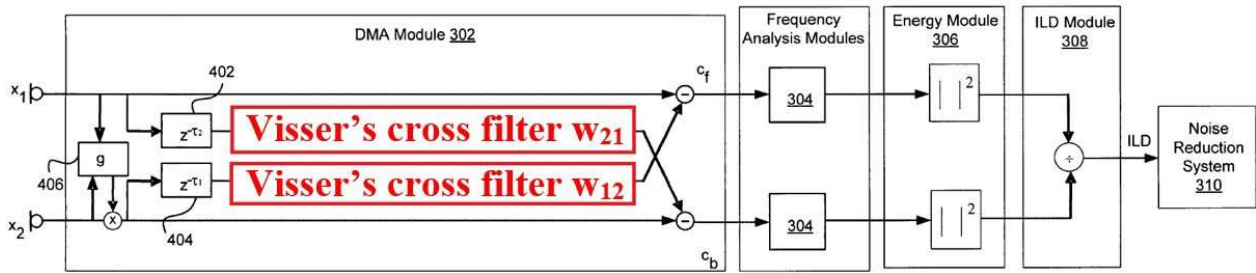


FIG. 4a

Ex. 1005, FIG. 4a (modified)

Further, a POSITA would have had a reasonable expectation of success in incorporating Visser's "cross filters" into Avendano's system. Ex. 1003, ¶¶ 121-125. For example, Avendano's "DMA Module 302" and Visser's "cross filters" are both configured to (i) receive similar types of input signals (e.g., signals from two physical microphones), and (ii) output signals for a similar purpose (e.g., to produce two processed signals in which one processed signal has a greater degree of speech than the other). Ex. 1005, 5:15-6:34, FIG. 4a; Ex. 1006, 17:36-44, FIG. 10; Ex. 1003,

¶¶121-123. Given these similarities, limited modifications would have been necessary to incorporate Visser’s “cross filters” into Avendano’s system and other operations in Avendano’s system would not have been impacted by the combination. Ex. 1003, ¶124.

Thus, the combination of Avendano and Visser would have been well within the grasp of a POSITA and a POSITA would have had a reasonable expectation of success in making the combination. *Id.*, ¶125.

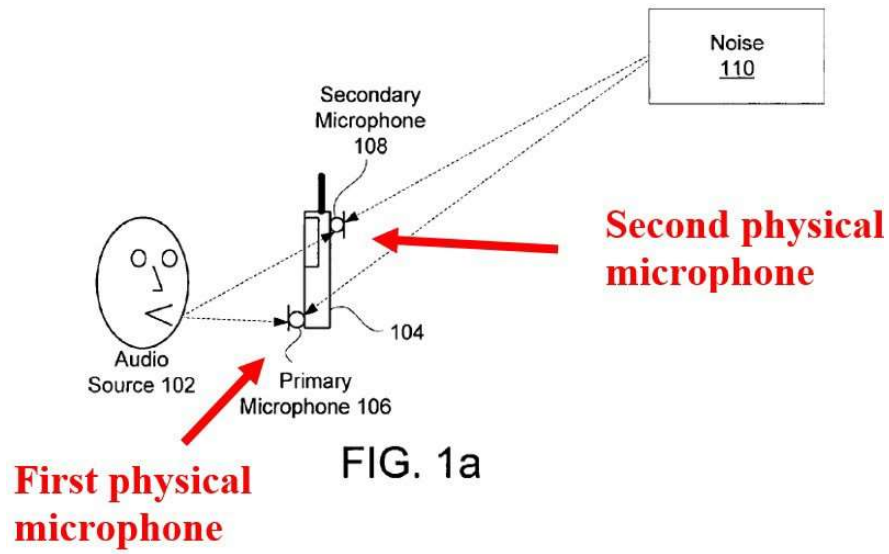
D. Claim 1

[1pre]: A method comprising:

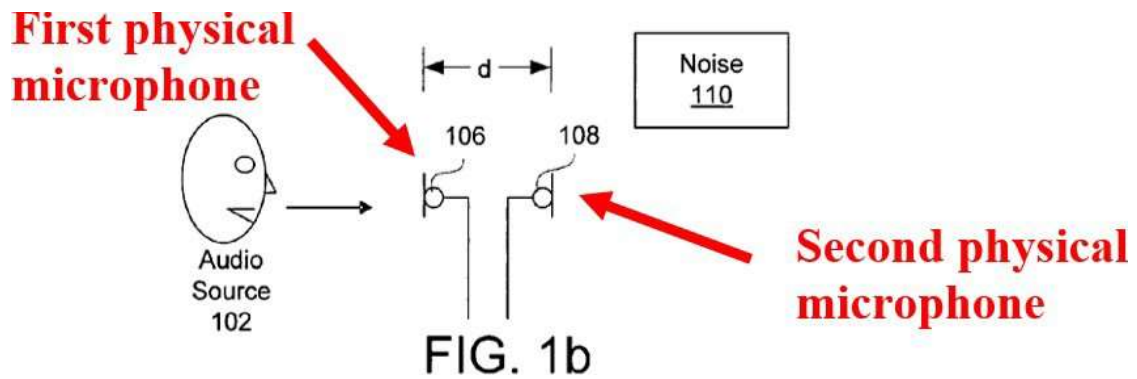
To the extent that the preamble of claim 1 is limiting, Avendano discloses a method. Ex. 1005, FIGS. 1a-7, 3:12-7:9, 9:43-10:25 (“an exemplary method for utilizing ILD of omni-direction microphones for noise suppression and speech enhancement.”); Ex. 1003, ¶126.

[1a]: forming a first virtual microphone by combining a first signal of a first physical microphone and a second signal of a second physical microphone

In Avendano, the first physical microphone is “primary microphone 106” and the second physical microphone is “secondary microphone 108.” Ex. 1005, 3:27-55, FIGS. 1a-1b; Ex. 1003, ¶127-131.



Ex. 1005, FIG. 1a



Ex. 1005, FIG. 1b

Further, in Avendano, the first signal is “primary acoustic signal (X_1)” and the second signal is “secondary acoustic signal (X_2).” *Id.* 4:20-27, FIGS. 4a-4b. Further, the first virtual microphone is “cardioid primary signal (C_f),” which is formed as a combination of “primary acoustic signal” and “secondary acoustic signal.” *Id.* 4:27-6:10, FIGS. 4a-4b.

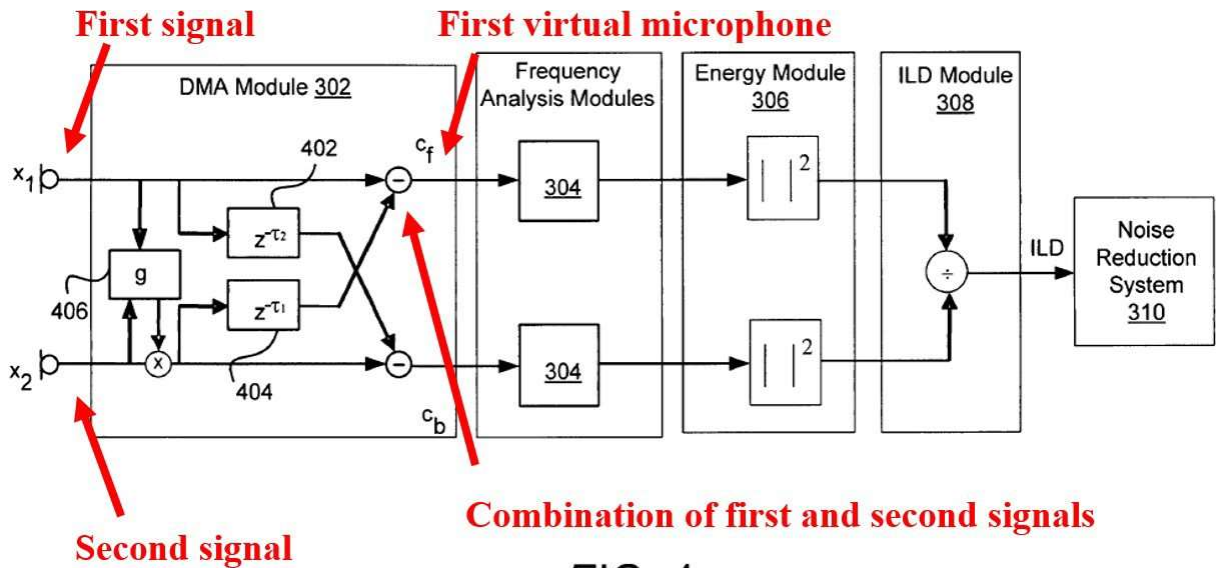


FIG. 4a

Ex. 1005, FIG. 4a

In exemplary embodiments, a cardioid primary signal (C_f) is mathematically determined in the frequency domain (Z transform) as

$$C_f = X_1 - z^{-t1} g X_2$$

First virtual microphone
Second signal
First signal

Ex. 1005, 5:25-30

[1b]: forming a filter that describes a relationship for speech between the first physical microphone and the second physical microphone, and

In Avendano, the filter is “delay node 402,” which describes a relationship for speech between the first physical microphone and the second physical microphone. Ex. 1005, 4:28-34, 5:19-35, FIGS. 4a-4b; Ex. 1003, ¶133-138.

In Avendano, “due to a space difference between the microphones,” and because the speed of sound propagation in air is widely known to be approximately 330 m/s, there will be a “difference in times of arrival of the signal from a speech source to the microphones.” Ex. 1005, 1:33-36; Ex. 1003, ¶135. Specifically, in Avendano, a first physical microphone is positioned closer to a source of speech (“audio source 102”), and a second physical microphone is positioned farther from the source of speech. Section III(D), [1a]; Ex. 1005, 3:45-49, FIGS. 1a-1b. Accordingly, the “relationship for speech” between Avendano’s two physical microphones is, at least in part, that one microphone will receive speech prior to the other microphone. Ex. 1003, ¶135. Avendano’s “delay node 402” delays the first signal, which would otherwise include speech content before the second signal. Ex. 1005, 5:15-35. Accordingly, Avendano’s “delay node 402” describes, at least in part, a temporal relationship for speech between the first and second signals. Ex. 1003, ¶136.

Further, Avendano's delay nodes are implemented using filters. Ex. 1005, 8-38-51 (“[t]o implement a fractional delay, allpass filters 416 and 418 ... are applied to the signals”).

This also is consistent with the '611 patent, which describes that a time delay can be implemented using a “delay filter.” Ex. 1001, FIG. 3 (“Delay filter $z^{-\gamma}$ ”).

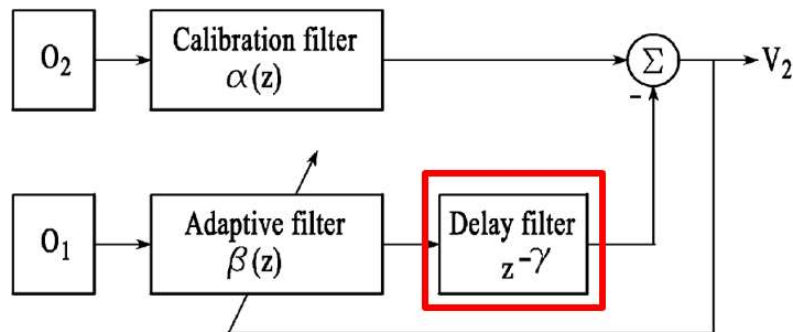


FIG.3

Ex. 1001, FIG. 3

To the extent that Avendano alone does not describe each of the features of [1b], it would have been obvious to modify Avendano's system in view of Visser to include these features. Ex. 1003, ¶¶139-144.

For example, it would have been obvious to incorporate Visser's “ICA or BSS processing function” into Avendano's system. Section III(C). According to this modification, Avendano's system would include Visser's cross filters w_{12} and w_{21} ,

each of which describes a relationship for speech between first and second physical microphones. Ex. 1003, ¶¶140-141.

According to one example modification, Visser's cross filter w_{12} would have been applied to Avendano's second signal x_2 (e.g., to further increase speech content in the first virtual microphone C_f), and Visser's cross filter w_{21} would have been applied to Avendano's first signal x_1 (e.g., to further decrease speech content in the second virtual microphone C_b). *Id.*

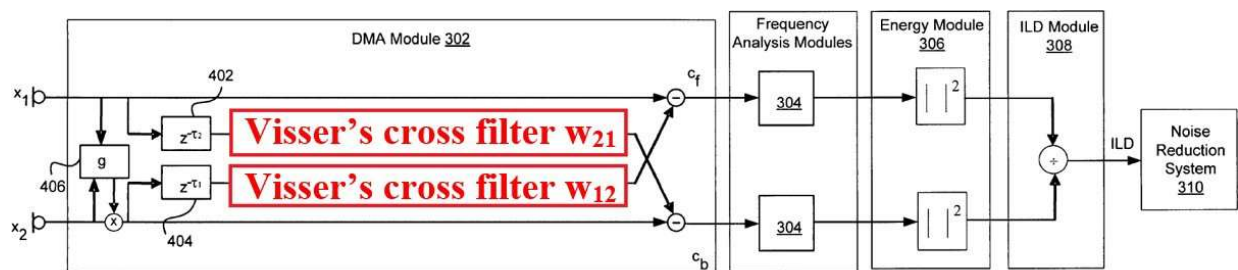


FIG. 4a

Ex. 1005, FIG. 4a (modified)

Further, Visser discloses that different physical microphones may have different responses to speech, depending on the location of the audio source relative to each physical microphone. Ex. 1006, 15:43-16:2. Visser's system leverages the different responses of the physical microphones to speech by generating cross filters w_{12} and w_{21} (e.g., using independent component analysis (ICA)) to form "speech channel 630" (which contains predominately speech) and "noise channel 640" (which contains predominately noise). *Id.*, 8:16-18, 16:3-28, 17:29-44, FIG. 10.

According to an “ICA process,” Visser’s filters are “adapt[ed] during operation,” such that the filters better separate speech from noise. *Id.*, 16:8- 11, 21:17-18:44, FIG. 12. Specifically, the “speech separation process ... may be adaptive and learn according to the specific acoustic environment,” adapting “to particular microphone placement, the acoustic environment, or a particular user's speech.” *Id.*, 9:8-12.

Accordingly, each of Visser’s cross filters w_{12} and w_{21} describes a relationship for speech between at least a first physical microphone and a second physical microphone. Ex. 1003, ¶144.

[1c]: forming a second virtual microphone by applying the filter to the first signal to generate a first intermediate signal, and summing the first intermediate signal and the second signal

In Avendano, the first intermediate signal is an output of “delay node 402,” which is generated by applying “delay node 402” (*i.e.*, a filter) to the first signal. Ex. 1005, 5:15-35, FIGS. 4a-4b; Ex. 1003, ¶¶145-152.

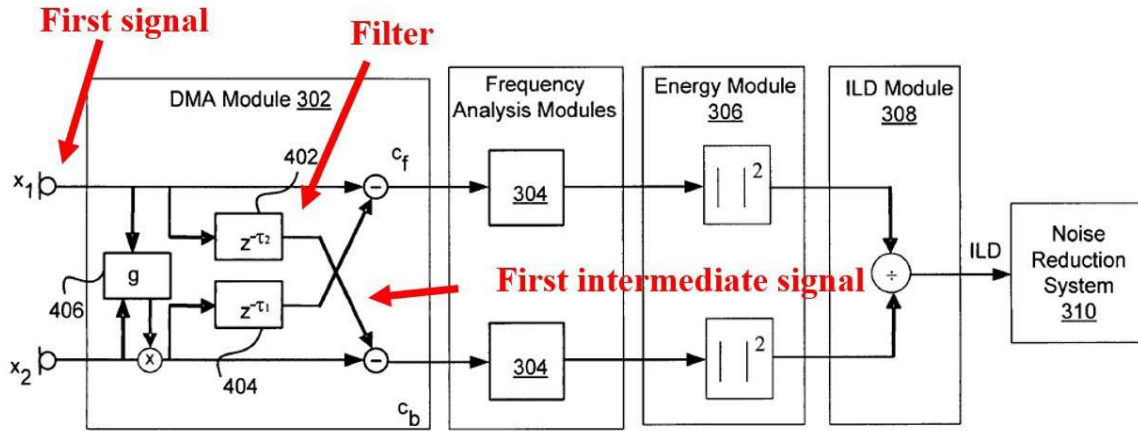


FIG. 4a

Ex. 1005, FIG. 4a

Further, in Avendano, a second virtual microphone is “cardioid secondary signal (C_b),” which is formed by subtracting the first intermediate signal from the second signal. *Id.*

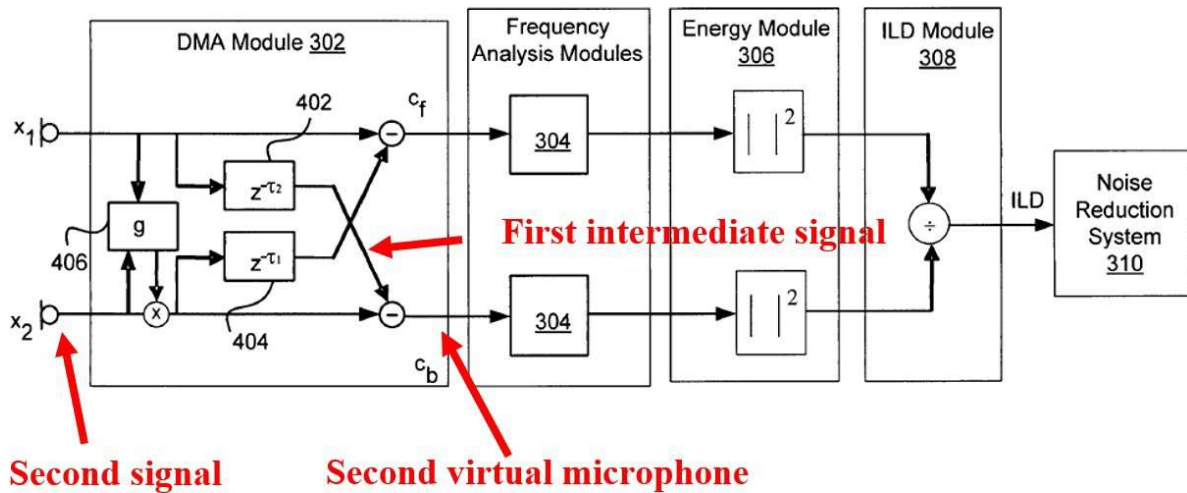


FIG. 4a

Ex. 1005, FIG. 4a

Avendano's FIG. 4b implements this subtraction by summing (i) the second signal x_2 , and (ii) the inverse of the first intermediate signal (represented using a summation node and a negative sign by a node input). Ex. 1003, ¶149.

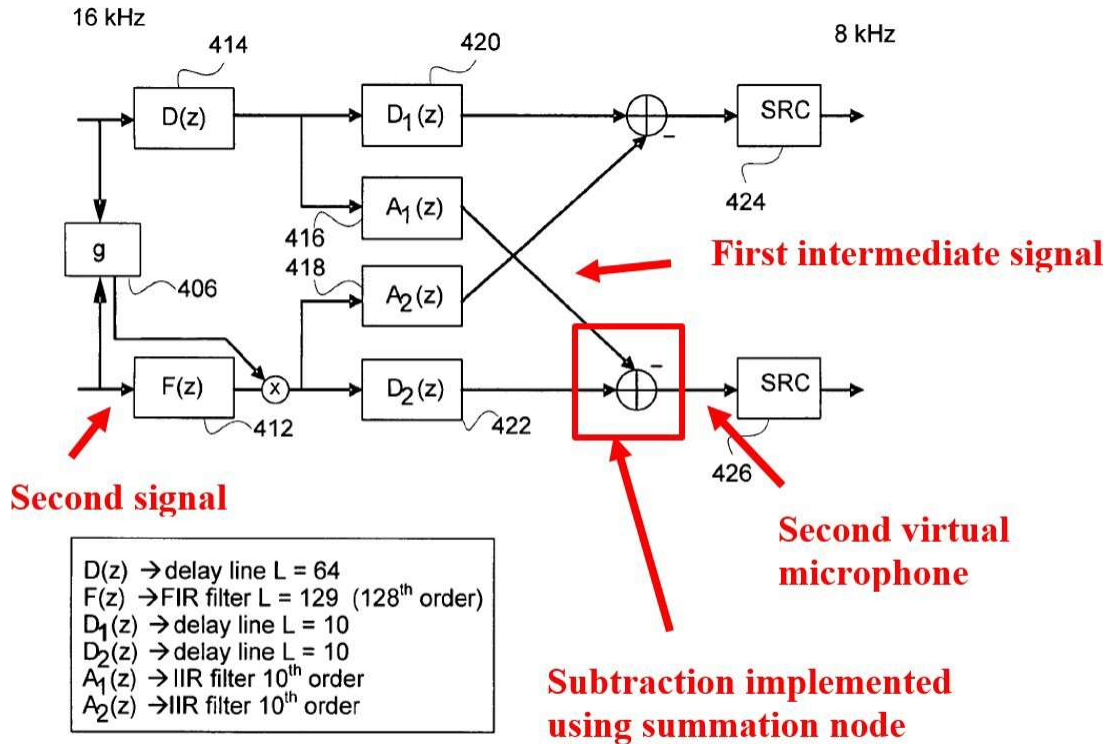


FIG. 4b
Ex. 1005, FIG. 4b

This is consistent with the '611 patent, which describes that “processing paths” are “summed to form virtual microphones,” and that “varying the magnitude and sign of the delays and gains of the processing paths leads to a wide variety of virtual microphones.” Ex. 1001, 21:26-52, FIGS. 25-26; Ex. 1003, ¶150.

Further, FIG. 3 of the '611 patent indicates that a summation (signified by “ Σ ”) includes summing a first value with the inverse of a second value (signified using a negative sign by an input to “ Σ ”) to form a second virtual microphone V_2 (*i.e.*, subtracting the second value from the first value). Ex. 1003, ¶151.

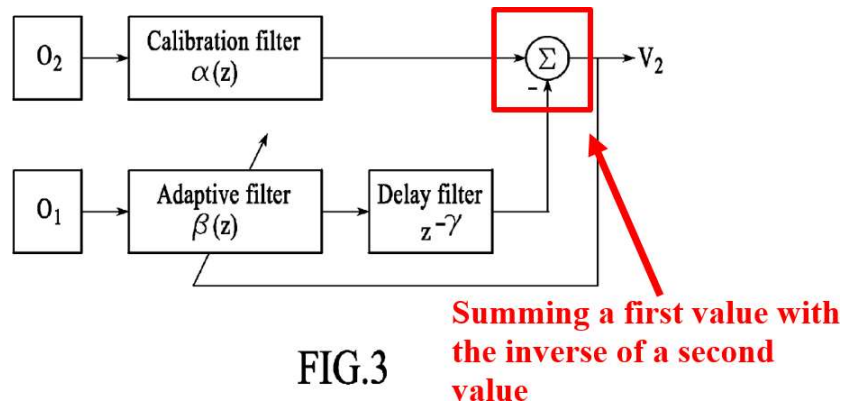


FIG.3

Ex. 1001, FIG. 3

Accordingly, as in the '611 patent, Avendano forms a second virtual microphone by applying the filter to the first signal to generate a first intermediate signal, and summing the first intermediate signal and the second signal. Ex. 1003, ¶152.

To the extent that Avendano alone does not describe each of the features in [1c], it would have been obvious to modify Avendano's system in view of Visser to include these features. *Id.*, ¶¶153-162.

As discussed above, it would have been obvious to incorporate Visser’s “ICA or BSS processing function” into Avendano’s system to further enhance the separation between speech and noise. Section III(C). According to one example modification, Visser’s cross filter w_{12} would have been applied to Avendano’s second signal x_2 (e.g., to further increase speech content in the first virtual microphone C_f), and Visser’s cross filter w_{21} would have been applied to Avendano’s first signal x_1 (e.g., to further decrease speech content in the second virtual microphone C_b). *Id.*

According to the modification, the first intermediate signal is an output of “cross filter w_{21} ,” which is generated by applying “cross filter w_{21} ” to a first signal x_1 . *Id.*

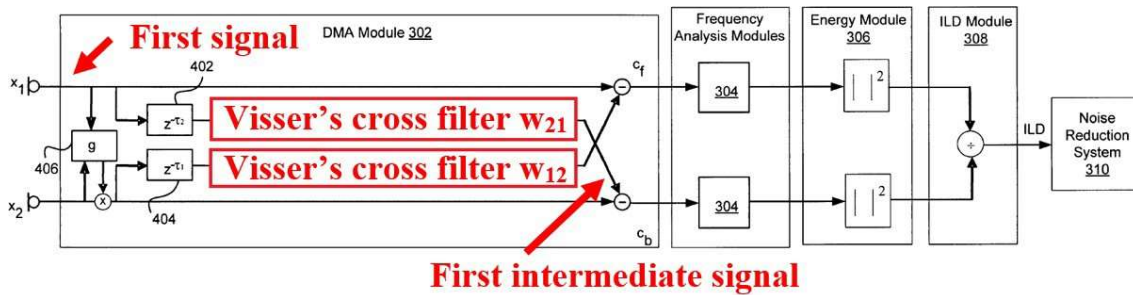
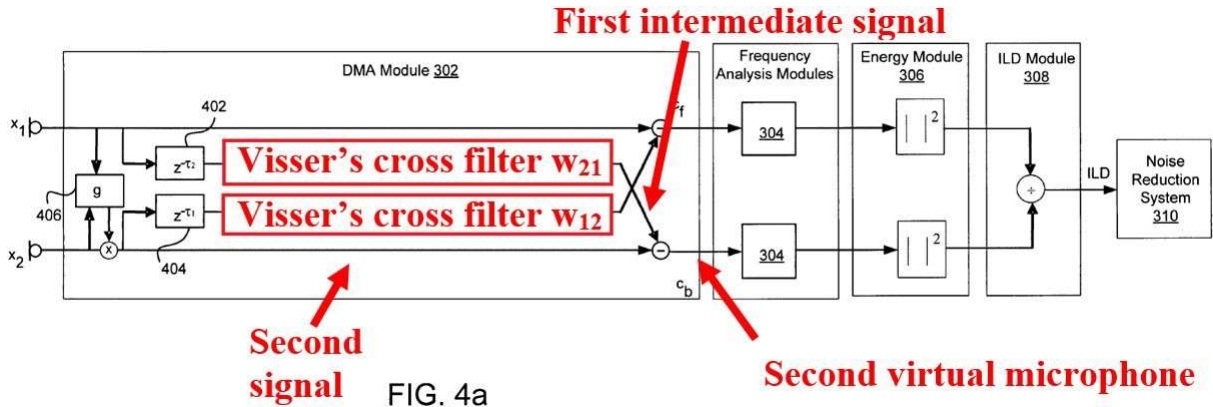


FIG. 4a

Ex. 1005, FIG. 4a (modified)

Further, according to the modification, a second virtual microphone is “cardioid secondary signal (C_b),” which is formed by subtracting the first intermediate signal from a second signal x_2 . Ex. 1005, FIG. 4a.



Ex. 1005, FIG. 4a (modified)

[1d]: generating an energy ratio of energies of the first virtual microphone and the second virtual microphone

In Avendano, the energy of the first virtual microphone is E_f , which is calculated by integrating the first virtual microphone over time. Ex. 1005, 5:49-61; Ex. 1003, ¶¶163-166.

$$E_f(t, \omega) = \int_{frame} |C_f(t', \omega)|^2 dt',$$

↑
↑

Energy of the first virtual microphone
First virtual microphone

Ex. 1005, 5:61

The energy of the second virtual microphone is E_b , which is calculated by integrating the second virtual microphone over time. *Id.*, 5:49-6:5.

$$E_b(t, \omega) = \int_{frame} |C_b(t', \omega)|^2 dt'.$$

Energy of the second virtual microphone **Second virtual microphone**

Ex. 1005, 6:5

Further, the energy ratio is “Inter-Level Difference” (ILD), which is determined by dividing the energy of the first virtual microphone by the energy of the second virtual microphone. *Id.*, 6:8-34.

$$ILD(t, \omega) = \frac{\int |C_f(t', \omega)|^2 dt'}{\int_{frame} |C_b(t', \omega)|^2 dt'}.$$

Energy of the first virtual microphone **Energy of the second virtual microphone**

Ex. 1005, 6:16

[1e]: detecting acoustic voice activity of a speaker when the energy ratio is greater than a threshold value

Avendano compares the energy ratio (ILD) to a “threshold” to determine the presence or absence of speech. Ex. 1005, 6:35-7:9; Ex. 1003, ¶¶167-171.

$$\lambda_I(t, \omega) = \begin{cases} \approx 0 & \text{if } ILD(t, \omega) < \text{threshold} \\ \approx 1 & \text{if } ILD(t, \omega) > \text{threshold} \end{cases}$$

Ex. 1005, 6:61

If the ratio (ILD) is “smaller than a threshold value (e.g., threshold=0.5) above which speech is expected to be,” a value λ_1 is set to zero (e.g., indicating an absence of speech). *Id.*, 6:58-7:3. However, if the ratio (ILD) “starts to rise (e.g., because speech is present within the large ILD region), λ_1 increases” (e.g., set to one, indicating a presence of speech). *Id.*

E. Claim 2

[2]: wherein the first virtual microphone and the second virtual microphone are distinct virtual directional microphones.

Avendano discloses a first virtual microphone (“cardioid primary signal (C_f)”) and a second virtual microphone (“cardioid secondary signal (C_b)”). Section III(D), [1a], [1c]. The first and second virtual microphones are distinct from one another and represent virtual directional microphones (e.g., directed to “two different directional patterns about the audio device”). Ex. 1005, FIG. 4a, 4:27-52; 5:15-35; Ex. 1003, ¶¶172-175.

F. Claims 3 and 4

[3]: wherein the first virtual microphone and the second virtual microphone have approximately similar responses to noise.

[4]: wherein the first virtual microphone and the second virtual microphone have approximately dissimilar responses to speech.

Avendano's first virtual microphone has a directional pattern directed to the front of "audio device 104," and the second virtual microphone has a directional pattern directed to the back of "audio device 104." Section III(D), [1a], [1c]; Ex. 1005, 4:47-52, 5:25-35, FIG. 6; Ex. 1003, ¶¶176-187.

Further, in Avendano, noise is typically generated in the "background" or "far field," rather than from a location near the audio device, and "may include reverberations and echoes." Ex. 1005, 2:49, 3:35-41, FIG. 1a.

A POSITA would have recognized that, according to Avendano's configuration, Avendano's virtual microphones would have approximately similar responses to noise. Ex. 1003, ¶¶179-182.

For example, although Avendano describes two virtual microphones having respective directional patterns, neither of these virtual microphones is directed to a "background" or "far field." Ex. 1005, 2:49, FIG. 1a. Accordingly, neither virtual microphone would be more sensitive to noise generated in the "background" or "far field" than the other. Ex. 1003, ¶180.

Further, according to Avendano's configuration, an audio device would not merely detect noise propagating directly from a noise source to the audio device, but would also detect reverberations and echoes of that noise from multiple other locations relative to the audio device. Ex. 1005, 3:35-41. Due to the non-directionally specific nature of these reverberations and echoes, neither virtual microphone would be more sensitive to detecting noise (and the reverberations and echoes thereof) than the other. Ex. 1003, ¶181.

A POSITA also would have recognized that, according to Avendano's configuration, Avendano's virtual microphones would have approximately dissimilar responses to speech. *Id.*, ¶¶183-187.

In Avendano, speech is typically generated near the front of the device. Ex. 1005, 3:42-55, FIG. 1a. According to this configuration, Avendano's virtual microphones would have approximately dissimilar responses to speech (*e.g.*, first virtual microphone is directed towards the source of speech, whereas the second virtual microphone is directed away from the source of speech). Ex. 1003, ¶187.

G. Claims 5 and 6

[5]: applying a calibration to at least one of the first signal and the second signal.

[6]: wherein the calibration compensates a second response of the second physical microphone so that the second response is equivalent to a first response of the first physical microphone.

Avendano applies a calibration (“gain factor, g ”) to at least the second signal “to equalize the signal levels” of the first and second signals. Ex. 1005, 5:36- 39, 9:54-59; FIGS. 4a-4b; Section III(D), [1a]; Ex. 1003, ¶¶188-195.

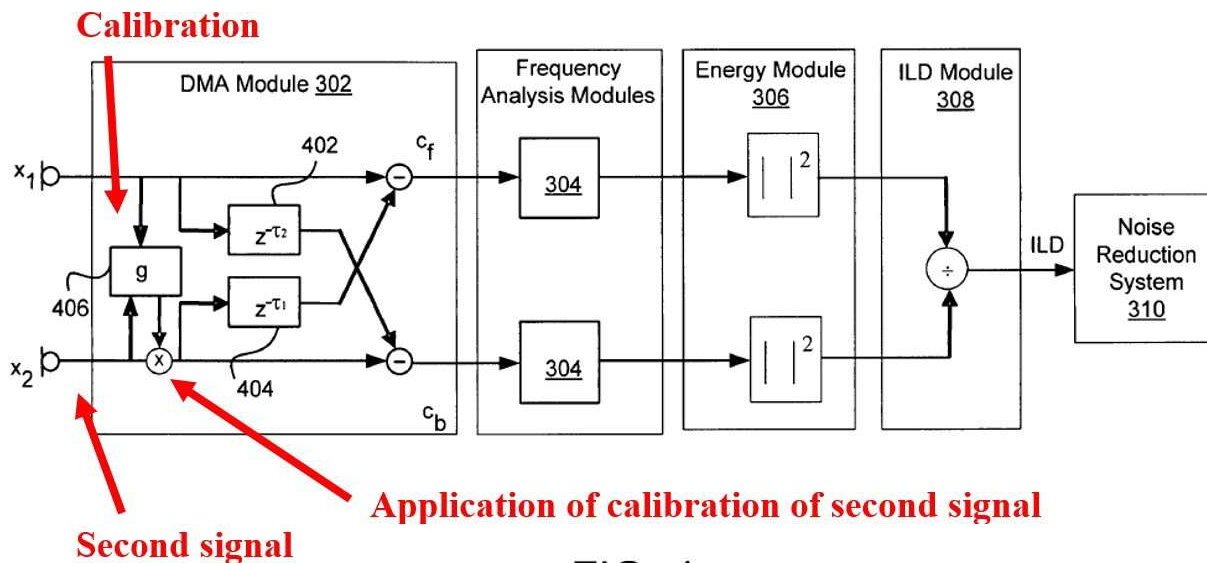


FIG. 4a

Ex. 1005, FIG. 4a

Accordingly, Avendano’s calibration compensates a second response of the second physical microphone so that the second response is equivalent to a first response of the first physical microphone. Ex. 1003, ¶195.

H. Claim 7

[7]: applying a delay to the first intermediate signal.

Avendano's first intermediate signal is an output of "delay node 402," which is generated by applying "delay node 402" (*i.e.*, a filter) to the first signal. Section III(D), [1c]; Ex. 1005, 5:15-35, FIGS. 4a-4b; Ex. 1003, ¶¶196-201.

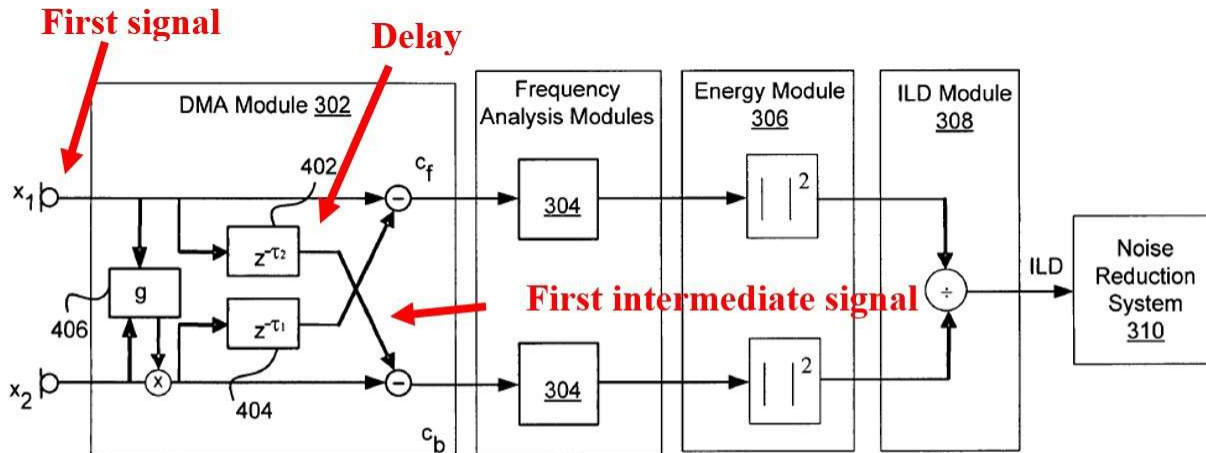


FIG. 4a

Ex. 1005, FIG. 4a

Accordingly, in Avendano, a delay has been applied to the first intermediate signal. Ex. 1003, ¶199.

Further, Visser's "cross filters w_{12} and w_{21} " can be "gain factors with only one filter coefficient per filter, for example a delay gain factor for the time delay between the output signal and the feedback input signal and an amplitude gain factor for amplifying the input signal." Ex. 1006, 17:58-62. Accordingly, in the combination of Avendano and Visser, a delay also would have been applied to the first intermediate signal by Visser's cross filter w_{21} . Ex. 1003, ¶201.

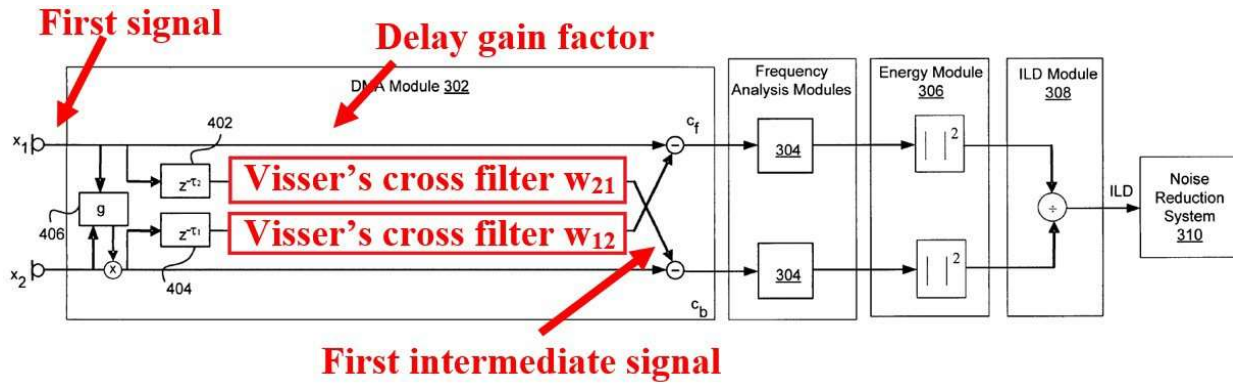


FIG. 4a

Ex. 1005, FIG. 4a (modified)

I. Claim 25

[25]: generating a vector of the energy ratio versus time.

Avendano's energy ratio (ILD) is generated as a vector over time (t). Ex. 1005, 6:8-34; Section III(D), [1d]; Ex. 1003, ¶¶202-204.

Energy ratio

$$ILD(t, \omega) = \frac{\int |C_f(t', \omega)|^2 dt'}{\int_{frame} |C_b(t', \omega)|^2 dt'}.$$

Ex. 1005, 6:16

J. Claim 26

[26]: wherein the first and second physical microphones are omnidirectional microphones.

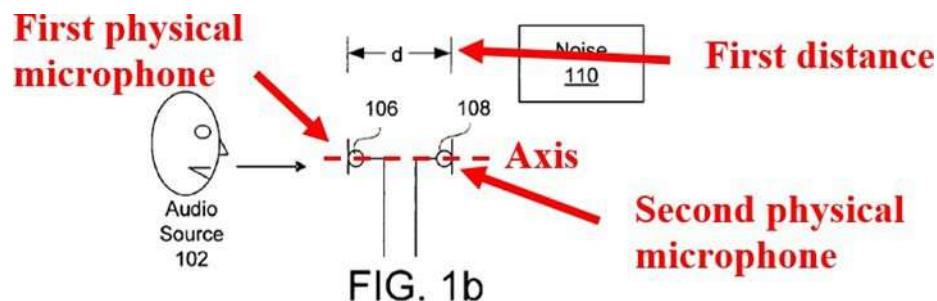
Avendano's physical microphones are "omni-directional microphone[s]." Ex. 1005, Abstract, 2:3-19, 3:33-35, 9:39-52; Ex. 1003, ¶¶205-207; Section III(D), [1a].

K. Claims 27 and 28

[27]: positioning the first physical microphone and the second physical microphone along an axis and separating the first physical microphone and the second physical microphone by a first distance.

[28]: wherein a midpoint of the axis is a second distance from a mouth of the speaker, wherein the mouth is located in a direction defined by an angle relative to the midpoint.

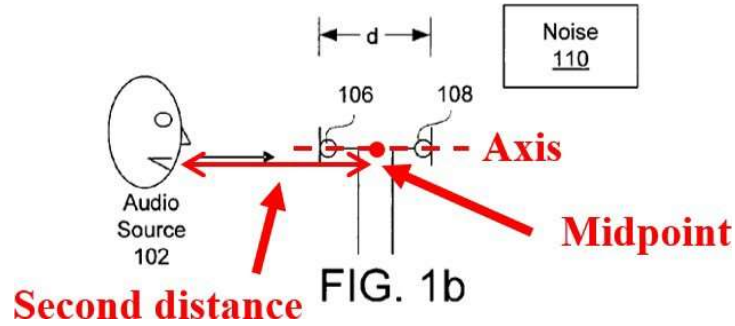
Avendano's first and second physical microphones are positioned along an axis, and separated by a first distance (d). Ex. 1005, FIG. 1b; Section III(D), [1a]; Ex. 1003, ¶¶208-214.



Ex. 1005, FIG. 1b

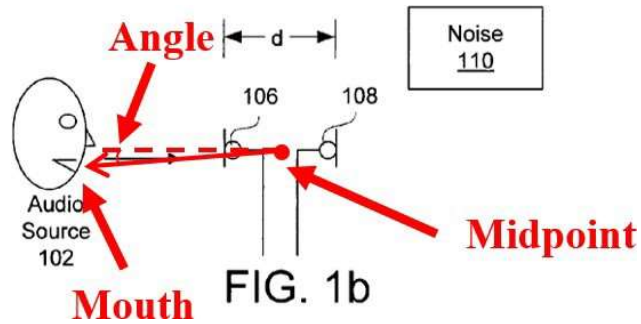
The midpoint of the axis is a second distance away from a mouth of a speaker.

Id.



Ex. 1005, FIG. 1b

The mouth of the speaker is located in a direction defined by an angle relative to the midpoint. *Id.*



Ex. 1005, FIG. 1b

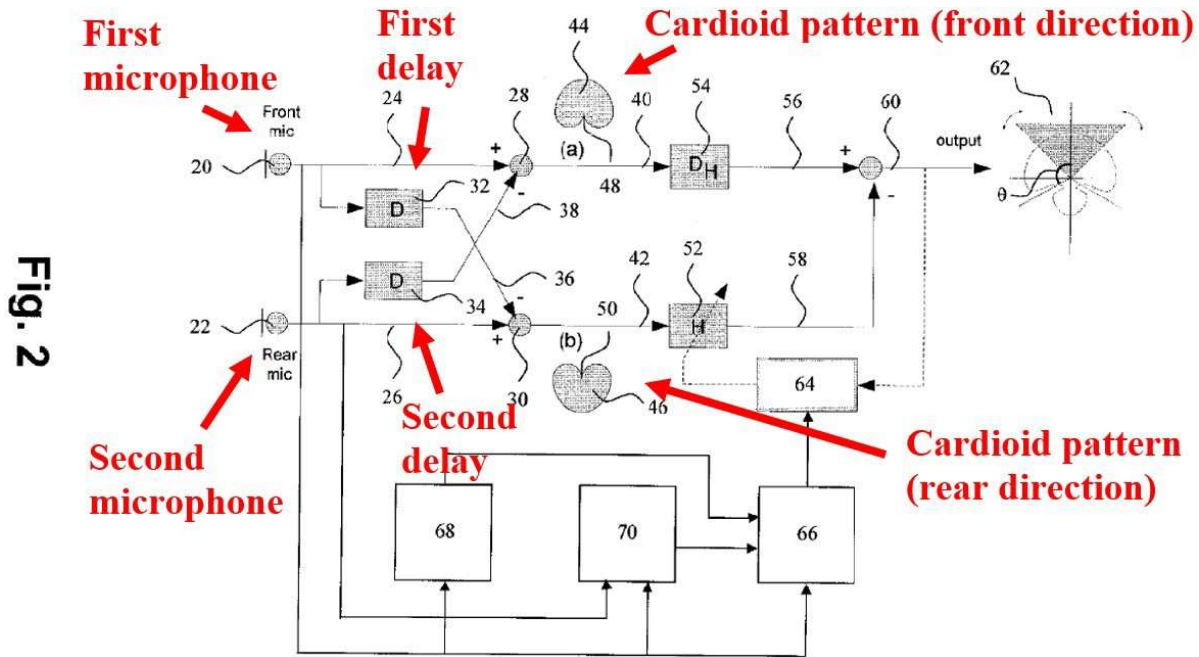
IV. GROUND 2: Avendano, Visser, And Bisgaard (Claims 8-16, 23, 24)

A. Bisgaard Overview

Bisgaard describes “a hearing instrument, such as a hearing aid, an implantable hearing prosthesis, a head set, a mobile phone, etc., with a signal processor for signal processing.” Ex. 1010, [0002]; Ex. 1011, 1:3-5.

Bisgaard’s system obtains signals from two “microphones 20, 22,” and processes the signals to obtain (i) “cardioid pattern 44” pointing towards a front of the device and (ii) “cardioid pattern 46” pointing towards a rear of the device.” Ex. 1010, [0041]-[0042], FIG. 2; Ex. 1011, 5:18-31, FIG. 2; Ex. 1003, ¶79.

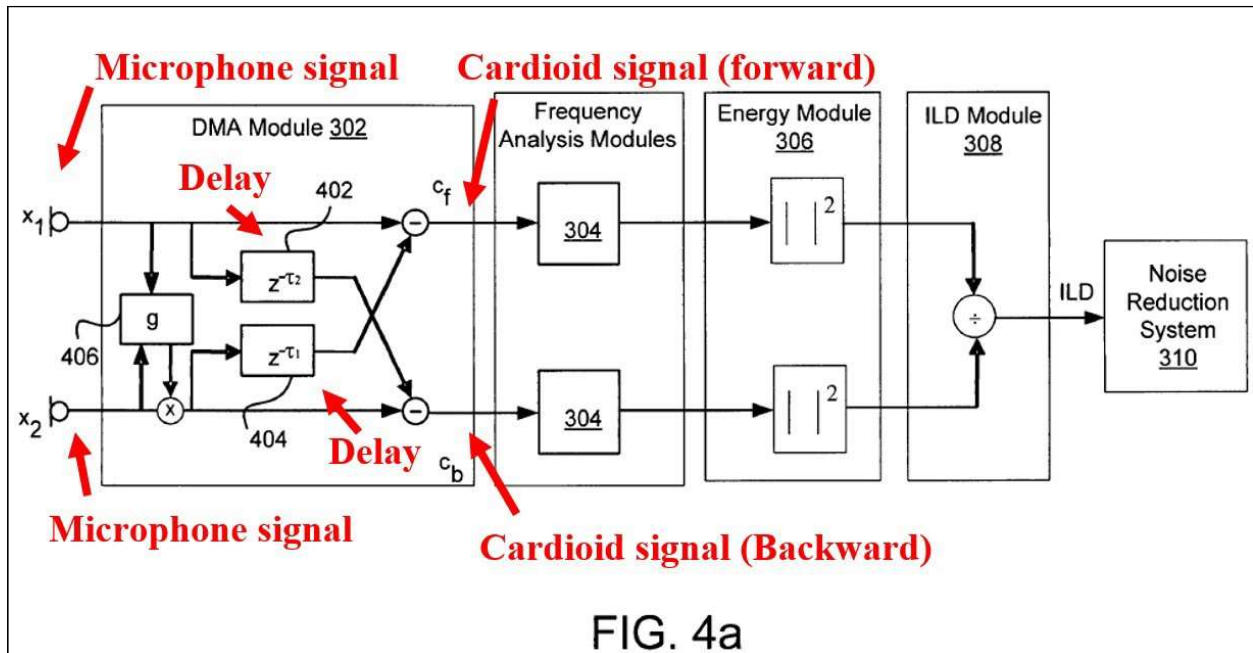
Further, Bisgaard includes “delay[s] 32, 34,” which “delay[] the digitized sound signal [received from microphones 20, 22] by the amount of time used by a sound signal to propagate in the 0° azimuth direction from the front microphone 20 to the rear microphone 22.” Ex. 1010, [0042], FIG 2; Ex. 1011, 5:21-24, FIG. 2; Ex. 1003, ¶¶78-81.



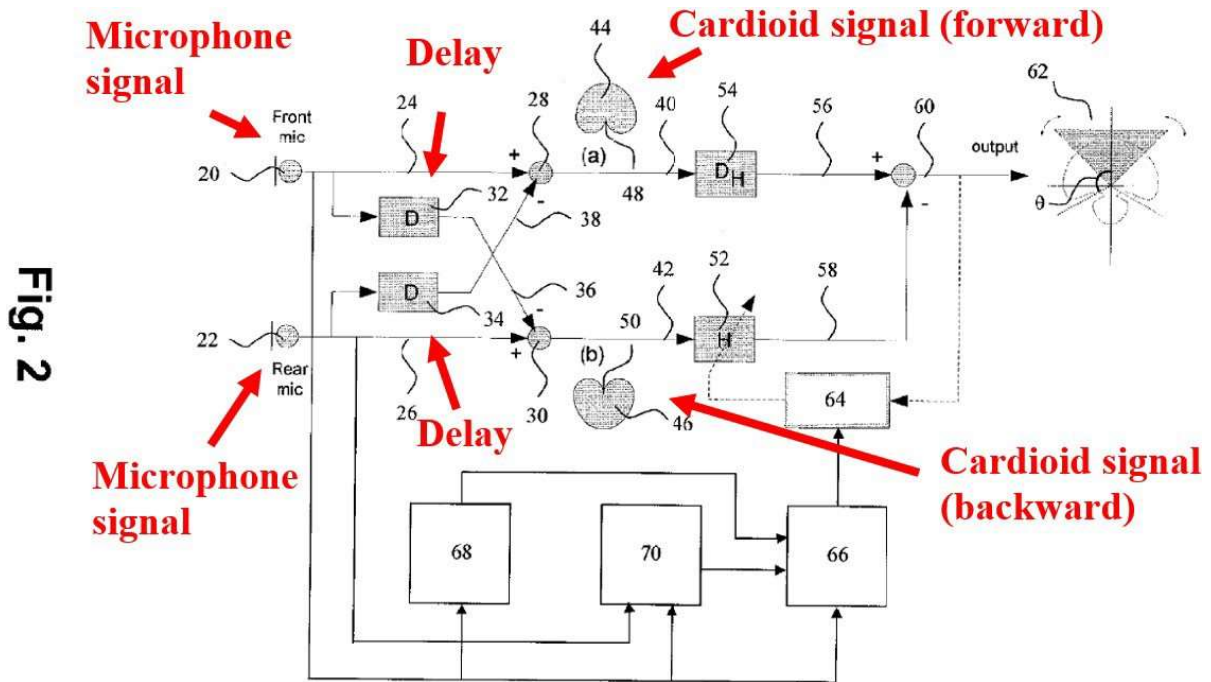
Ex. 1010, FIG. 2

B. Combination of Avendano, Visser, and Bisgaard

A POSITA would have found it obvious to combine Avendano and Visser with Bisgaard. Ex. 1003, ¶¶215-222. Like Avendano, Bisgaard receives signals from two microphones, and processes the signals to obtain (i) a cardioid signal directed in a forward direction (*e.g.*, towards a user), and (ii) another cardioid signal directed in a backwards direction (*e.g.*, away from the user). Ex. 1005, 4:42-5:35, FIGS. 4a-4b, 6; Ex. 1010, [0041]-[0042], FIG. 2; Ex. 1011, 5:18-42, FIG. 2. Further, like Avendano, Bisgaard forms each of the cardioid signals by delaying the signal from one microphone, and combining it with the signal from the other microphone. *Id.*



Ex. 1005, FIG. 4a



Ex. 1010, FIG. 2

A POSITA would have understood that Avendano's "delay nodes" and Bisgaard's "delays" serve the same purpose, namely delaying signals from physical microphones to produce directional cardioid signals (*e.g.*, by combining the delayed microphone signal with another microphone signal). Ex. 1003, ¶219. Accordingly, a POSITA would have looked to Bisgaard for details related to implementing such a delay in Avendano's system. *Id.*

From this disclosure, and as an example of applying Bisgaard's teaching regarding delays to Avendano's system, a POSITA would have found it obvious to configure each of Avendano's "delay nodes" to "delay[] the digitized sound signal" received by each of Avendano's physical microphones "by the amount of time used by a sound signal to propagate in the 0° azimuth direction from the front microphone 20 to the rear microphone 22," as taught by Bisgaard (*e.g.*, to account for the delay in the arrival of sound between the two microphones). Ex. 1010, [0042]; Ex. 1011, 5:22-24; Ex. 1003, ¶220.

Further, given the similarities between Avendano's "delay nodes" and Bisgaard's "delays," a POSITA would have had a reasonable expectation of success in incorporating Bisgaard's "delays" in Avendano's system. Ex. 1003, ¶221. For example, as discussed above, Avendano's "delay nodes" and Bisgaard's "delays" are configured to receive similar input signals and generate similar outputs. Ex. 1005, 4:42-5:35, FIGS. 4a-4b, 6; Ex. 1010, [0041]-[0042], FIG. 2; Ex. 1011, 5:18-

42, FIG. 2. Given these similarities, limited modifications would have been necessary to incorporate Bisgaard's "delays" into Avendano's system and other operations in Avendano's system would not have been impacted by the combination. Ex. 1003, ¶221.

Thus, the combination of Avendano, Visser, and Bisgaard would have been well within the grasp of a POSITA and a POSITA would have had a reasonable expectation of success in making the combination. *Id.*, ¶222.

C. Claim 8

[8]: wherein the delay is proportional to a time difference between arrival of the speech at the second physical microphone and arrival of the speech at the first physical microphone.

As Avendano describes, "due to a space difference between the microphones, the difference in times of arrival of the signals from a speech source to the microphones may be utilized to localize the speech source." Ex. 1005, 1:33-36. Further, "embodiments may use a combination of energy level differences and time delays to discriminate speech." *Id.*, 3:56-58. Namely, Avendano delays a first signal x_1 by "delay node 402," and delays a second signal x_2 by "delay node 404." Section III(A); Ex. 1005, 15:35, FIG. 4a; Ex. 1003, ¶¶223-230.

A POSITA would have recognized that, to account for the "space difference between the microphones," the delays would need to be proportional to a time

difference between the arrivals of the speech at the first and second physical microphones. Ex. 1003, ¶225. For example, a POSITA would have recognized that, if the “space difference between the microphones” were to increase, the delay would likewise proportionally increase to account for the increase in propagation time, and vice versa. *Id.*

To the extent that Avendano alone does not describe the features of [8], it would have been obvious to modify Avendano’s system in view of Bisgaard to include these features. *Id.*, ¶¶226-230. A POSITA would have found it obvious to combine Avendano and Visser with Bisgaard (*e.g.*, by incorporating Bisgaard’s “delays” into Avendano’s system). Section IV(B).

Bisgaard’s “delays” are used to account for “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone.” Ex. 1010, [0042]; Ex. 1011, 5:22-24.

A POSITA would have recognized that, to account for “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone,” the sound signal would be delayed by an amount of time that is proportional to a time difference between arrival of the speech at one physical microphone and arrival of the speech at another physical microphone. Ex. 1003, ¶229-230.

For example, a POSITA would have recognized that, if “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone” were to increase, the delay would likewise proportionally increase to account for the increase in propagation time, and vice versa. Ex. 1010, [0042]; Ex. 1011, 5:22-24; Ex. 1003, ¶230.

D. Claim 9

[9]: wherein the forming of the first virtual microphone comprises applying the filter to the second signal.

In the combination, Visser’s cross filter w_{12} would be applied to Avendano’s second signal x_2 . Section III(C)-(D); Ex. 1003, ¶¶231-233. Further, the filtered second signal would be used to form Avendano’s first virtual microphone C_f (e.g., by combining the filtered second signal with Avendano’s first signal x_1). Id.

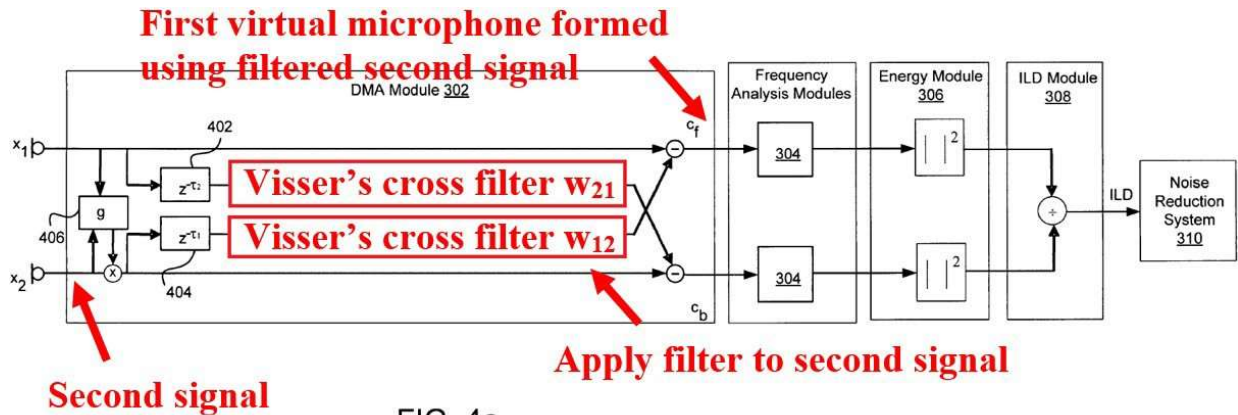


FIG. 4a
Ex. 1005, FIG. 4a (modified)

E. Claim 10

[10]: wherein the forming of the first virtual microphone comprises applying the calibration to the second signal.

Avendano applies a calibration (“gain factor, g ”) to the second signal x_2 .
Section III(G); Ex. 1003, ¶¶234-236.

Further, Avendano’s calibrated second signal x_2 is used to form the first virtual microphone C_f (e.g., by combining the calibrated second signal x_2 with the first signal x_1). Ex. 1005, 4:47-52, 5:25-35.

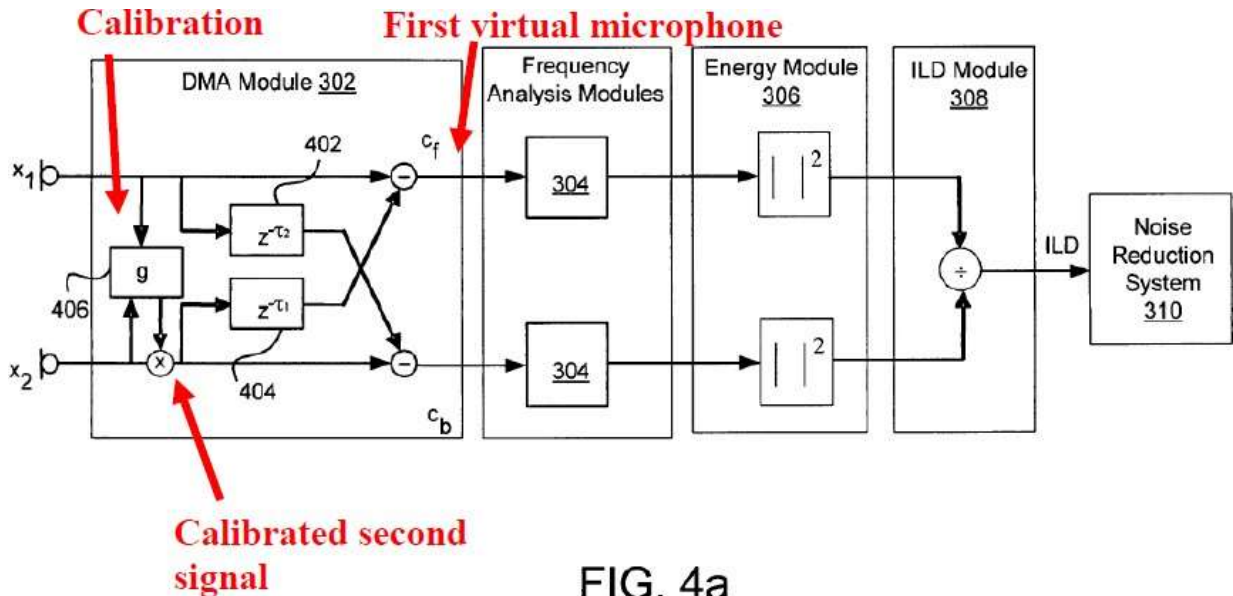


FIG. 4a
Ex. 1005, FIG. 4a

F. Claim 11

[11]: wherein the forming of the first virtual microphone comprises applying the delay to the first signal.

In Avendano's system, "DMA Module 302" forms the first virtual microphone C_f . Ex. 1005, 4:20-5:39, FIGS. 3, 4a. Further, in a "practical implementation of the DMA module 302" shown in FIG. 4b, a delay is applied to the first signal x_1 via "delay node 414." *Id.*, 8:38-51, FIG. 4b; Ex. 1003, ¶¶237-239.

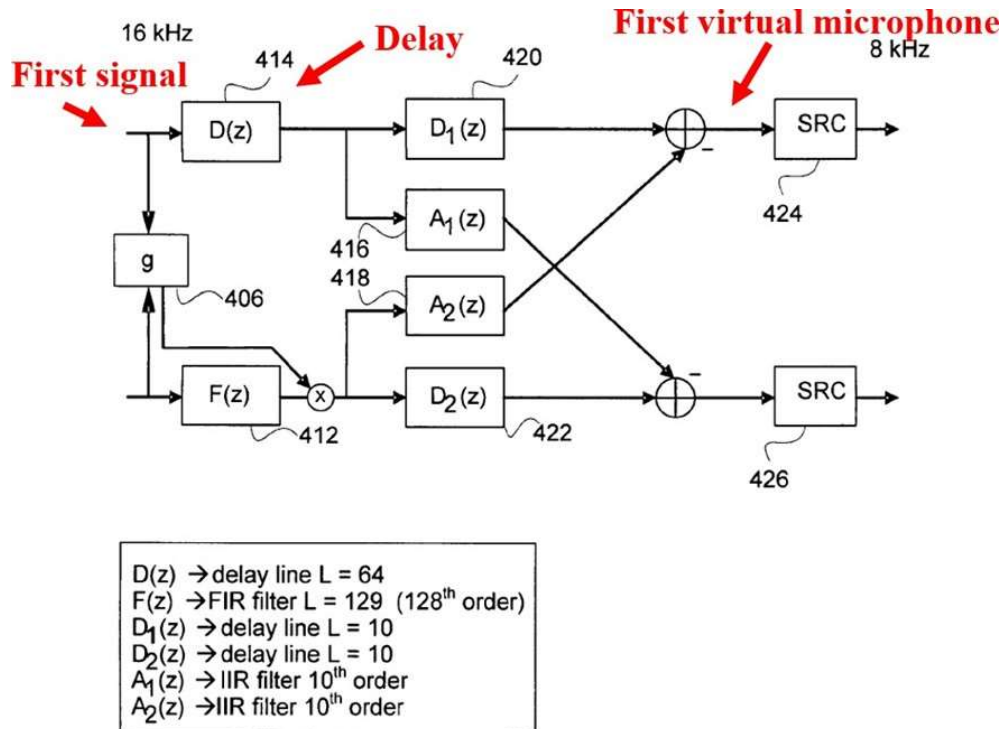


FIG. 4b

Ex. 1005, FIG. 4b

G. Claim 12

[12]: wherein the forming of the first virtual microphone by the combining comprises subtracting the second signal from the first signal.

Avendano's first virtual microphone C_f is formed by subtracting second signal x_2 from first signal x_1 . Ex. 1005, 5:15-35, FIG. 4a; Ex. 1003, ¶¶240-241.

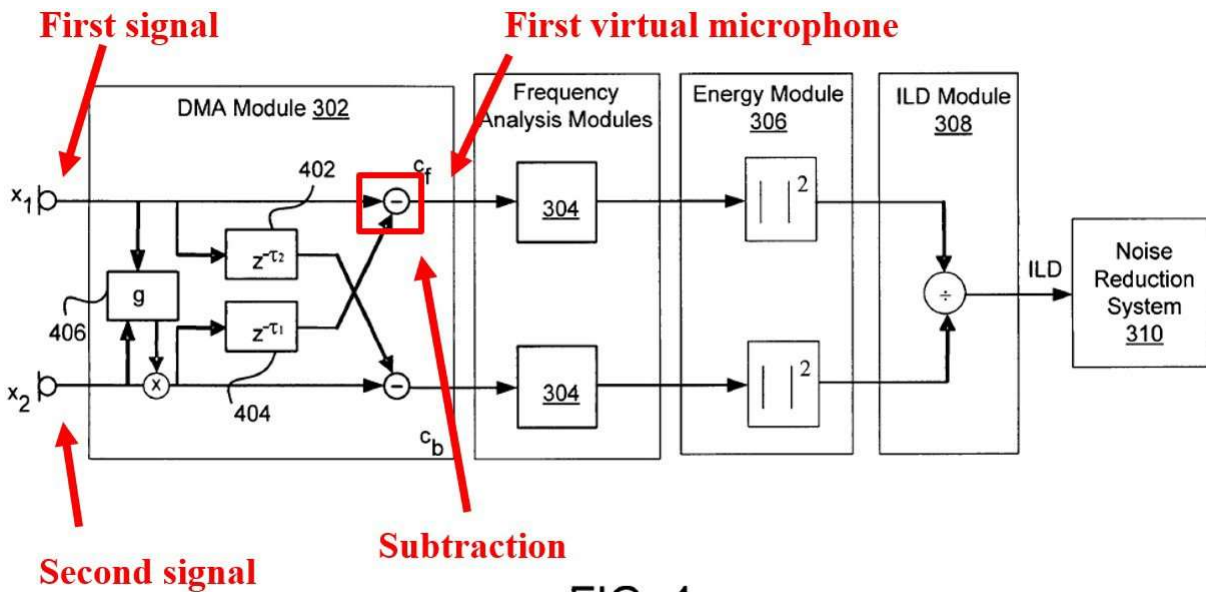


FIG. 4a

Ex. 1005, FIG. 4a

H. Claim 13

[13]: wherein the filter is an adaptive filter.

Visser's cross filters w_{12} and w_{21} are adaptive filters. Ex. 1006, 8:16-18, 16:3-28, 16:57-25:9, FIGS. 5, 7, 10-13; Ex. 1003, ¶¶242-244. For example, Visser's cross filters w_{12} and w_{21} are adapted in an "ICA or BSS processing function." Ex. 1006, 17: 29-32. An "ICA process" includes filters with "adaptable and adjustable filter

coefficient[s],” which are adapted using “learning stage 752.” *Id.*, 16:61-17:4; 21:17-58, FIG. 12.

I. Claim 14

[14]: adapting the filter to minimize a second virtual microphone output when only speech is being received by the first physical microphone and the second physical microphone.

Visser’s cross filters w_{12} and w_{21} are adapted to filter signals received from two physical microphones (“[i]nput signals X_1 and X_2 ”) to generate (i) a first virtual microphone that would “contain[] predominantly desired signals,” namely speech (“speech channel”), and (ii) a second virtual microphone that would “contain[] predominantly noise signals” (“noise channel”). Section IV(H); Ex. 1006, 17:36-44, 16:61-17:4, 21:17-58; Ex. 1003, ¶¶245-250.

Accordingly, Visser’s cross filters w_{12} and w_{21} are adapted such that when the physical microphones only receive speech (*i.e.*, without noise), the second virtual microphone—which would contain “predominantly noise signals”—would be minimized. Ex. 1003, ¶250.

J. Claim 15

[15]: wherein the adapting comprises applying a least-mean squares process.

Visser's "adaptive least mean square (NLMS) algorithm" is used to "build[] a linear filter model." Ex. 1006, 23:13-25. Based on Visser's disclosure, a POSITA would have recognized that adaptive filters (e.g., Visser's cross filters w_{12} and w_{21}) can be adapted by applying a least-mean squares process (e.g., Visser's "adaptive least mean square (NLMS) algorithm"). Ex. 1003, ¶¶251-254.

K. Claim 16

[16]: generating coefficients of the filter during a period when only speech is being received by the first physical microphone and the second physical microphone.

Visser generates coefficients of the filter during a period when only speech is being received. Ex. 1003, ¶¶255-261. For example, Visser "turn[s] off" an "ICA module" for adapting filters "when desired speech is not present, ... thereby enabling adaptation only when such adaptation will be able to achieve a separation improvement." Ex. 1006, 9:12-43. As Visser explains, this "allows the ICA process to achieve and maintain good separation quality even after prolonged periods of desired speaker silence and avoid algorithm singularities due to unfruitful separation efforts for addressing situations the ICA stage cannot solve," "adds significant robustness to the methodology," and conserves "processing and battery power." *Id.*

L. Claim 23

[23]: wherein the filter is a static filter.

Avendano's "delay node 402" describes a relationship for speech between the first and second physical microphones. Section III(D), [1b]; Ex. 1005, 4:28-34, 5:19-35, FIGS. 4a-4b; Ex. 1003, ¶¶262-274. A POSITA would have recognized that, if the first and second physical microphones do not move relative to one another or to the source of sound, then the delay in time between the arrival of speech at one physical microphone and the arrival of speech at the other physical microphone would not change. Ex. 1003, ¶264. Accordingly, a POSITA would have found it obvious to implement Avendano's delay as a static filter. *Id.*

To the extent that Avendano alone does not describe a static filter, the combination of Avendano and Visser (and Bisgaard) would have rendered this feature obvious. *Id.*, ¶¶265-271. A POSITA would have found it obvious to implement Visser's cross filters w_{12} and w_{21} as static filters in at least some circumstances. *Id.*, ¶¶267-271.

Visser's "separation process could use ... an application specific adaptive filter process using some degree of a priori knowledge about the acoustic environment to accomplish substantially similar signal separation." Ex. 1006, 16:23-27. For example, Visser's filters can have "filter values ... or taps," and in cases that the device "has only a limited range of operating conditions" (*e.g.*, limited changes in "the distance from each microphone to the speaker's mouth"), "default values" can be selected for the "taps ... to account for the expected operating arrangement." *Id.*,

22:7-25, FIG. 12. As Visser explains, “the default values may adapt over time and according to environment conditions.” *Id.*

However, a POSITA would have recognized that such an adaption would not necessarily be required (*e.g.*, if the device is not expected to deviate from the “limited range of operating conditions” and/or sufficient “a priori knowledge about the acoustic environment” is known). *Id.*, 16:23-27, 22:7-25, FIG. 12; Ex. 1003, ¶¶268-271. Accordingly, a POSITA would have found it obvious to implement the delay as a static filter under these circumstances. *Id.*

To the extent that the combination of Avendano and Visser does not describe a static filter, the combination of Avendano, Visser, and Bisgaard would have rendered this feature obvious. Ex. 1003, ¶¶272-274.

Specifically, Bisgaard provides details regarding how to implement Avendano’s delay, such as configuring the delay to account for “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone.” Sections IV(B)-(C); Ex. 1010, [0042]; Ex. 1011, 5:22-24; Section III(D), [1b]; Ex. 1005, 4:28-34, 5:19-35, FIGS. 4a-4b.

A POSITA would have recognized that, if the first and second microphones do not move relative to one another or to the source of sound, then “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone” would not change. Ex. 1003, ¶274.

Accordingly, a POSITA would have found it obvious to implement the delay as a static filter. *Id.*

M. Claim 24

[24]: wherein the forming of the filter comprises: determining a first distance as distance between the first physical microphone and a mouth of the speaker; determining a second distance as distance between the second physical microphone and the mouth; and forming a ratio of the first distance to the second distance.

A POSITA would have recognized that, to account for “the amount of time used by a sound signal to propagate in the 0° azimuth direction from [a] front microphone ... to [a] rear microphone,” the sound signal would be delayed by an amount of time that is proportional to a time difference between arrival of the speech at one physical microphone and arrival of the speech at another physical microphone. Section IV(C); Ex. 1010, [0042]; Ex. 1011, 5:22-24; Ex. 1003, ¶¶275-279.

Further, a POSITA would have recognized that, given constant environment conditions, sound would propagate through an environment at a constant speed. Ex. 1003, ¶278. A POSITA would have recognized that an environment’s conditions are unlikely to change during the short time period during which sound propagates between two closely positioned microphones, and thus the speed of sound is unlikely

to change during this time period. *Id.* Thus, a POSITA would have recognized that the propagation time of sound from a source of a sound to a destination would be proportional to the distance between the source of the sound and the destination. *Id.*

Accordingly, to determine a time difference between arrivals of the speech at two physical microphones, a POSITA would have found it obvious to (i) determine a first distance between the first physical microphone and a mouth of the speaker (*i.e.*, the source), (ii) determine a second distance between the second physical microphone and the mouth, and (iii) form a ratio of the first distance to the second distance (*e.g.*, representing a proportional relationship between the first distance and the second distance). *Id.*, ¶279.

V. GROUND 3: Avendano, Visser, Bisgaard, And Hou (Claims 17- 19)

A. Hou Overview

Hou describes “[i]mproved approaches to matching sensitivities of microphones in multi-microphone directional processing systems ... so that directional noise suppression is robust.” Ex. 1008, Abstract. Specifically, Hou describes “a two-microphone directional processing system 500” for “compensat[ing] (or correct[ing]) for the relative difference in sensitivity between ... mismatched first and second microphones” and producing an “output signal ... hav[ing] robust directionality despite a mismatch between the first and second microphones.” *Id.*, 5:25-56, FIG. 5.

Hou's system receives "a first electronic sound signal" from "microphone 502," and "estimates the minimum for the first electronic sound signal" using "first minimum estimate unit 508." *Id.*, 5:27-41, FIG. 5.

Further, Hou's system receives "a second electronic sound signal" from "microphone 504," delays the "second electronic sound signal" using "delay unit 516," and "estimates the minimum for the second electronic sound signal" using "second minimum estimate unit 510." *Id.*

Further, "divide unit 512 produces a quotient by dividing the first minimum estimate by the second minimum estimate," where "[t]he quotient represents a scaling amount that is sent to a multiplication unit 514." *Id.*, 5:42-45.

Further, "[t]he second electronic sound signal is then multiplied with the scaling amount to produce a compensated sound signal." *Id.*, 5:45-47. As Hou describes, "[t]he compensated sound signal is thus compensated (or corrected) for the relative difference in sensitivity between the mismatched first and second microphones 502 and 504." *Id.*, 5:47-50.

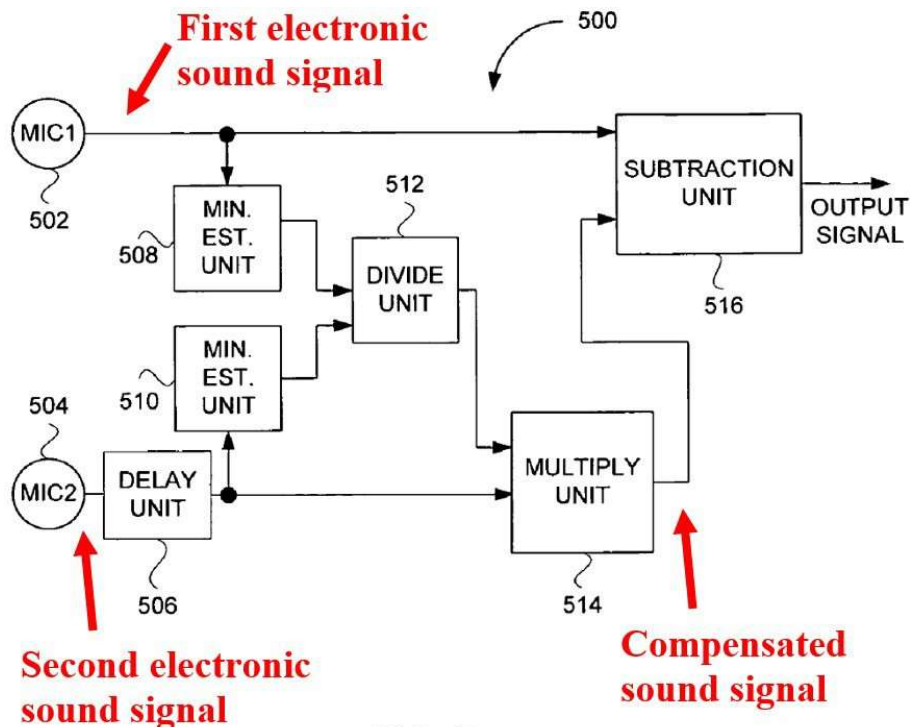


FIG. 5

Ex. 1008, FIG. 5

Further, “subtraction unit 516 then subtracts the compensated electronic sound signal from the first electronic sound signal to produce an output signal ... hav[ing] robust directionality despite a mismatch between the first and second microphones 502 and 504.” *Id.*, 5:50-56.

As Hou describes, generating a compensated sound signal ensures that “directional noise suppression is not affected by microphone mismatch, ... the drift of microphone sensitivity over time, ... [or] the non-uniform distribution of sound pressure in real-life application.” *Id.* 9:28-41; Ex. 1003, ¶¶82-91.

B. Combination of Avendano, Visser, Bisgaard, and Hou

A POSITA would have found it obvious to combine Avendano, Visser, and Bisgaard with Hou. Ex. 1003, ¶¶280-294.

First, both Avendano and Hou come from the same field of endeavor of enhancing speech and attenuating noise. Ex. 1005, Abstract, 1:24-26, 3:13-26, 3:42-60; Ex. 1008, Abstract, 2:44-52, 4:40-49; Ex. 1003, ¶282.

Second, both Avendano and Hou describe enhancing speech and attenuating noise using similar techniques, including equalizing the signal levels of two signals and generating directional signals based on the equalized signals. Ex. 1003, ¶¶283-288.

For example, Avendano's "primary microphone 106 is much closer to [an] audio source 102 than the secondary microphone 108," and thus "the intensity level is higher for the primary microphone 106 resulting in a larger energy level during a speech/voice segment." Ex. 1005, 3:27-49; Section III(A). To account for the differences in intensity levels, Avendano applies "gain factor, g" to a second signal "to equalize the signal levels" of the first and second signals.

Ex. 1005, 5:36-39, 9:54-59; FIGS. 4a-4b. Such equalization is beneficial, as "systems can suffer loss of performance when the microphone signals have different levels." *Id.*, 5:37-39. Further, Avendano uses the equalized signal to generate

directional signals, such as “cardioid primary signal (C_f)” (front) and “cardioid secondary signal (C_b)” (back). *Id.*, 41-52, 5:25-35, FIGS. 4a-4b, 6.

Hou’s “two-microphone directional processing system 500” serves a similar function as Avendano’s “gain factor,” namely “compensat[ing] (or correct[ing]) for the relative difference in sensitivity between ... mismatched first and second microphones.” Ex. 1008, 5:47-50; Ex. 1003, ¶287. Similarly, Hou explains that compensation or correction is beneficial, as “[t]he sensitivity of the microphones of the sound pick up system must be matched in order to achieve good directionality.” Ex. 1008, 1:48-2:2.

Like Avendano, Hou processes the compensated or corrected signal to generate a directional signal (*e.g.*, “an output signal ... hav[ing] robust directionality”) to aid in “directional noise suppression.” *Id.*, 5:50-56, 9:28-41.

A POSITA would have recognized that Avendano teaches the importance of equalizing signal levels of signals generated by two respective microphones, and that Hou discloses a specific “two-microphone directional processing system 500” for performing precisely this function. Ex. 1003, ¶289. Accordingly, a POSITA would have been motivated to implement Hou’s signal compensation or correction components in Avendano to further improve the equalization of the signal output by the microphones (*e.g.*, to generate a signal “hav[ing] robust directionality” and to

enhance “directional noise suppression”). Ex. 1008, 5:50-56, 9:28-41; Ex. 1003, ¶289.

As an example modification, Hou’s “first minimum estimate unit 508,” “second minimum estimate unit 510,” and divide unit 512,” would be implemented in Avendano’s “gain module 406.” Ex. 1003, ¶290. In particular, Hou’s “first minimum estimate unit 508” would receive Avendano’s signal x_2 , and Hou’s “second minimum estimate unit 510” would receive Avendano’s signal x_1 . *Id.* Further, Hou’s “divide unit 512” would calculate a gain as the quotient of the output of “first minimum estimate unit 508” and the output of “second minimum estimate unit 510.” *Id.* Further still, the calculated gain would be applied to Avendano’s signal x_1 by “multiplication unit 514.” *Id.*

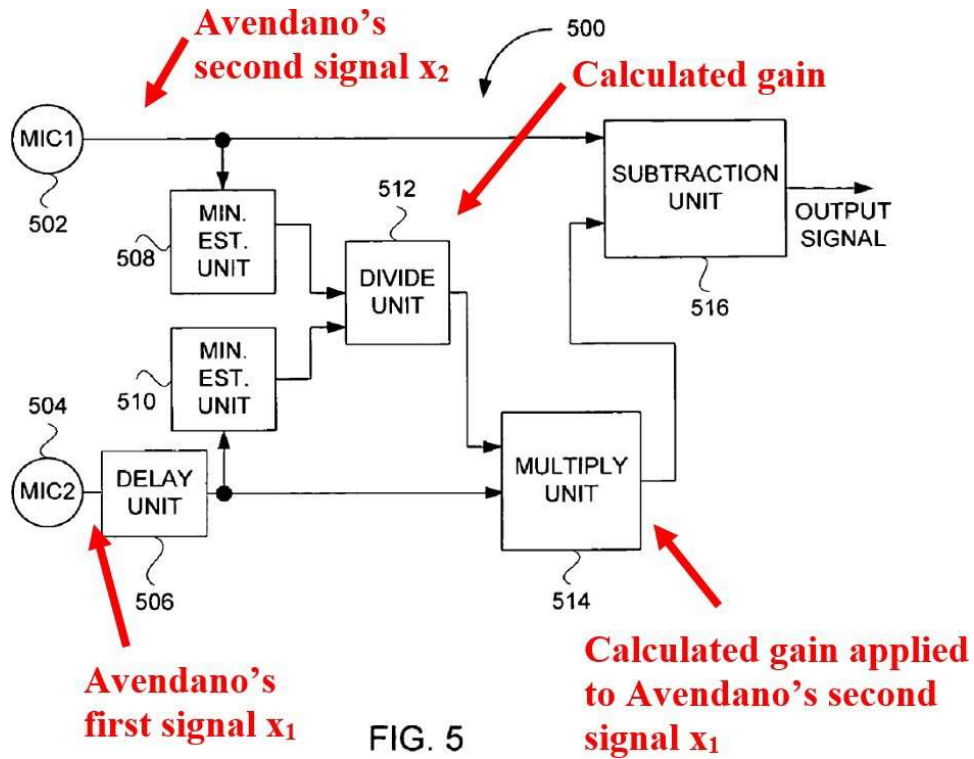


FIG. 5

Ex. 1008, FIG. 5

Although Hou multiplies one of the signals by the calculated gain, a similar equalization effect could be achieved by multiplying the other signal by the inverse of the calculated gain. Ex. 1003, ¶291.

This is analogous to Avendano's system in which a gain is calculated based on the signals x_1 and x_2 , and is applied to the signal x_2 (e.g., by a multiplexer). Ex. 1005, 4a; Ex. 1003, ¶292.

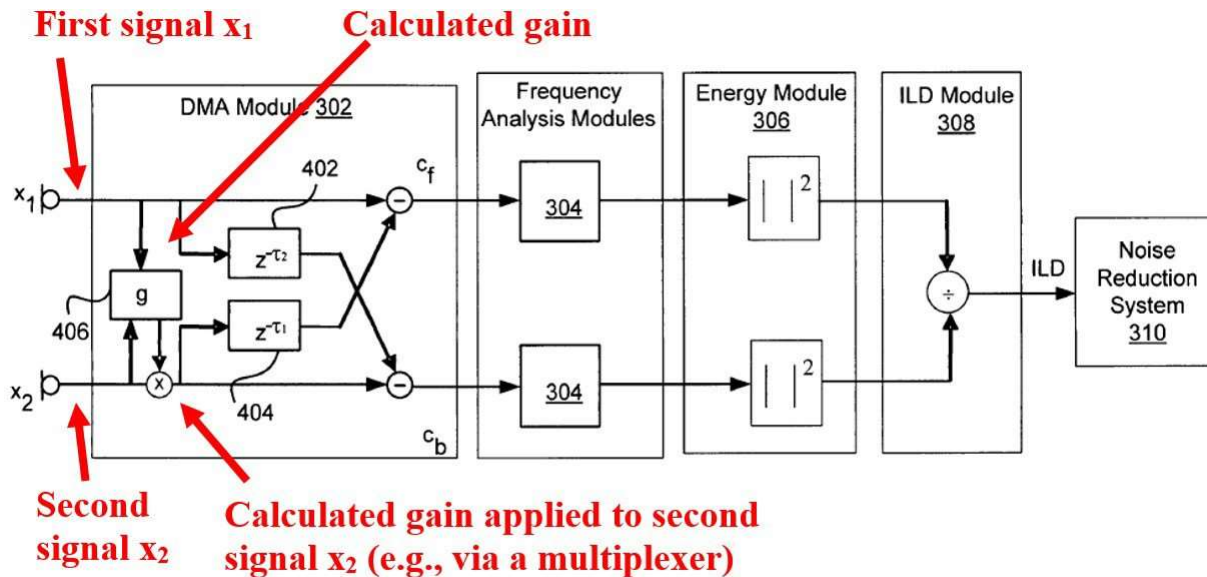


FIG. 4a

Ex. 1005, FIG. 4a

Given the similarities between Avendano’s “gain factor, g ” and Hou’s “two-microphone directional processing system 500,” a POSITA would have had a reasonable expectation of success in incorporating Hou’s “two-microphone directional processing system 500” in Avendano’s system. Ex. 1003, ¶293. For example, Avendano’s “gain factor, g ” and Hou’s “two-microphone directional processing system 500” receive similar input signals (*e.g.*, from two microphones), and generate similar outputs (*e.g.*, “equalized” or “compensated” signals). Ex. 1005, 5:36-39, 9:54-59; FIGS. 4a-4b; Ex. 1008, 5:25-50. Given these similarities, limited modifications would have been necessary to incorporate Hou’s “two-microphone directional processing system 500” into Avendano’s system and other operations in

Avendano's system would not have been impacted by the combination. Ex. 1003, ¶293.

Thus, the combination of Avendano and Hou would have been well within the grasp of a POSITA and a POSITA would have had a reasonable expectation of success in making the combination. *Id.*, ¶294.

C. Claim 17

[17]: wherein the forming of the filter comprises: generating a first quantity by applying a calibration to the second signal; generating a second quantity by applying the delay to the first signal; forming the filter as a ratio of the first quantity to the second quantity.

Hou's "first minimum estimate unit 508" would receive Avendano's second signal x_2 (e.g., from Avendano's second physical microphone 108), and apply a calibration to the second signal x_2 . (e.g., estimate a minimum of the second signal) to generate a first quantity. Section V(B); Ex. 1008, 5:38-39, FIG. 5; Ex. 1003, ¶¶295-299.

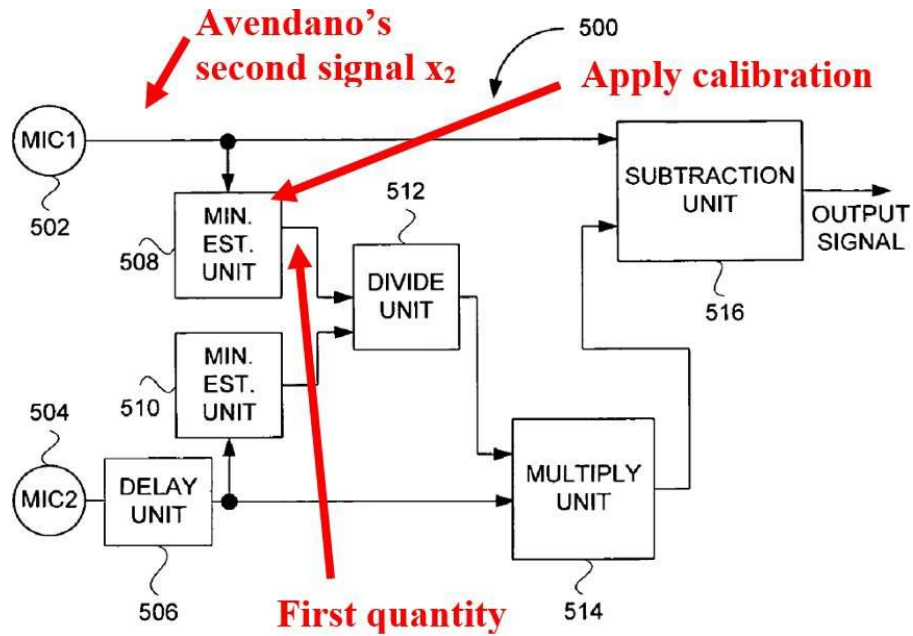
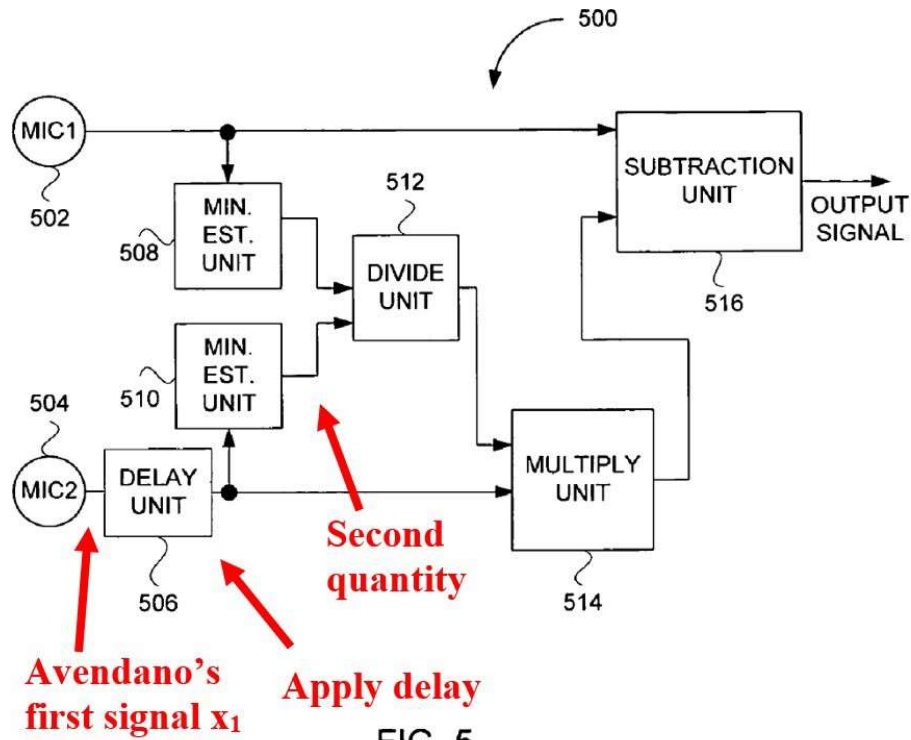


FIG. 5

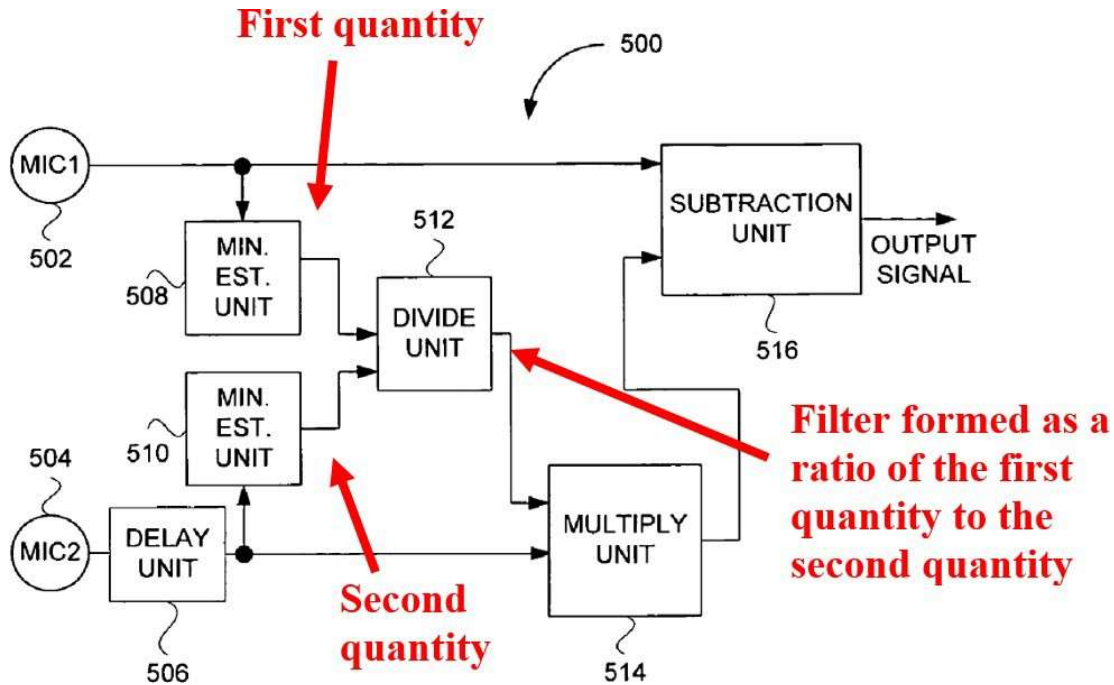
Ex. 1006, FIG. 5

Further, Hou's "delay unit 506" would receive Avendano's first signal x_1 (e.g., from Avendano's first physical microphone 106), and apply a delay to that first signal x_1 to generate a second quantity. Section V(B); Ex. 1008, 5:36-38, FIG. 5; Ex. 1003, ¶298.



Ex. 1006, FIG. 5

Further, Hou's "divide unit 512" forms a filter as a ratio of the first quantity to the second quantity. Section V(B); Ex. 1008, 5:42-45, FIG. 5. Specifically, Hou's "divide unit 512 produces a quotient by dividing the first minimum estimate by the second minimum estimate," "[t]he quotient represent[ing] a scaling amount that is sent to a multiplication unit 514." Ex. 1008, 5:42-45.



Ex. 1006, FIG. 5

D. Claims 18 and 19

[18]: wherein the generating of the energy ratio comprises generating the energy ratio for a frequency band.

[19]: wherein the generating of the energy ratio comprises generating the energy ratio for a frequency subband.

Avendano’s “frequency analysis module 304 takes the cardioid signals and mimics the frequency analysis of the cochlea (i.e., cochlear domain) simulated by a filter bank.” Ex. 1005, 4:55-5:2; Ex. 1003, ¶¶300-311; Section III(D), [1d]. Specifically, Avendano’s “frequency analysis module 304 separates the cardioid signals into frequency bands,” and performs “sub-band analysis on the acoustic

signal [to] determine[] what individual frequencies are present in the complex acoustic signal during a frame.” *Id.*

“Once the frequencies are determined,” Avendano’s “energy module 306” calculates the energy ratio (“ILD”) based on signals that are output by “frequency analysis module 304.” Ex. 1005, 5:3-10. Specifically, the energy ratio is calculated “based on bandwidth of the cochlea channel” (*e.g.*, frequency band(s) or subband(s) for which “frequency analysis module 304 ... mimics the frequency analysis of the cochlea”). *Id.*; Ex. 1003, ¶¶303, 309.

Accordingly, Avendano generates the energy ratio for a frequency band or subband. Ex. 1003, ¶¶304, 310.

This is also reflected in Avendano’s equation for the energy ratio (“ILD”), which indicates that the energy ratio is generated as a vector for a frequency band or subband (ω):

Energy ratio 

$$\boxed{ILD(t, \omega)} = \frac{\int |C_f(t', \omega)|^2 dt'}{\int_{frame} |C_b(t', \omega)|^2 dt'}$$

Ex. 1005, 6:16

**VI. GROUND 4: AVENDANO, VISSER, BISGAARD, HOU, AND
FREQUENCY ART (BYRNE, BURNETT, AND/OR BERGLUND)
(CLAIMS 20-22)**

A. Byrne Overview

Byrne describes “[t]he long-term average speech spectrum (LTASS) ... for 12 languages,” including “[LTASS] values for five samples of English.” Ex. 1009, Abstract, 2112-2120, FIG. 1. For these samples, the average of the sound pressure level (SPL)—corresponding to an intensity of sound—peaked at approximately 500 Hz, and included significant spectral components in frequency ranges that include 500 Hz (*e.g.*, in frequencies greater than 200 Hz, between 250- 1250 Hz, and between 200-3000 Hz). *Id.*; Ex. 1003, ¶¶92-95.

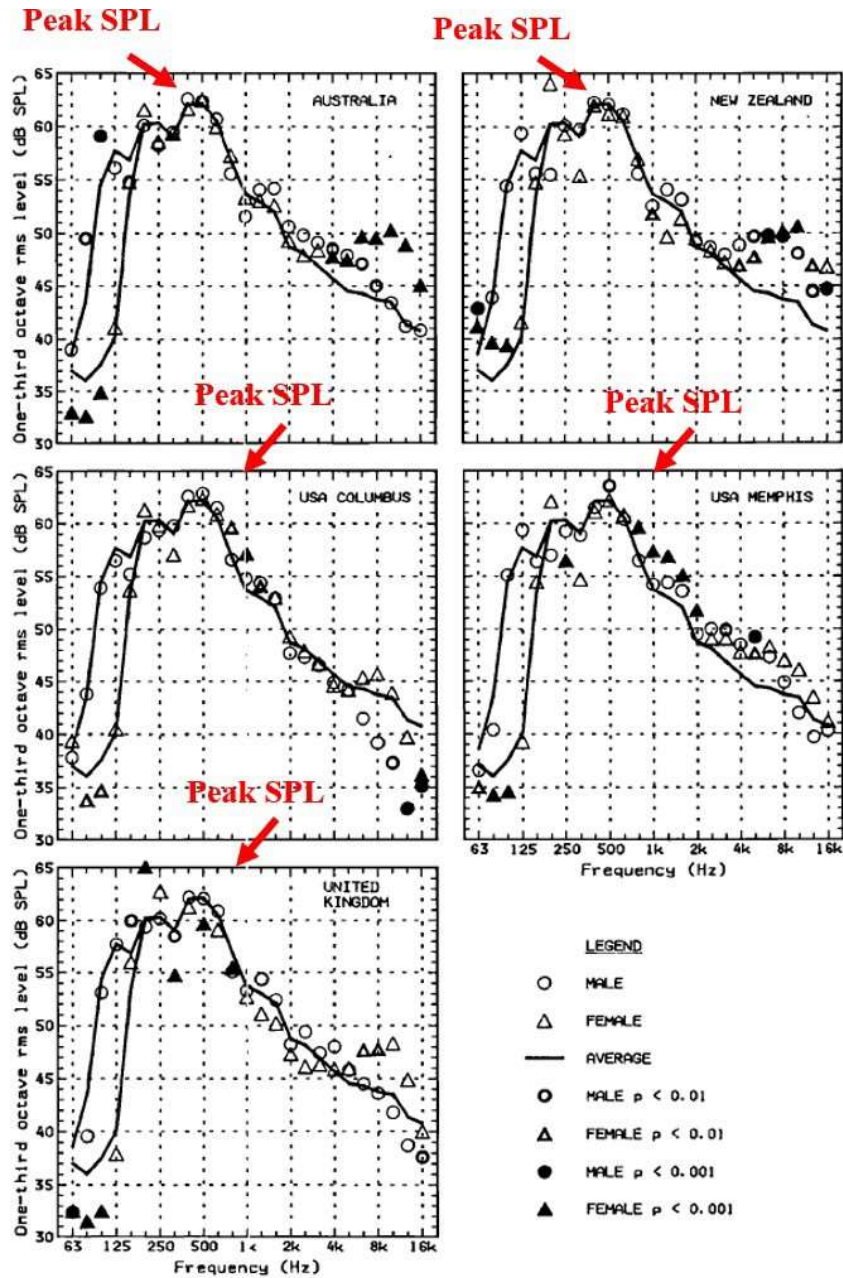


FIG. 1. Male and female long-term average speech spectrum (LTASS) values for five samples of English. Solid line shows LTASS average across 17 speech samples (all samples except Arabic), males and females separately for frequencies below 160 Hz, combined for higher frequencies.

Ex. 1009, FIG. 1

Further, Byrne's LTASS values vary for different types of speech (*e.g.*, speech uttered by different genders, according to different languages, etc.). *Id.*, FIGS. 1-9, 2112-2120.

B. Burnett Overview

Burnett describes “[s]ystems and methods ... for detecting voiced and unvoiced speech in acoustic signals having varying levels of background noise.” Ex. 1012, Abstract. Burnett's system “group[s] ... utterances by their spectral characteristics,” which enables the system to “work better in noisy environments.” *Id.*, [0063]. Specifically, Burnett's system “bandpass[es] the information from [two microphones] Mic 1 and Mic 2 so that it is possible to see which bands in the Mic 1 data are more heavily composed of noise and which are more weighted with speech.” *Id.* Example frequency bands include (i) 500-4000 Hz corresponding to “k” in “kick” (ii) 1700-4000 Hz corresponding to “sh” in “she,” (iii) 300-2500 Hz corresponding to “/i/ (‘ee’),” and (iv) 900-1200 Hz corresponding to “/a/ (‘ah’).” *Id.*, [0064]; Ex. 1003, ¶¶96-97.

C. Berglund Overview

Berglund describes “[t]he source of human exposure to low-frequency noise and its effects.” Ex. 1013, Abstract.

For example, Berglund's Figure 3 shows the spectral components of common noises experienced by passengers of “road transportation vehicles”:

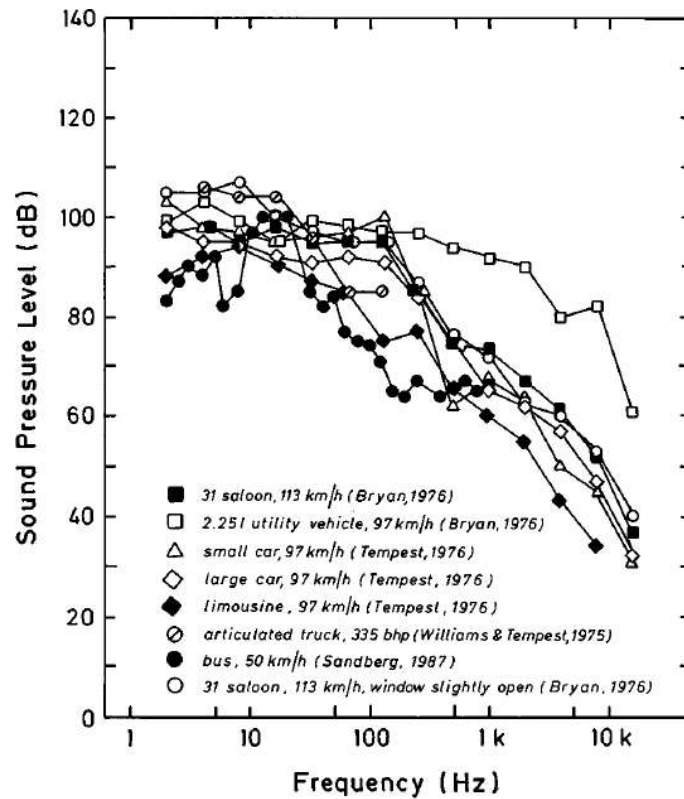


FIG. 3. Passenger noise exposure in road transportation vehicles as a function of frequency.

Ex. 1013, FIG. 3

Berglund's sound pressure level (SPL)—corresponding to an intensity of sound—for noise produced by “road transportation vehicles” is highest in the “low frequency noise” range (*e.g.*, less than approximately 250 Hz). *Id.*, 2985-2987, FIGS. 1, 3; Ex. 1003, ¶¶98-100.

D. Combination of Avendano, Visser, Bisgaard, Hou, and Frequency Art (Byrne, Burnett, and/or Berglund)

A POSITA would have found it obvious to combine Avendano, Visser, Bisgaard, and Hou with one or more of Byrne, Burnett, and Berglund. Ex. 1003, ¶¶312-323.

A POSITA would have recognized that the processes described by Avendano (and in combination with Visser, Bisgaard, and Hou) are configured to analyze audible speech in specific frequency range(s). Ex. 1003, ¶314; Section III(D); Section V(D); Ex. 1005, 4:55-5:10. Accordingly, a POSITA would have found it obvious to perform these processes in frequency range(s) that are known to have significant spectral components of speech in order to detect the presence of speech. Ex. 1003, ¶314. For example, in view of Byrne, a POSITA would have found it obvious to perform these processes in a frequency range that includes at least 500 Hz (*e.g.*, the frequency of sound expected to have the highest intensity for human speech) to more readily detect the presence of English language speech. *Id.* Example frequency ranges include (i) frequencies greater than 200 Hz, (ii) frequencies between 250-1250 Hz, and (iii) frequencies between 200-3000 Hz. *Id.*

A POSITA also would have found it obvious to filter out frequency range(s) that are known to have significant spectral components of noise to reduce the likelihood of false positives in a speech detection process. *Id.*, ¶315. For example, in

view of Berglund, a POSITA would have found it obvious to filter out frequencies corresponding to common noises experienced by passengers of “road transportation vehicles,” such as in the “low frequency noise” range (*e.g.*, less than approximately 250 Hz). Ex. 1013, FIG. 3; Ex. 1003, ¶315.

A POSITA also would have recognized that the frequency range may vary, depending on the type of speech that is being detected (*e.g.*, the gender of the speaker, the spoken language, etc.), the type of noise that is being filtered out, and the expected spectral distributions thereof. Ex. 1003, ¶316.

Further, a POSITA would have found it obvious to tune the frequency range to accommodate the specific operational requirements of the system. *Id.*, ¶¶317- 323. For example, a POSITA would have recognized that the time and/or computational resources required to perform the processes described by Avendano, Visser, Bisgaard, and Hou would depend on the range of frequencies that are analyzed. *Id.*, ¶318. Accordingly, a POSITA would have found it obvious to balance (i) analyzing a larger range of frequencies to conduct a more comprehensive analysis, and (ii) analyzing a smaller range of frequencies to improve speed, accuracy, and/or efficiency. *Id.*

As another example, a POSITA also would have recognized that, as taught by Burnett, analyzing certain frequency ranges may be beneficial in improving the separation of “voiced and unvoiced speech from background acoustic noise.” Ex.

1012, [0063]. As Burnett describes, a system can “group [a speaker’s] utterances by their spectral characteristics” by “bandpass[ing] the information from [two microphones] so that it is possible to see which bands in [one of the microphones] are more heavily composed of noise and which are more weighted with speech.” *Id.* In particular, a system can use bandpass filters to capture signals in certain frequency ranges to detect unvoiced speech (*e.g.*, 500-4000 Hz to detect “k” in “kick,” and 1700-4000 Hz to detect “sh” in “she”), and in other frequency ranges to detect voiced speech (*e.g.*, 300-2500 Hz to detect “ee,” and 900-1200 Hz to detect “ah”). *Id.* In view of Burnett, a POSITA would have found it obvious to tune the frequency range, depending on the specific unvoiced speech and/or voiced speech that is to be detected. Ex. 1003, ¶319-322.

E. Claim 20

[20]: wherein the frequency subband includes frequencies higher than approximately 200 Hertz (Hz).

A POSITA would have found it obvious to perform the processes described by Avendano, Visser, Bisgaard, and Hou in frequency range(s) known to have significant spectral components of speech, while filtering out frequency range(s) that are known to have significant spectral components of noise. Section VI(D); Ex. 1003, ¶¶324-329. This would have resulted in performing the processes in a frequency subband that includes frequencies higher than approximately 200 Hz. *Id.*

For example, as taught by Byrne, the spectral components of English language speech are predominantly in a frequency range greater than 200 Hz. Ex. 1009, FIG. 1; Ex. 1003, ¶326.

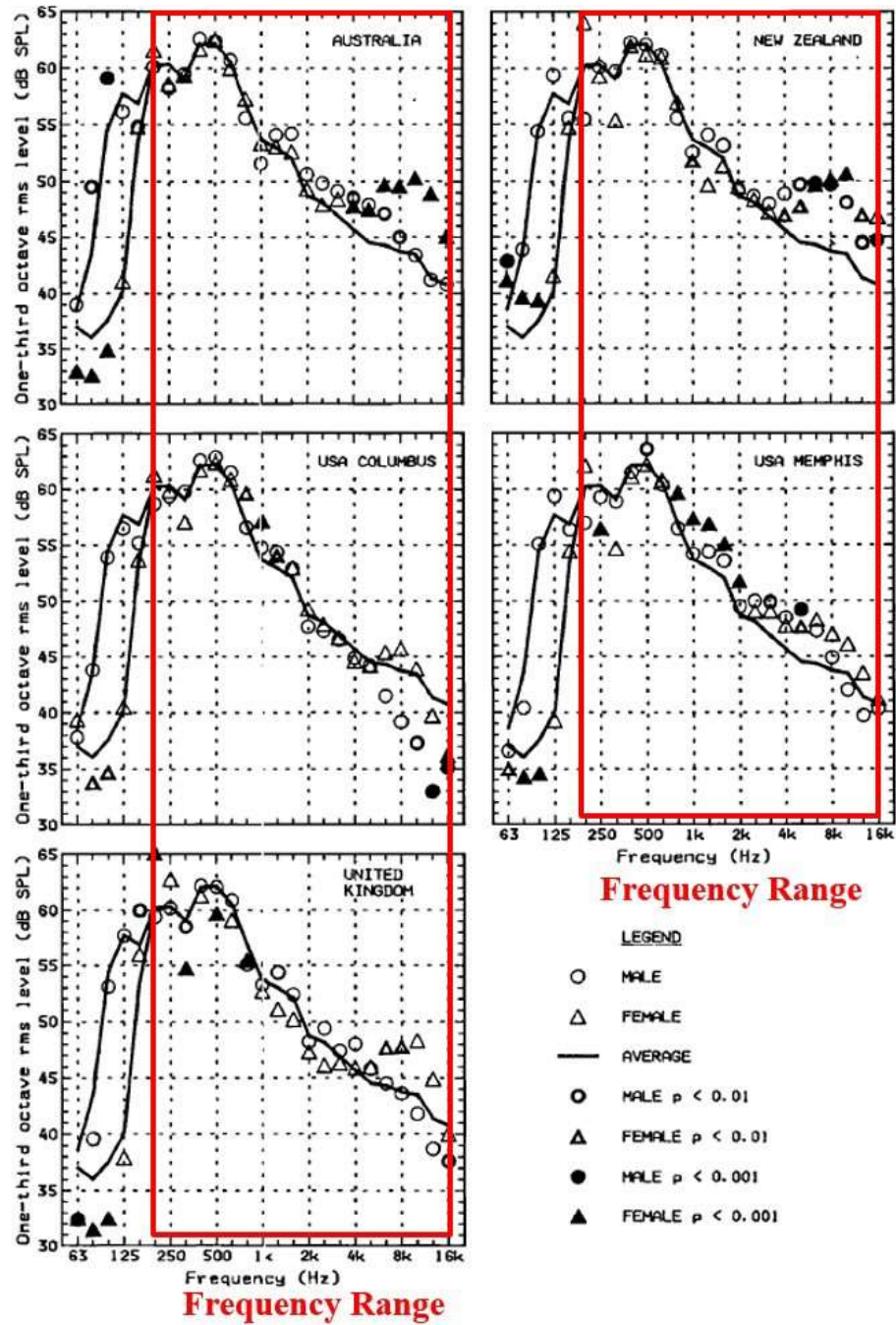


FIG. 1. Male and female long-term average speech spectrum (LTASS) values for five samples of English. Solid line shows LTASS average across 17 speech samples (all samples except Arabic), males and females separately for frequencies below 160 Hz, combined for higher frequencies.

Ex. 1009, FIG. 1

Further, as taught by Berglund, the spectral components of common noises experienced by passengers of “road transportation vehicles” is highest in the “low frequency noise” range (*e.g.*, less than approximately 250 Hz). Ex. 1009, 2985-2987, FIGS. 1, 3; Ex. 1003, ¶327.

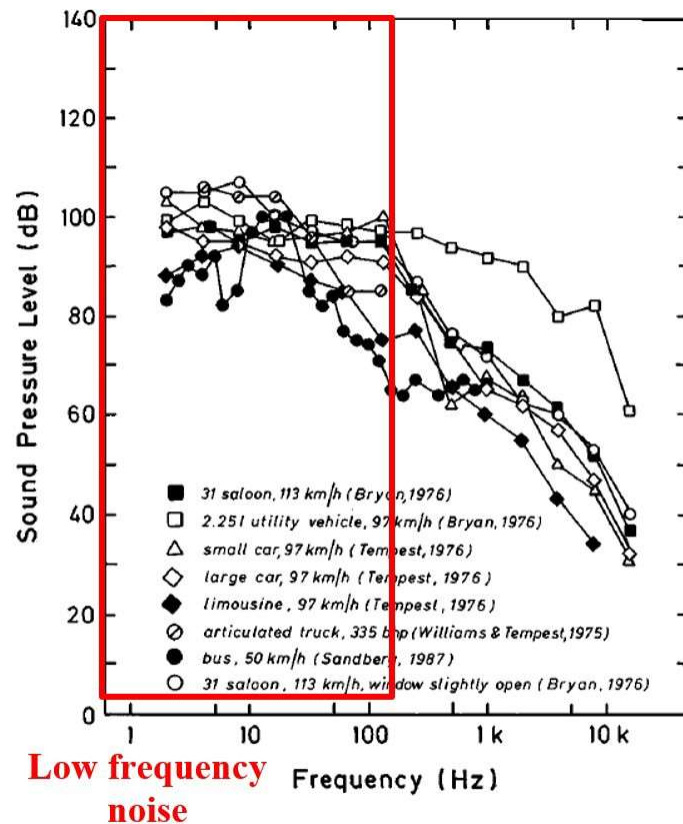


FIG. 3. Passenger noise exposure in road transportation vehicles as a function of frequency.

Ex. 1013, FIG. 3

Further, in view of Burnett, a POSITA would have found it obvious to perform the processes in frequency range(s) known to have “voiced and unvoiced speech” to improve the separation of “voiced and unvoiced speech from background acoustic noise.” Ex. 1012, [0063]; Ex. 1003, ¶328-329. As described by Burnett, these

frequency ranges would include frequencies higher than approximately 200 Hz (*e.g.*, 500-4000 Hz to detect “k” in “kick,” 1700-4000 Hz to detect “sh” in “she,” 300-2500 Hz to detect “ee,” and 900-1200 Hz to detect “ah”). Ex. 1012, [0064].

F. Claims 21 and 22

[21]: wherein the frequency subband includes frequencies in a range from approximately 250 Hz to 1250 Hz.

[22]: wherein the frequency subband includes frequencies in a range from approximately 200 Hz to 3000 Hz.

When evaluating whether a prior range anticipates a species, courts evaluate whether the claimed species is “critical.” *ClearValue, Inc. v. Pearl River Polymers, Inc.*, 668 F.3d 1340, 1344-45 (Fed. Cir. 2012). “[W]here there is a range disclosed in the prior art, and the claimed invention falls within that range, the burden of production falls upon the patentee to come forward with evidence of teaching away, unexpected results, or other pertinent evidence of nonobviousness.” *E.I. DuPont de Nemours & Co. v. Synvina C.V.*, 904 F.3d 996, 1006 (Fed. Cir. 2018) (internal citations omitted).

There is no evidence in the ’611 patent that either of (i) “a range from approximately 250 Hz to 1250 Hz” or (ii) “a range from approximately 200 Hz to 3000 Hz” is critical. For example, in the specification, the **only** explanation regarding these frequency ranges reads as follows:

The ratio R can be calculated for the entire frequency band of interest, or can be calculated in frequency subbands. One effective subband discovered was 250 Hz to 1250 Hz, another was 200 Hz to 3000 Hz, but many others are possible and useful.

Ex. 1001, 6:42-46.

As the '611 patent concedes, the frequency ranges 250-1250 Hz and 200-3000 Hz are but two of “many” frequency ranges that are “possible” and “useful.” *Id.* The '611 patent does not explain how using these frequency ranges would produce any new and unexpected results relative to the prior art. Ex. 1003, ¶¶332-334, 347-349. In fact, the '611 patent does not explain how using these frequency ranges would be more advantageous than using any other specific frequency ranges or even “the entire frequency band of interest.” *Id.* Further, the '611 patent does not provide any guidance regarding the selection of “possible” and “useful” frequency ranges. *Id.* Also, nothing in the prosecution history of the '611 patent asserts any new and unexpected results relative to the prior art that are associated with these frequency ranges. *Id.*

Further still, there is no evidence in the disclosure or the prosecution history of the '611 patent that the prior art taught away from using either of these frequency ranges, or that any other pertinent secondary considerations are associated with either of these frequency ranges. *Id.*

To the extent that the Patent Owner contends that the '611 patent's description of frequency ranges of certain "voiced and unvoiced speech" (*i.e.*, Ex. 1001, 17:18-36) demonstrates the criticality of either of the claimed ranges, Petitioner respectfully disagrees, as the description does not correspond to either of the claimed ranges.

For example, the '611 patent uses bandpass filters to capture signals in certain frequency ranges to detect unvoiced speech (*e.g.*, 500-4000 Hz to detect "k" in "kick," and 1700-4000 Hz to detect "sh" in "she"), and in other frequency ranges to detect voiced speech (*e.g.*, 300-2500 Hz to detect "ee," and 900-1200 Hz to detect "ah"). *Id.*, 17:29-36.

However, none of these frequency ranges corresponds to the claimed frequency range of 250-1250 Hz. Instead, the claimed frequency range would arbitrarily include the entirety of the frequency range for some types of speech (*e.g.*, 900-1200 Hz ("ah")), exclude a portion of the frequency range for other types of speech (*e.g.*, 300-2500 Hz ("ee") and 500-4000 Hz ("k" in "kick")), and exclude the entirety of the frequency range for yet other types of speech (*e.g.*, 1700-4000 Hz ("sh" in "she")).

Likewise, none of these frequency ranges corresponds to the claimed frequency range of 200-3000 Hz. Instead, the claimed frequency range would arbitrarily include the entirety of the frequency range for some types of speech (*e.g.*,

900-1200 Hz (“ah”) and 300-2500 Hz (“ee”)), and exclude a portion of the frequency range for other types of speech (*e.g.*, 500-4000 Hz (“k” in “kick”) and 1700-4000 Hz (“sh” in “she”)).

The ’611 patent does not describe any criticality associated with these arbitrarily selected frequencies. Nor does any criticality exist for any of these arbitrarily selected frequencies. Ex. 1003, ¶¶330-359.

However, even if this description in the ’611 patent were to demonstrate the criticality of either of the claimed ranges, Burnett nevertheless includes this same description verbatim. Ex. 1012, [0063]-[0064]. As such, the claimed ranges would still be disclosed by Burnett.

Further, Avendano in combination with Byrne, Burnett, and/or Berglund (and in the further combination with Visser, Bisgaard, and Hou) also discloses the claimed ranges. Ex. 1003, ¶¶335-344, 350-359.

By virtue of the inclusive term “including,” claims 21 and 22 do not preclude a frequency subband from including frequencies in addition to those between approximately 250-1250 Hz and 200-3000 Hz, respectively. A POSITA would have found it obvious to perform the processes described by Avendano, Visser, Bisgaard, and Hou in a frequency subband that includes at least frequencies higher than approximately 200 Hz to more readily detect the presence of English language speech. Sections VI(D)-(E). Such a frequency subband would include frequencies in

a range from approximately 250-1250 Hz and 200-3000 Hz, among others. Ex. 1003, ¶¶339, 354.

Nevertheless, to the extent that claims 21 and 22 preclude a frequency subband from including frequencies other than those between approximately 250-1250 Hz and 200-3000 Hz, respectively, as described above, a POSITA also would have found it obvious to perform the processes described by Avendano, Visser, and Bisgaard in these frequency subbands. Ex. 1003, ¶¶340-344, 355-359.

Frequencies in a Range from Approximately 250 Hz to 1250 Hz:

A POSITA would have found it obvious to perform the processes in frequency range(s) known to have significant spectral components of speech (*e.g.*, to more readily detect the presence of human speech). Section VI(D); Ex. 1003, ¶336.

As taught by Byrne, English language speech is expected to have significant spectral components in a frequency range between 250 Hz and 1250 Hz, as this frequency range would include the frequencies of sound expected to have the highest intensities for human speech (*e.g.*, frequencies of and around 500 Hertz). Ex. 1009, FIG. 1; Ex. 1003, ¶337-338.

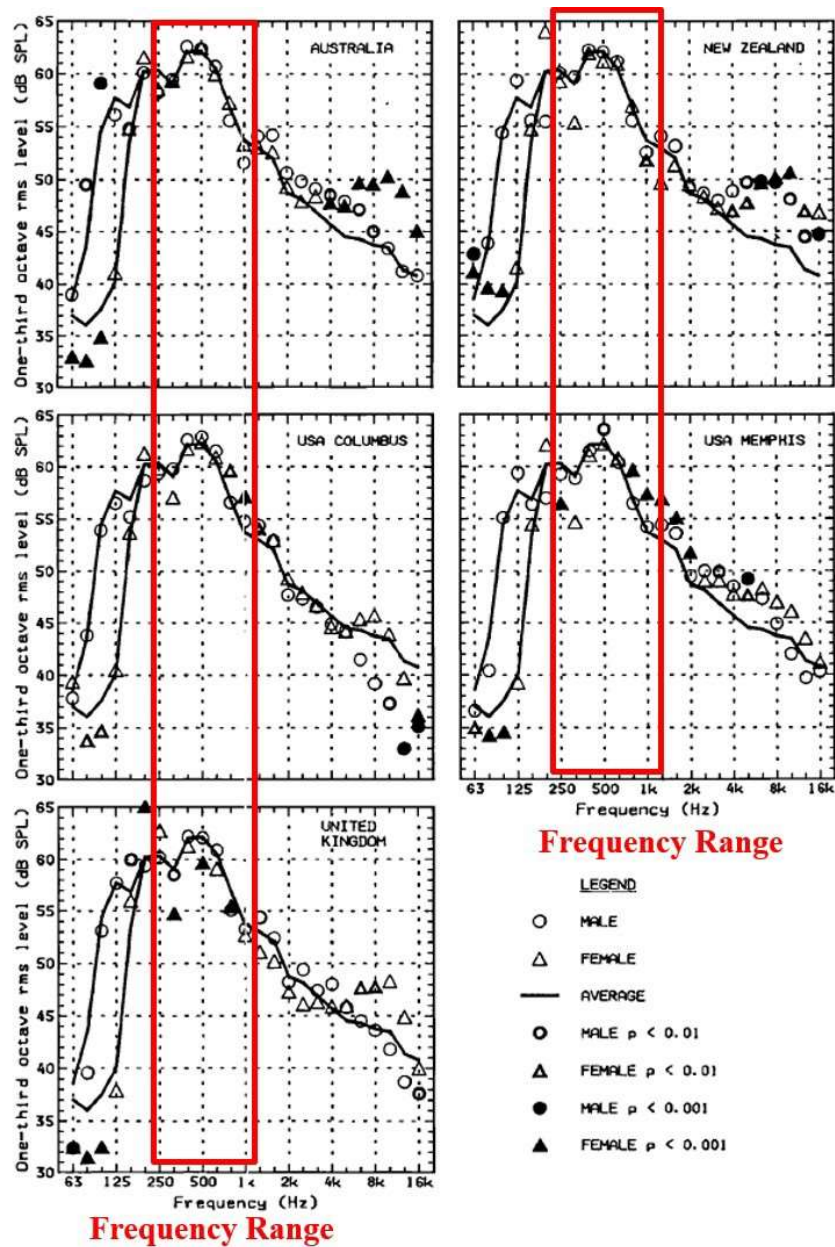


FIG. 1. Male and female long-term average speech spectrum (LTASS) values for five samples of English. Solid line shows LTASS average across 17 speech samples (all samples except Arabic), males and females separately for frequencies below 160 Hz, combined for higher frequencies.

Ex. 1009, FIG. 1

As taught by Burnett, certain types of speech can be detected in a frequency range of 250-1250 Hz (e.g., “ah” in frequencies 900-1200 Hz. Ex. 1012, [0064].

Further, a POSITA would have found it obvious to filter out frequency range(s) that are known to have significant spectral components of noise, such as in a “low frequency noise” range less than approximately 250 Hz. Section VI(D); Ex. 1013, 2985-2987, FIGS. 1, 3; Ex. 1003, ¶341.

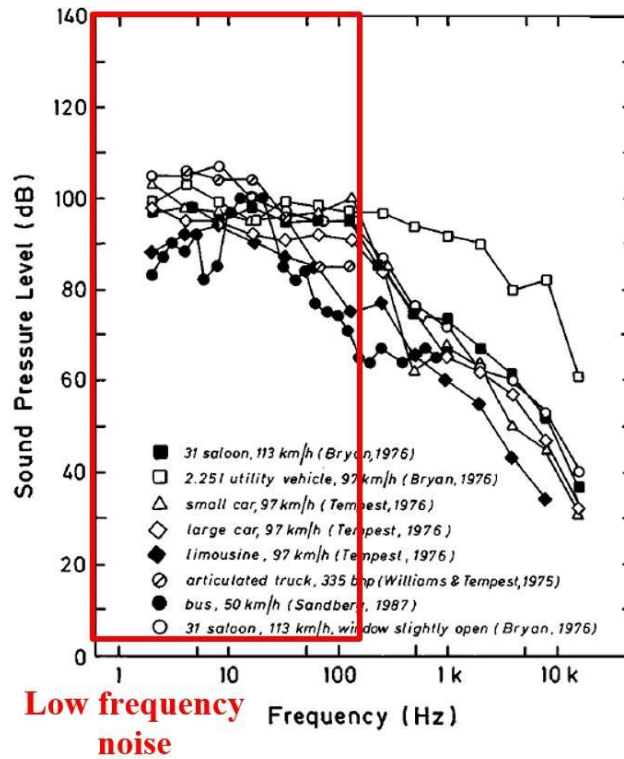


FIG. 3. Passenger noise exposure in road transportation vehicles as a function of frequency.

Ex. 1013, FIG. 3

Further, a POSITA would have found it obvious to tune the frequency range, depending on the specific use case at hand. Section VI(D). For example, a POSITA would have found it obvious to prioritize analysis of frequency range(s) known to have significant spectral components of speech, such as a frequency range between

approximately 250 Hz and 1250 Hz (*e.g.*, to improve the speed, accuracy, and/or efficiency in detecting voice activity). Ex. 1003, ¶342.

Additionally, a POSITA would have found it obvious to use other frequency ranges, depending on the type of speech that is being detected, the type of noise that is being filtered out, and the expected spectral distributions thereof. *Id.*

Frequencies in a Range from Approximately 200 Hz to 3000 Hz:

A POSITA would have found it obvious to perform the processes in frequency range(s) known to have significant spectral components of speech (*e.g.*, to more readily detect the presence of human speech). Section VI(D); Ex. 1003, ¶351.

As taught by Byrne, human speech is expected to have significant spectral components in a frequency range between 200 Hz and 3000 Hz, as such a frequency range would include the frequencies of sound expected to have the highest intensities for human speech (*e.g.*, frequencies of and around 500 Hertz). Ex. 1009, FIG. 1; Ex. 1003, ¶352.

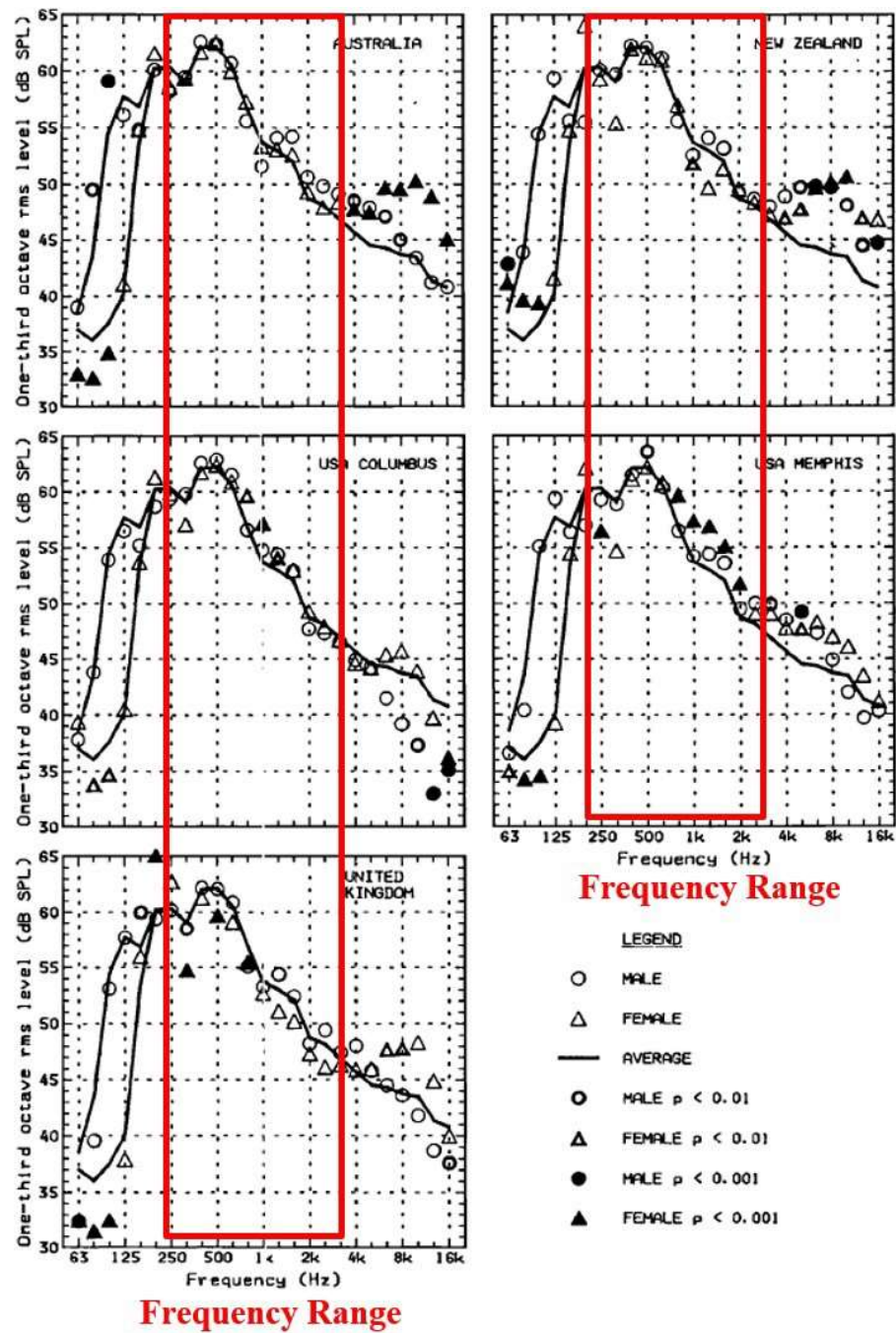


FIG. 1. Male and female long-term average speech spectrum (LTASS) values for five samples of English. Solid line shows LTASS average across 17 speech samples (all samples except Arabic), males and females separately for frequencies below 160 Hz, combined for higher frequencies.

Ex. 1009, FIG. 1

As taught by Burnett, certain types of speech can be detected in a frequency range of 200-3000 Hz (*e.g.*, “ee” in frequencies 300-2500 Hz). Ex. 1012, [0064].

Further, a POSITA would have found it obvious to filter out frequency range(s) that are known to have significant spectral components of noise, such as in a “low frequency noise” range less than approximately 250 Hz. Section VI(D); Ex. 1013, 2985-2987, FIGS. 1, 3; Ex. 1003, ¶356.

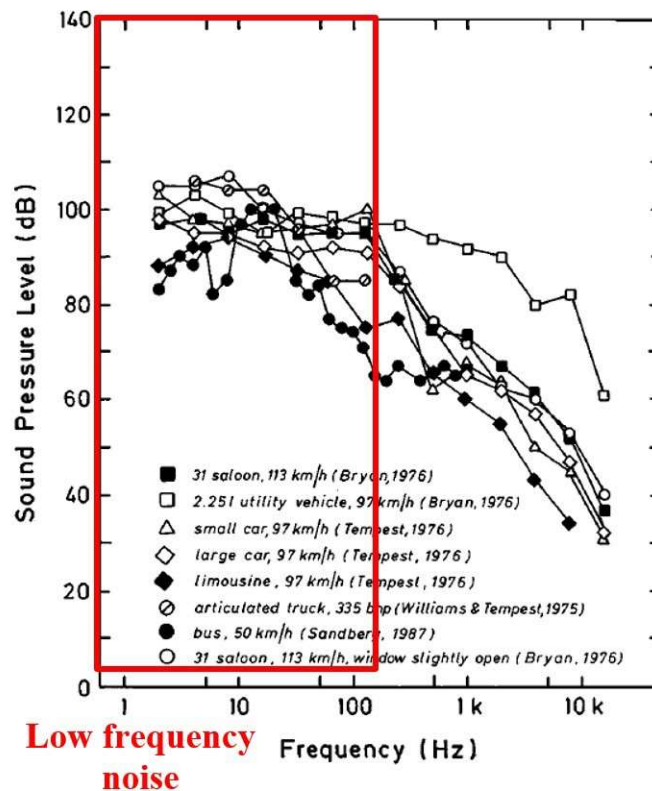


FIG. 3. Passenger noise exposure in road transportation vehicles as a function of frequency.

Ex. 1013, FIG. 3

Further, a POSITA would have found it obvious to tune the frequency range, depending on the specific use case at hand. Section VI(D). For example, a POSITA would have found it obvious to prioritize analysis of frequency range(s) known to have significant spectral components of speech, such as a frequency range between approximately 200 Hz and 3000 Hz (*e.g.*, to improve the speed, accuracy, and/or efficiency in detecting voice activity). Ex. 1003, ¶357. Additionally, a POSITA also would have found it obvious to use other frequency ranges, depending on the type of speech that is being detected, the type of noise that is being filtered out, and the expected spectral distributions thereof. *Id.*

VII. INSTITUTION IS APPROPRIATE HERE

A. The *Fintiv* Factors Support Institution

1. Factor 1: Potential Stay

On September 23, 2021, Patent Owner sued Petitioner for infringement of the '611 patent in *Jawbone Innovations, LLC v. Google LLC*, No. 6:21-cv-00985 (W.D. Tex.). On February 1, 2023, the case was transferred from the Western District of Texas to the Northern District of California ("NDCA"), *Jawbone Innovations, LLC v. Google LLC*, 3:23-cv-00466 (N.D. Cal.) ("the Litigation"). Ex. 1020. On March 13, 2023, Google filed a motion to stay the Litigation pending *inter partes* reviews. On April 27, 2023, the Court granted Google's motion to stay, vacating all pending

dates and deadlines pending final resolution of Google’s IPR proceedings. Ex. 1021. Thus, this factor weighs against discretionary denial.

2. Factor 2: Proximity of Trial to FWD

The NDCA has not set a trial date. In fact, the Court vacated all dates in the stay order. Thus, this factor weighs against discretionary denial. *Global Tel*Link Corp. v. HLFIP Holding, Inc.*, IPR2021-00444, Paper 14, at 16-19 (P.T.A.B. Jul. 22, 2021).

3. Factor 3: Investment in Parallel Proceeding

Patent Owner filed its complaint on September 23, 2021 and served its infringement contentions on January 13, 2022. Aside from those contentions, the parties have invested little in the Litigation. Fact discovery opened but is now stayed pending IPR. No expert reports have been served.

The NDCA has not established a case schedule. An initial case management conference has not yet taken place in the NDCA due to the Court’s stay order vacating all dates and deadlines, including the date for the initial case management conference. Minimal investment has been made in the case because fact discovery remains in its infancy. Also, the “remaining investment” significantly outweighs any investment made thus far, which weighs against discretionary denial. *Samsung Elecs. Am. Inc. v. Snik LLC*, IPR2020-01428, Paper 10, at 11 (P.T.A.B. Mar. 9, 2021).

4. Factor 4: Overlapping Issues

If this IPR is instituted, Petitioner cannot pursue in the Litigation any invalidity grounds raised or that could have been reasonably raised in this IPR. Thus, duplicative efforts or conflicting decisions are unlikely.

Petitioner hereby stipulates that, if this Petition is instituted, Petitioner will not pursue the grounds (i.e., Grounds 1-4 in the Amazon IPR based on Avendano, Visser, Bisgaard, Hou, and Frequency Art) identified in this petition in district court. Thus, duplicative efforts or conflicting decisions are unlikely. This factor weighs against discretionary denial. *Sand Revolution II, LLC v. Cont'l Intermodal Grp.-Trucking LLC*, IPR2019-01393, Paper 24, at 11-12 (P.T.A.B. Jun. 16, 2020) (informative).

5. Factor 5: Parties in Parallel Proceedings

The parties are the same, but trial in the Litigation will not start before this IPR reaches a final written decision based on the NDCA's stay order pending IPR. Thus, this factor is neutral. *Google LLC v. Jawbone Innovations, LLC*, IPR2022-00630, Paper 10, at 17 (P.T.A.B. Sept. 13, 2022).

6. Factor 6: Other Circumstances

To Petitioner's knowledge, other than Amazon's co-pending IPR petition on the same grounds that Petitioner is seeking to join in its accompanying motion for joinder, the '611 patent has never been compared to this combination of references

by the PTO, any court, or any jury. Despite this, Patent Owner has asserted the '611 patent against Petitioner and nine other defendants. Given Amazon's co-pending IPR petition has been instituted and placed the patentability of claims 1-28 of the '611 patent at issue, the Board should institute review to resolve the cloud over these claims.

The petition's merits are particularly strong, as the Board recognized by instituting Amazon's co-pending IPR petition (which this petition copies). This strongly favors institution. *Apple Inc. v. Fintiv, Inc.*, IPR2020-00019, Paper 11, at 18 (P.T.A.B. Mar. 18, 2020).

Accordingly, the *Fintiv* analysis favors institution.

B. Google's Previously Filed IPR Petition Does Not Warrant Discretionary Denial (*General Plastic*)

The *General Plastic* factors on balance weigh in favor of institution and joinder with the Amazon IPR. While Google previously challenged the '611 patent, this case is distinguishable from *Apple Inc. v. Uniloc 2017 LLC*, IPR2020-00854, Paper 9 (P.T.A.B. Oct. 28, 2020) (precedential) for the reasons explained below.

1. General Plastic Factors 1-3

Factors 1-3 seek to preclude a joinder petitioner from using a prior Board decision or preliminary response as a "roadmap" to cure deficiencies in a second filed petition. *Code200, UAB v. Bright Data Ltd.*, IPR2022-00861, Paper 18, at 5

(P.T.A.B. Aug. 23, 2022) (“*Code200*”) (precedential). When “the later petition is not refined based on lessons learned from later developments,” road-mapping concerns are “minimized.” *Id.*; *Cisco Sys., Inc. v. Centripetal Networks, Inc.*, IPR2022-01151, Paper 12 at 48 (P.T.A.B. Jan. 4, 2023) (“*Cisco*”).

Google previously filed a petition challenging claims 1-28 of the ’611 patent that was denied based on claim construction. *Google LLC v. Jawbone Innovations, LLC*, IPR2022-00604, Paper 12 (P.T.A.B. Oct. 6, 2022). But Google’s joinder petition does not map any prior Board decision or Patent Owner paper. Google instead seeks to join an already-instituted case without altering the instituted grounds or evidence. *Code200* at 5.

Google could not have road-mapped any prior paper in its joinder petition. Google’s petition is a copy of an Amazon petition that, in turn, is a copy of an Apple petition filed on June 3, 2022. *Apple Inc. v. Jawbone Innovations, LLC*, IPR2022-01085, Paper 2 (P.T.A.B. Jun. 3, 2022). The Apple petition was filed before Patent Owner’s preliminary response to Google’s first petition (July 11, 2022) and before the Board’s subsequent denial of institution (October 6, 2022). *See Google*, IPR2022-00604, Paper 8 (P.T.A.B. Jul. 11, 2022); *Google*, IPR2022-00604, Paper 12 (P.T.A.B. Oct. 6, 2022). With no papers available when Apple filed its petition, there was nothing to map. Google was also not aware of Byrne, Bisgaard, or Berglund asserted in the Amazon IPR until Apple filed its petition. The Board in

Google LLC v. Express Mobile, Inc. granted joinder under similar circumstances for a petitioner whose first petition was denied on the merits, finding that “there is no reason to conclude that Petitioner used the filings or decision . . . to obtain an unfair advantage in this [joinder] [p]etition, which was essentially prepared by someone else.” IPR2022-00790, Paper 15 at 8 (P.T.A.B. Sept. 27, 2022) (“*Express Mobile*”).

Google similarly has not “strategically stage[d] [its] prior art and arguments in multiple petitions” to gain an unfair advantage over Patent Owner. *General Plastic* at 17. Google is not a real party in interest in either the Apple or Amazon petition. While Google was aware of Avendano, Visser, Hou, and Burnett asserted in Apple’s petition, none of that art overlaps with Google’s first petition and Google simply seeks to join an existing proceeding.

It was reasonable for Google to wait until Apple’s petition was instituted before seeking joinder in that case. Apple settled before institution, *Apple Inc. v. Jawbone Innovations, LLC*, No. IPR2022-01085, Paper 14 (P.T.A.B. Jan. 9, 2023), so Google had no opportunity to seek joinder in that proceeding during the appropriate joinder window, *see* 35 U.S.C. § 315(c); 37 C.F.R. § 42.122(b). Google was instead obligated to wait until Amazon’s copycat petition was instituted. *Express Mobile* at 10. Like the situation in *Express Mobile*, Google was time-barred from filing a petition after the institution decision in its first proceeding, so Google’s delay was not gamesmanship or strategic staging of prior art but the reality that

Google could not be granted joinder absent institution in the Apple (and then Amazon) proceeding. *Id.* at 10-11.

Because Google has not road-mapped any prior Board decision or Patent Owner paper, Google was not aware of all art asserted in the Amazon IPR until Apple filed its petition, and because Patent Owner would not be prejudiced by Google seeking solely to maintain the *inter partes* proceeding instituted against claims 1-28 of the '611 patent, the Board should find factors 1-3, taken together, do not weigh in favor of exercising discretion to deny institution. *Express Mobile* at 7-9.

This case also stands apart because Google is not filing multiple, subsequent petitions from lessons learned, or piling on to multiple IPR challenges to the same patent. Google instead is seeking to join the Amazon IPR to ensure the existing proceeding reaches a final written decision. This adds no additional burden to Patent Owner and requires no additional resources from the Board. *Express Mobile* at 9.

2. General Plastic Factors 4 and 5

The Board considers *General Plastic* factors 4 (length of time from knowledge of art) and factor 5 (explanation for time between petitions) “to assess and weigh whether a petitioner should have or could have raised the new challenges earlier.” *Intel* at 11 (quoting *General Plastic* at 18). As explained above, and like the case in *Express Mobile*, it was reasonable for Google to wait and seek joinder until

after the Amazon IPR was instituted because Google could not join the proceeding before institution.

It was also reasonable for Google to file its joinder petition and seek to maintain the Amazon IPR based on the institution decision in Google's first petition. In Google's first petition, the Board denied institution on a claim construction issue that arose after the petition was filed. *Google*, IPR2022-00604, Paper 12 at 8-13. This claim interpretation issue arose at institution (October 6, 2022), after Apple's petition was filed (June 3, 2022) and after Google was time-barred. *See General Plastic* at 10-11. Once institution was denied, Google reasonably waited until it could join Apple's petition (e.g., when the Amazon IPR was instituted). *Express Mobile* at 10-11.

3. *General Plastic* Factors 6 and 7

Factor 6 (finite resources) and factor 7 (one-year deadline) weigh in favor of institution because Google's petition introduces no new issues that are not already in the Amazon IPR and no changes to the existing schedule. As the Board found in *Intel*, "instituting this Petition will [not] significantly affect the resources of the Board or our ability to issue a final determination within the one-year statutory timeline." *Intel* at 14. The Board should permit joinder because it already "found the challenges reasonably likely to be successful" and it will "continue expending resources to decide the merits of the [petition] regardless of joinder." *Id.*

The Director in *Code200* moreover found that “the one-year statutory time period may be adjusted for a joined case under 35 U.S.C. § 316(a)(11),” if necessary, and the “the Board’s mission ‘to improve patent quality and restore confidence in the presumption of validity that comes with issued patents’” favors resolving a pending IPR petition. *Code200* at 6. Google seeks to join the Amazon IPR as an understudy to ensure the case reaches a final written decision. This will not unduly burden the Board or Patent Owner. *Ericsson* at 13.

C. Discretionary Denial Under § 325(d) is Also Not Appropriate

The Office has not previously considered “the same or substantially the same prior art or arguments.” 35 U.S.C. §325(d). Here, all grounds rely on Avendano, which the PTO never considered during original examination of the ’611 patent.

Further, though this Petition presents the same grounds as in Amazon IPR (IPR2023-00286), Petitioner is filing this petition to preserve its ability to maintain an IPR on the merits in the event that Amazon terminates its involvement. If Amazon does not terminate and Petitioner is joined to the Amazon IPR, Petitioner will take an understudy role in the Amazon IPR. Thus, the proposed grounds will not be cumulative of references previously considered by the Board.

VIII. MANDATORY NOTICES UNDER 37 C.F.R. § 42.8(a)(1)

A. Real Party-In-Interest Under 37 C.F.R. § 42.8(b)(1)

Google LLC is the real party-in-interest for this petition.⁶

B. Related Matters Under 37 C.F.R. § 42.8(b)(2)

To the best knowledge of Petitioner, the '611 patent is or has been involved in the following district court litigations and petitions for *inter partes* review:

Name	Number	Court	Filed
<i>Jawbone Innovations, LLC v. Amazon.com, Inc.</i>	3:22-cv-06727	N.D. Cal.	Nov. 29, 2021 (transferred from E.D. Tex. on Nov. 1, 2022)
<i>Jawbone Innovations, LLC v. Amazon.com, Inc</i>	2:21-cv-00435	E.D. Tex.	Nov. 29, 2021
<i>Jawbone Innovations, LLC v. Apple Inc.</i>	6:21-cv-00984	W.D. Tex.	Sept. 23, 2021
<i>Jawbone Innovations, LLC v. Google LLC</i>	6:21-cv-00985	W.D. Tex.	Sept. 23, 2021
Petition for <i>Inter Partes</i> Review	IPR2022-00604	PTAB	Feb. 22, 2022

⁶ Google LLC is a subsidiary of XXVI Holdings Inc., which is a subsidiary of Alphabet Inc. XXVI Holdings Inc. and Alphabet Inc. are not real parties in interest to this proceeding.

Name	Number	Court	Filed
Petition for <i>Inter Partes</i> Review	IPR2022-00889	PTAB	May 16, 2022
Petition for <i>Inter Partes</i> Review	IPR2022-01085	PTAB	Jun. 3, 2022
Petition for <i>Inter Partes</i> Review	IPR2022-01495	PTAB	Sept. 2, 2022
Petition for <i>Inter Partes</i> Review	IPR2023-00286	PTAB	Nov. 28, 2022
Petition for <i>Inter Partes</i> Review	IPR2023-00285	PTAB	Nov. 28, 2022
<i>Jawbone Innovations, LLC v. Google LLC</i>	3-23-cv-00466	N.D. Cal.	Feb. 1, 2023
<i>Jawbone Innovations, LLC v. Meta Platforms, Inc. d/b/a Meta</i>	6-23-cv-00158	W.D. Tex.	Feb. 28, 2023
<i>Jawbone Innovations, LLC v. ZTE Corp.</i>	2-23-cv-00082	E.D. Tex.	Feb. 28, 2023
<i>Jawbone Innovations, LLC v. Panasonic Holdings Corp.</i>	2-23-cv-00081	E.D. Tex.	Feb. 28, 2023
<i>Jawbone Innovations, LLC v. Guangdong Oppo Mobile Telecomm.</i>	2-23-cv-00079	E.D. Tex.	Feb. 28, 2023
<i>Jawbone Innovations, LLC v. LG Elecs., Inc.</i>	2-23-cv-00078	E.D. Tex.	Feb. 28, 2023
<i>Jawbone Innovations, LLC v. HTC Corp.</i>	2-23-cv-00077	E.D. Tex.	Feb. 28, 2023

Name	Number	Court	Filed
<i>Jawbone Innovations, LLC v. Sony Elecs., Inc.</i>	2-23-cv-01161	D.N.J.	Feb. 28, 2023

C. Lead And Back-Up Counsel Under 37 C.F.R. § 42.8(b)(3)

Petitioner provides the following designation of counsel.

Lead Counsel	Back-Up Counsel
Erika H. Arner (Reg. No. 57,540) erika.arner@finnegan.com Finnegan, Henderson, Farabow, Garrett, & Dunner LLP 1875 Explorer Street, Suite 800 Reston, VA 20190-6023 Tel: 571-203-2700 Fax: 202-408-4400	Daniel C. Cooley (Reg. No. 59,639) daniel.cooley@finnegan.com Alexander M. Boyer (Reg. No. 66,599) alexander.boyer@finnegan.com Finnegan, Henderson, Farabow, Garrett & Dunner, LLP 1875 Explorer Street Suite 800 Reston, VA 20190-6023 Tel: 571-203-2700 Fax: 202-408-4400 Kevin D. Rodkey (Reg. No. 65,506) kevin.rodkey@finnegan.com Finnegan, Henderson, Farabow, Garrett & Dunner, LLP 271 17th Street, NW Suite 1400 Atlanta, GA 30363 Tel: 404-653-6400 Fax: 202-408-4400 Christina Ji-Hye Yang (Reg. No. 79,103) christina.yang@finnegan.com Finnegan, Henderson, Farabow, Garrett & Dunner, LLP 901 New York Avenue NW Washington, DC 20001-4413 Tel: (202) 408-4000 Fax: (202) 408-4400

D. Service Information

Please address all correspondence to lead and back-up counsel at the addresses shown above and Google-v-Jawbone-IPRs@finnegan.com. Petitioner also consents to electronic service by e-mail.

E. Conclusion

Petitioner has established a reasonable likelihood of prevailing with respect to the challenged claims and requests the Board institute *inter partes* review and cancel each challenged claim as unpatentable.

Date: July 7, 2023

Respectfully submitted,

/Daniel C. Cooley/

Daniel C. Cooley, Back-up Counsel
Reg. No. 59,639

CERTIFICATE OF COMPLIANCE

The undersigned hereby certifies that the foregoing **Petition for *Inter Partes* Review** contains 13,983 words, excluding those portions identified in 37 C.F.R. § 42.24(a), as measured by the word-processing system used to prepare this paper.

/Daniel C. Cooley/
Daniel C. Cooley, Back-up Counsel
Reg. No. 59,639

CERTIFICATE OF SERVICE

The undersigned certifies that the foregoing **Petition for *Inter Partes* Review** was served on July 7, 2023, by FedEx Priority Overnight at the following address of record for the subject patent. The associated **Exhibits 1001-1006, 1008-1021** and the **Power of Attorney** were also served on July 7, 2023.

Mark Leonardo
NUTTER MCCLENNEN & FISH LLP
Seaport West
155 Seaport Boulevard
Boston, MA 02210

A courtesy copy has also been mailed to litigation counsel for Patent Owner at:

Peter Lambrianakos
Fabricant LLP
411 Theodore Fremd Avenue,
Suite 206 South
Rye, New York 10580

/Lisa C. Hines/

Lisa C. Hines
Senior Litigation Legal Assistant
FINNEGAN, HENDERSON, FARABOW,
GARRETT & DUNNER, LLP