

استخراج بازنمائی‌های مقاوم مبتنی بر مدل فیزیولوژی گوش انسان برای سیستم‌های بازشناسی خودکار گفتار

کارو لوکس^۲
lucas@ece.ut.ac.ir

خسرو حسین‌زاده^۱
hoseinzadeh@ce.sharif.edu

سید کمال‌الدین غیائی^۱
ghiathi@mehr.sharif.edu

امین فاضل دهکردی^۱
afazel@mehr.sharif.edu

^۱ دانشکده کامپیوتر دانشگاه صنعتی شریف
^۲ دانشکده برق و کامپیوتر دانشگاه تهران

چکیده: همیشه موجودات زنده بهترین الگوها برای تکنولوژی بشر بوده‌اند. بر خلاف سالهای متمادی تحقیق و مطالعه، هنوز به طور واضح معلوم نیست که سیستم شنوایی انسان به چه شکل با دامنۀ وسیعی از اصوات عمل می‌کند [1] و همچنین مشخص نیست که انسان به چه شکل گفتار را در حضور صداهای دیگر می‌تواند بفهمد [2]. یک سیستم بازشناسی گفتار که بتواند به شکل موثری کارایی سیستم شنوایی انسان را داشته باشد باید تضمین کند که اطلاعات پردازش شده گفتار دارای بازنمائی‌هایی شبیه به سیستم شنوایی انسان می‌باشد. معمولاً در بسیاری از سیستم‌های بازشناسی گفتار امروزی بازنمائی‌های استخراج شده بیشتر بر اساس سیستم تولید گفتار می‌باشد و به نظر می‌رسد که مدل کردن فیزیولوژی گوش انسان می‌تواند در استخراج ویژگی‌های مؤثرتر مفید باشد. در این مقاله با الهام گرفتن از فیزیولوژی گوش انسان قسمتی از گوش داخلی انسان مدل شده است. نتایج حاصل از آزمایشات بر روی دادگان اعداد متصل انگلیسی نشان دهنده آن است که استفاده از استخراج ویژگی‌های مبتنی بر یافته‌های فیزیولوژیک گوش انسان می‌تواند دقت سیستم‌های بازشناسی گفتار را به خصوص در شرایط نویزی بالاتر ببرد.

کلمات کلیدی: بازشناسی مقاوم گفتار، بازنمائی‌های مقاوم، فیزیولوژی گوش انسان.

۱- مقدمه

بازشناسی خودکار گفتار یکی از تکنولوژی‌های در حال ظهور برای افزایش رابطۀ بین انسان و ماشین می‌باشد. با وجود اینکه امروزه محصولاتی در این زمینه به وجود آمده است ولی هنوز استفاده همه منظوره از این تکنولوژی عملی نشده است. مهمترین مشکل وجود عدم تطابق بین محیط آموزش و محیط آزمایش می‌باشد. همه سیستم‌های بازشناسی خودکار گفتار دارای یک طبقه‌بندی کننده می‌باشند که بر روی سیگنالهای گفتار (معمولاً فاقد نویز) آموزش داده می‌شود با این هدف که بتواند محتوای آوایی گفتار را برای بازشناسی گفتار (معمولاً نویزی) یاد بگیرد. محیط

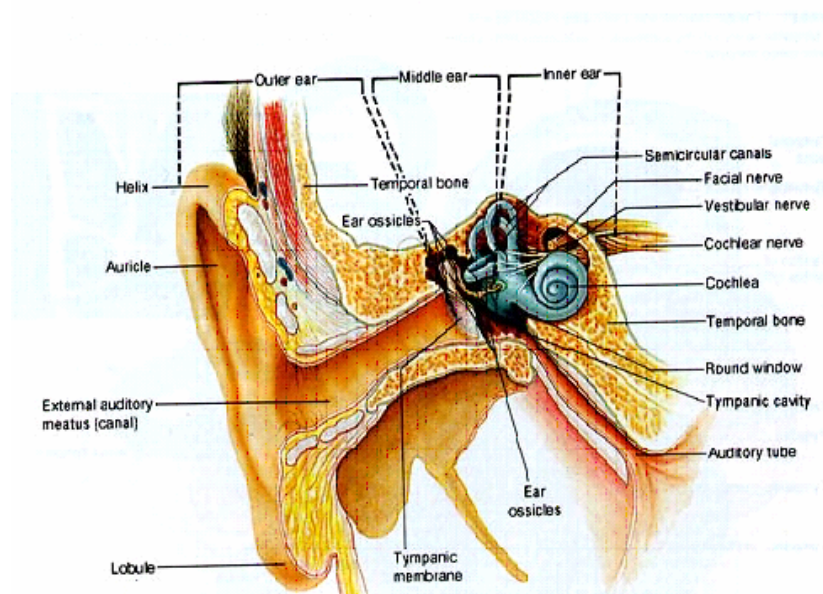
آزمایش ممکن است با محیط آموزش تفاوت‌هایی از نظر تنوع گوینده‌ها، نویز پس‌زمینه و اثرهای کانال انتقال گفتار داشته باشد.

همانطور که گفته شد مشکل بزرگ سیستم‌های بازشناسی گفتار مقاوم نبودن آنها نسبت به نویز است که باعث شده نتوان به طور فراگیری از این سیستم‌ها استفاده کرد. از آنجائی که بازشناسی گفتار در انسان در مقایسه با سیستم‌های بازشناسی گفتار امروزی در مقابل نویز بسیار مقاوم می‌باشد، به نظر می‌رسد که یکی از راه‌های مقاوم کردن این سیستم‌ها الهام گرفتن از بیولوژی موجود در طبیعت می‌باشد. بسیاری از سیستم‌های بازشناسی از بازنمائی MFCC استفاده می‌کنند [3] که یک الگوریتم مبتنی بر فیلتر بانک می‌باشد که فیلترهای آن بر روی فرکانسهای لگاریتمی در دامنه Mel قرار گرفته‌اند. از طرفی باندهای فرکانسی هر یک از این فیلترها به وسیله مکان هر یک از آنها تعیین می‌گردند و از روی یافته‌های بیولوژیکی نمی‌باشند.

در این مقاله یک روش جدید با الهام گرفتن از گوش انسان و مدل کردن حرکت‌های مکانیکی غشای پایه^۱ و سلولهای مژکدار^۲ برای دریافت فرکانسهای مختلف صدا بر اساس یافته‌های فیزیولوژی ارائه می‌شود.

۲- فیزیولوژی سیستم شنوایی انسان

گوش اندام حسی تبدیل کننده صدا است. گوش انسان از سه بخش گوش بیرونی، گوش میانی و گوش درونی تشکیل شده است (شکل ۱). گوش بیرونی امواج صوتی را به پرده صماخ، که در حد فاصل گوش بیرونی و گوش میانی است، راه می‌دهد. امواج صوتی پرده صماخ را به ارتعاش در می‌آورند. ارتعاشهای پرده صماخ از طریق سه استخوان کوچک موجود در گوش میانی به گوش درونی انتقال می‌یابند [4].

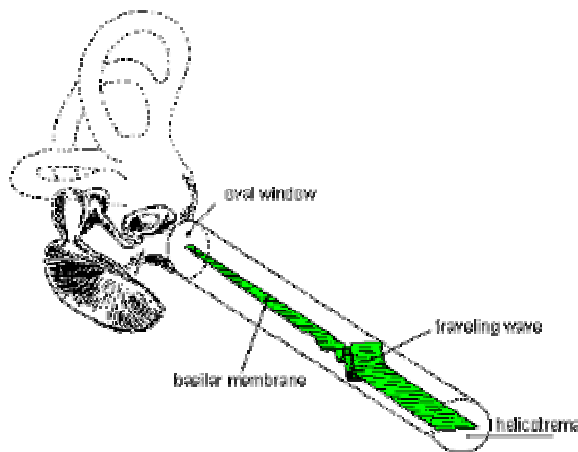


¹ Basilar Membrane

² Hair Cells

شکل (۱). ساختار کلی گوش

گوش درونی حفره^۱ پر از مایعی است که حلزونی گوش^۱ نامیده می‌شود. حلزونی گوش انسان دارای یک ساختار مخروطی شکل است که فنروار در اطراف یک هسته^۲ استخوانی پیچیده شده است. نیروی مکانیکی تقویت شده که از طرف گوش میانی به گوش داخلی فرستاده می‌شود بلافاصله توسط حلزونی گوش به فشار هیدرولیکی تبدیل می‌شود. این فشار هیدرولیکی، حرکت حاصل را به مجرای حلزونی گوش^۲ و اندام کورتی^۳ میرساند. در سراسر طول حلزون، نوعی غشاء امتداد یافته است که غشای پایه نامیده می‌شود. این غشاء پرده^۳ خم‌پذیری است که هرگاه صدائی به پرده^۳ صماخ برخورد کند آن را به ارتعاش در می‌آورد. ارتعاش غشای پایه به صورت جابجائی موج‌مانندی است که آن غشاء را تغییر شکل می‌دهد.



شکل (۲). حرکت موج جابجائی در طول غشای پایه

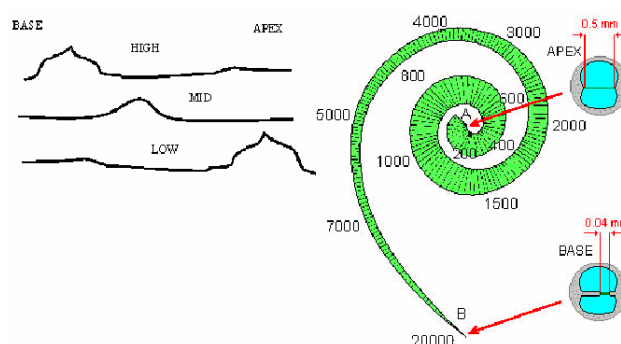
هنگامی که موج جابجائی در طول غشای پایه پیش می‌رود (شکل ۲)، بر دامنه^۲ آن افزوده می‌شود تا اینکه در نقطه‌ای به اوج برسد. اما وقتی که موج جابجائی غشای پایه از نقطه^۲ اوج جلوتر می‌رود، دامنه^۲ آن کاهش می‌یابد. محل اوج دامنه^۲ موج جابجائی در غشای پایه، بستگی به فرکانس صوت دارد. اصوات با فرکانس بالا سبب می‌شوند که موج جابجائی در قاعده^۲ حلزون به اوج خود برسند در حالیکه اصوات با فرکانس پایین سبب می‌شوند که موج جابجائی در رأس حلزون به اوج خود برسد (شکل ۳). بنابراین، طرح مکانیکی غشای پایه چنان است که به اصوات دارای فرکانسهای مختلف با تغییر دادن شکل خود پاسخ می‌دهد [4]. تغییر شکل غشای پایه را سلولهای مژکدار به پتانسیل کار تبدیل می‌کنند. وقتی که غشای پایه مرتعش می‌شود و به اوج تغییر شکل خود می‌رسد، سلولهای مژکدار نسبت به جسم آن سلول خم می‌شوند و از این طریق عصب شنوائی،

¹ Cochlea

² Cochlea duct

³ Organ of Corti

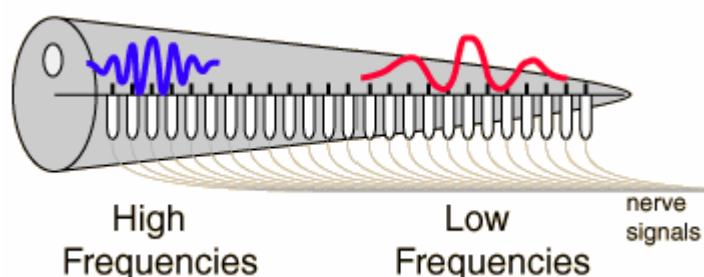
پتانسیل کار تولید می کند که به مغز منتقل می شود و اطلاعات مربوط به فرکانس صوت را به مغز می رساند.



شکل (۳) حساسیت حلزون گوش نسبت به فرکانسهای مختلف

۳- مدلسازی حرکت های مکانیکی غشاء پایه و سلول های مژکدار

از آنجائی که غشاء پایه و سلول های مژکدار با همکاری یکدیگر فرکانسهای مختلف صدا را به انرژی مکانیکی تبدیل می کنند با تقریب خوبی می توان هر قسمت از آن را به صورت گسسته به شکل یک دیافراگم در نظر گرفت. به این دیافراگم ها در واقع می توان به شکل سیستم های نوسانی میرا که توسط یک نیروی خارجی تحریک می گردند، نگاه کرد.



شکل (۴) مدل کردن حرکت های مکانیکی غشاء پایه و سلول های مژکدار با دیافراگم

مدل هر یک از این سیستم های نوسانی میرا را می توان با رابطه (۱) نمایش داد [5]، [6]،

[7] و [8]،

$$\ddot{x} + \frac{b}{m} \dot{x} + \omega_0^2 x = \frac{F}{m} \cos \omega t \quad (۱)$$

که در این رابطه b نشان دهنده فاکتور میرائی، ω_0 فرکانس طبیعی هر کدام از این دیافراگم ها و F نیروی خارجی است که توسط صوت با فرکانس ω بر هر یک از آنها اعمال می شود. از آنجائی که این مدلسازی باید به صورت گسسته انجام شود، در اینجا به تعداد ۲۵۶ عدد از این مدلها در نظر گرفته شده است. از آنجائی که هر یک از این مدلها باید به فرکانس خاصی حساس باشند و در آن فرکانس حالت تشدید داشته باشند، برای هر یک از آنها یک باند فرکانسی خاص در نظر گرفته شده است. این باند فرکانسی تعیین کننده مقدار b بر اساس رابطه زیر می باشد،

$$b = \frac{m\omega_0}{Q} \quad (2)$$

که در رابطه فوق Q فاکتوری است که تعیین کننده میزان انرژی از دست رفته است و همچنین k را به عنوان ضریبی تعریف می کنیم که به وسیله رابطه $k = m\omega_0^2$ به جنس هر یک از دیپازونها وابسته است.

برای استفاده از این مدل در استخراج بازنمائی ها باید به ازای تعدادی نمونه صوت شتاب هریک از دیپازونها را به دست بیاوریم و از روی این شتاب و میزان انحراف آنها انرژی را برای تک تک آنها محاسبه کنیم. این انرژی به ازای تک تک دیپازونها می تواند به عنوان بردار بازنمائی برای سیستم بازشناسی گفتار مورد استفاده قرار بگیرد.

شتاب هر یک از این مدلها را بر اساس سرعت و مکان قبلی و همچنین نیروی اعمال شده جدید که از طرف نمونه های صوت بر آنها وارد می شوند، طبق رابطه زیر بدست می آید

$$a = \frac{F - bv_{pr} - kx_{pr}}{m} \quad (3)$$

از روی این شتاب سرعت $v = v_{pr} + at$ و مکان $x = x_{pr} + at$ جدید به دست می آید. سپس انرژی از طریق رابطه $E = \frac{1}{2}mv^2 + \frac{1}{2}kx^2$ به دست می آید و از آن برای استخراج بازنمائی ها استفاده می کنیم.

۴- نتایج حاصل از پیاده سازی

برای نشان دادن تأثیر بازنمائی های مبتنی بر مدل فیزیولوژیک گوش آزمایشگاهی برای مقایسه با بازنمائی های MFCC طراحی شد. برای آزمایش از سیستم بازشناسی گفتار HTK¹ استفاده شد [9]. با استفاده از سیستم بازشناسی مذکور مدل سازی اعداد انگلیسی با استفاده از مدل مخفی مارکوف انجام داده شد. HMM های به کار برده شده از نوع چگالی پیوسته و با تلفیق های گوسی می باشند و پرش بین حالت های آن فقط به صورت چپ به راست مجاز می باشد. هر عدد زبان انگلیسی با یک HMM دارای ۱۶ حالت و ۳ تلفیق گوسی مدل شد. برای آموزش این سیستم از دادگان اعداد متصل زبان انگلیسی که شامل ۱۸ گوینده که هر گوینده ۷۷ رشته عددی را بیان می کند، استفاده شد. هنگام آزمایش از ۲۰۰ رشته عددی که به وسیله گوینده های غیر از گوینده های داده ای آزمون بیان شده اند، استفاده شد. دادگان تست به چهار دسته ۵۰ تائی تقسیم شدند و به هر دسته چهار نوع نویز مختلف با SNR های متفاوت اضافه شد. در جدول (۱) خلاصه نتایج حاصل از آزمایش ها آمده است. همانطور که در جدول مشخص است در اغلب موارد آزمایش بهبودهای در نتیجه حاصل از بازشناسی به دست آمد.

¹ Hidden Markov Toolkit

Word Error Rate								
Noise	Feature	20 db	15 db	10 db	5 db	0 db	-5 db	Overall
Car	MFCC	13.29	24.05	53.16	80.38	93.67	93.67	59.7
	Bio Model	13.29	23.42	55.06	74.05	82.28	93.67	56.96
Exhibition	MFCC	37.75	48.34	59.6	77.48	90.73	91.39	67.55
	Bio Model	33.77	36.42	64.24	78.15	86.09	89.4	64.68
Babble	MFCC	26.95	44.31	72.46	88.02	94.61	99.4	70.96
	Bio Model	23.95	34.73	63.47	76.05	83.23	88.62	61.68
Subway	MFCC	46.75	60.36	71.01	88.17	89.94	90.53	74.46
	Bio Model	47.93	57.4	72.78	83.43	87.57	91.12	73.37

جدول (۱) نتایج حاصل از مقایسه بازنمایی های MFCC با بازنمایی های مبتنی بر مدل فیزیولوژیک

۵- جمع بندی و نتیجه گیری

در این مقاله مدلی بر مبنای یافته های فیزیولوژیک برای استخراج ویژگی جهت سیستم های بازشناسی گفتار معرفی شده است. نتایج بدست آمده حاکی از آن است که مدل ارائه شده برای غشای پایه و سلولهای مژکدار دارای قابلیت هایی برای بهبود نتایج بازشناسی گفتار در حضور نویز می باشد. همچنین می توان بیان داشت که رویکردهای فیزیولوژیک و مدل کردن گوش انسان نسبت به استخراج ویژگی های مرسوم و سنتی دارای نتایج بهتری می باشد.

۶- مراجع

- [1] D.W. Massaro, *"Speech Perception by Ear and Eye: a Paradigm for Psychological Inquiry"*, Lawrence Erlbaum Associates, Hillsdale, New Jersey, 1987.
- [2] B. Dodd & R. Campbell, Eds., *"Hearing by Eye: The Psychology of Lip-Reading"*, Lawrence Erlbaum Associates, Hillsdale, New Jersey, 1987.
- [3] M. D. Skowronski and J. G. Harris *"INCREASED MFCC FILTER BANDWIDTH FOR NOISE-ROBUST PHONEME RECOGNITION"*, Proc. IEEE Int. Conf. Acoustic, Speech, Signal Processing, (Orlando, FL, U.S.A.), May 2002.
- [4] R. Rhoades and R. Pflanzner, *Human Physiology*, 3rd Edition, Saunders College Publishing, 1996.
- [5] R. Chabay and B. Sherwood, *Matter and Interactions*, John Wiley & Sons, 2002.
- [6] C. H. Edwards, D. E. Penney, *Differential Equations and Boundary Value Problems: Computing and Modeling*, 2nd Edition, Prentice Hall, 1999.
- [7] J. Stewart, *Calculus, Early Transcendentals*, 4th Edition, Brooks Cole, 1999.
- [8] H. D. Young, R. A. Freedman, *University Physics with Modern Physics with Mastering Physics*, 11th Edition, Benjamin Cummings, 2004.
- [9] S. Young, G. Everman, D. Kershaw, G Moore, J. Odell, D. Ollason, D. Povey, V. Valtchev, P. Woodland, *The HTK Book (for HTK Version 3.2)*, Cambridge University, 2002.