



BIOENERGETIC FEATURES OF PHONATION - THE SUBSTRATE OF GLOBAL MUSICALITY

*Dionysios Politis¹, Georgios Kyriafinis², Jannis Constantinidis², Nektarios Paris³

¹ Software and Interactive Technologies Lab, School. of Informatics, Aristotle University of Thessaloniki, Greece.

² 1st Academic ENT Dept., AHEPA General Hospital, Aristotle University of Thessaloniki, Greece.

³ Dept. of Music Science and Art, University of Macedonia, Greece.

ABSTRACT

The energy levels of phonation constitute indicators revealing the active effects of bodily, kinesthetic and audio communication power levels interrelation with the climatic and cultural habitat of distinct population groups. The transformation of the bioenergetic factors into characteristic expression patterns and rhythmic pitched articulation sequences within the discourse of speech or music communication is examined in both unrestrained vocalized expression of feelings, freely or loudly, and also in organized choir singing; the latter retains strong elements based on the historical, cultural and social development cues of the ethnic groups staging public performances. The generic character of these voicing features is correlated with linguistic and bioenergetic profiling of living communities.

KEYWORDS: Singing Voice, Linguistics, Ethnomusicology, Otorhinolaryngology and Physiology.

I. INTRODUCTION: THE LINGUISTIC DIVERGENCE

Human language is studied intensively by IT systems as it is a focal point for global communication in its synchrony and diachrony, in both theoretical and practical terms; for example, the excessive spread of mobile telephony with smartphones exhorting to over 3 billion people means that the operating environment has to be superfluously adapted to the speaking languages of a global audience. It is estimated that currently some 7,200 languages are spoken daily, and many of them are voiced in a manner not particularly changed for centuries or even millennia (Anderson, 2012).

Of course, the most influential languages spoken world wide, by as much as 97% of contemporary people, are less than 300. Unfortunately, this reduced number, that leaves many dialects, regional languages or vernacular forms without official recognition, does not occur as a linguistic research finding; on the opposite direction, it has to do with the languages supported by computers, electronic devices and smartphones (Loh and Kanai, 2016). It is exactly this domain, the one of mobile devices, the closest AI support technology to one's everyday transactions, fully exerting the HCI communication potential by advancing multichannelled neurocognition experiences in each user's mother tongue (Wanderley and Orio, 2002). For instance, very popular apps in everyday use are standalone or on-line programs for reading/viewing the news in one's native language and working environment, translation tools aiding on the fly conversations held at any possible pairs of supported languages, or listening/viewing to music paraphernalia (Statista, 2019). Moreover, in their written communication users rarely make orthographic mistakes, not due to their improved grammatical or syntactic potential, but merely by using AI tools for spelling and grammar autocorrection.

This magnitude of oral and written communication channels comprises, nevertheless, in neurocognitive terms, a huge range, even though that "merely" some 300 languages are supported out of 7200; even further, the mechanism for how such a subtlety, in detecting and analyzing patterns of the oral-aural channel of communication functions, i.e. the phonological domain, has not been fully explained (Politis et al., 2016). In any case, its range is indeed astonishing, as languages spoken by hundreds of millions of people are producing variants that become dialects of widely spread languages, e.g. Spanish or Portuguese as a mother tongue in various countries, well distanced geographically each within another (Boas, 1940). Sometimes these idiomatic divergent propensities are promoted to independent languages, usually not due to well-founded scientific grounds, but rather based on political relationships.

Languages are deployed in three levels, as far as neural processes affect mental calculation stereotypes and working memory modules.

- I. The first, the most fundamental, is the spoken level. Languages are using extensively the oral-aural communication channel and their developmental mechanism relies heavily on the existence of a perceptive acoustic signal triggering mechanism (Kyriafinis et al., 2017).
- II. If acoustic stimuli, however, serve for the instigation of discernible cognitive models, it is the second upper level that makes the linguistic substrate to become the acting mechanism for advanced cognitive achievements in literature, poetry, philosophy, etc. Languages that have enlarged

their domain with words having many notional shades are expected to develop semantic complicacies that promote high-level thinking. Conversely, uncontrolled enlargement with superficially picked-up terms, as a side-effect of uncontrolled globalization, lurks the creation of semantically non-associative clusters of words that rather lead to a perilous unclarity in homonymity, and even further, to obscureness and muzziness of meaning when a language extends its self to new areas of human activity (Rappleve, 2012). Schooling, an intense, long-lasting educational activity, aims to develop the grammatical and syntactical potential of its trainees to high levels (Department of Education, 2010). Unfortunately, this top-down activity may also become interconnected with affluent political biases, since it is the most dynamic, unrejoined investment of Governance over the youth zone of its subjects (Carney, 2012). Even further, this level is associated with the writing systems each language may use. There have been reported several occasions, at least in the last couple of centuries, where languages shift from on alphabet to the other (Lightner, 1972). However, the main rift, in neurocognitive terms does not stand between alphabets, as it does between alphabetic and ideographic variants of speech communication: while the former compose words out of their phonological constituent units, the latter employ numerous pictorial representations of concepts or things (Morris and Adamson, 2010). Neuroscientists recently have determined how differently the cerebrum responds when humans read translationally "equivalent" propositions in these two systems. Not surprisingly, these two ways of linguistic communication, more or less equally dominant in numbers of their speakers, predispose somehow the accentuation patterns to be used: "dynamic" for the former, "tonal" for the latter.

- III. The third level, the most difficult to represent mathematically, develops as a sophisticated superstratum of linguistic exploitation (Chomsky, 2000). It examines concepts beyond the semantic, morphosyntactic or lexicostatic level of linguistic development. In many occasions it surpasses the etymological analysis of voiced communication and seeks "particles" of speech, usually in a multilingual representation, that have association with certain neuroscientific patterns or features of cerebral activity (Takayama, 2009). They may, as such, compose linguistic atlases based on discernible communication characteristics distinctive of sex, phylum, or topicalization for religious and other such reasons.

In acoustic terms, the prosodic systems used by each language serve as an essential gateway for the music deployed culturally in that region. For starters, languages are categorized into two, mutually exclusive to a great extent variants, as far as their accentuation is concerned: the tonal or contour languages, and the dynamic ones (Ramat, 1988).

For reasons that are beyond the scope of this research, the former have prevailed in Far East, while the latter in what we call "West".

II. THE HISTORICAL - CULTURAL PERSPECTIVE: MODES OF LISTENING AND STYLES OF MUSIC

While the linguistic continuum reprimands unification and remains prolifically scattered, the musical sphere of communication does not seem to be so in both its diachrony and synchrony. As a matter of fact, musical culture and technology

spreads equivalently with flourishing advanced states of social development, as superior civilizations engulf these elements in their complex legal, scientific, political and religious organizations (Huron, 2001).

Ancient civilizations endorsed their musical perception, in such a way, that it has been more or less scripturally documented, as is the case of the Sanskrit Vedas, 15th to 10th century BC (Pingle, 1962). In the Chinese antiquity, somewhere in the 10th century BC, even a governmental State Music Bureau was founded, supervising imperial music production and distribution (Hays, 2016). In Egypt, music practices (Lentz, 1961) have been depicted in paintings of enormous artistic added value (like the Lute and double pipe players from a painting found in the Theban tomb of Nebamun, c. 1350 BC), while some musical inscriptions have been more or less deciphered (Sachs, 1921; Hickmann, 1960). And when it comes to Ancient Greek Music, a variety of documents, treatises and inscriptions cater for a rather sufficient preamble for its reproduction and theoretical foundations (Spyridis, 1999; Pöhlmann and West, 2001).

However, in antiquity there seemed to be considerable restrictions on how the “West” perceived its universe of discourse (Fig. 1).

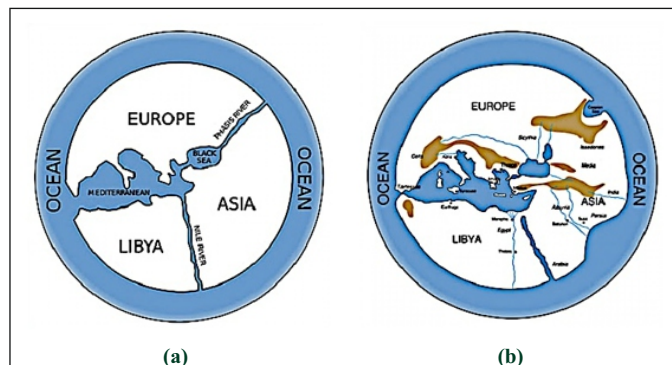


Fig 1. (a) The world, according to the pre-Socratic philosopher Anaximander of Miletus (c.610-546 BC) **(b)** His disciple, historian and geographer Hecataeus of Miletus (c.550-476 BC), one of the first to mention the Celtic people, presented a more augmented picture, based on the Greco-Persian wars communiqué. Pictures cropped from Wikimedia.

By the time of Eratosthenes of Cyrene (276-195 BC), the “father of Geography”, the Indo-European cultural continuum was more or less moulded (Wellesz, 1957), due to the conquests of Alexander the Great (Barker, 2009). Even the Britannic isles were clearly depicted (Fig. 2).

Out of these images, it becomes evident that the ancient civilizations in the “West” developed around the Mediterranean basin, characterized by its mild climate, stretching to India via Middle East (Vendries, 1999). The first hint for an equally populous civilized world far East, appears in Ptolemy’s Geography, circa 150 AD (Fig. 3).

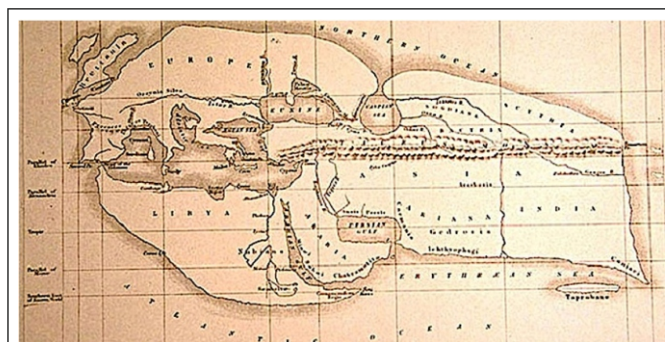


Fig. 2. 19th century reconstruction of Eratosthenes’ map for the known world, according to his contemporary Alexandrian scholars, 3rd century BC. Image cropped from Wikimedia.



Fig. 3. Reconstituted in the 15th century Ptolemy’s world map of the 2nd century AD, indicating “Sinae” (China) far right, beyond the island of “Taprobane” (Ceylon or Sri-Lanka) and the “Aurea Chersonesus” (Southeast Asian Peninsula). Image cropped from Wikimedia.

Some centuries later, the Arab geographer and historian Muhammad al-Idrisi (1100-1165 AD), living in Palermo, Sicily, incorporated knowledge from merchants and travellers to Africa, Indian Ocean and the Far East, managing thus to depict the entirety of the Eurasian continent. However, information on the sub-Saharan African continent seems to be missing (Fig. 4).

Also, the huge bulk of the Himalaya seems to hinder cross-cultural osmosis between Far East and the pioneer, after the 12th century, region in scientific discovery, the “West” (Hammann and Harwood, 1976).

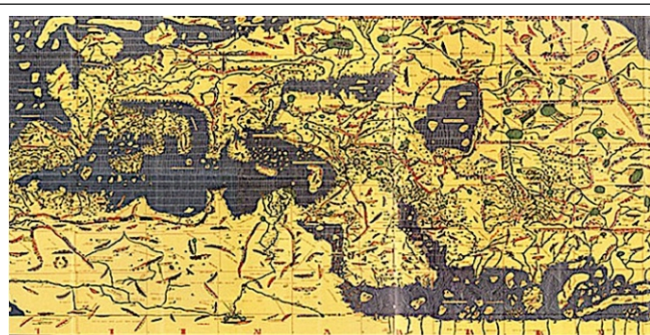


Fig. 4. The Tabula Rogeriana, formed by Muhammad al-Idrisi for Roger II of Sicily in 1154 AD. The Eurasian continuum is nearly completely depicted. Image cropped from Wikimedia.

Nevertheless, the well-structured reality of this “Parmenidian” unity, that constitutes the wonder of contemporary musical culture did not appear as a perfunctory courtesy (Gregory and Varney, 1996); on the contrary, it evolved as a concluding summary of influences out of a series of contributions from acknowledged scholars, scientists and musicians (Rothstein, 1995).

Indeed, the ethnological, historical and linguistic continuity of the Western part of the Roman Empire waned as successive waves of Germanic tribes started moving to the appealing warm waters of the Mediterranean basin. Visigoths, Ostrogoths, Vandals to name the first; then Burgundians, Swabians; finally, Franks and Lombards in the 6th century. And afterwards, Angli, Normans, Danes. The same time, in the beginning of the 7th century, Arab invaders and Berber pirates move towards the Iberian Peninsula and Southern France, filling in the vacuum of power and peoples.

As a new ethnological substrate was formed, an intensive effort for the Christianization of the “West”, up to its utmost Northern districts commenced (Elias, 2000). On the contrary, missions to the Arabic South were restricted. Monastic communities were formed, with the Roman districts of Gaul like Provence, Brittany and Arles excelling in providing instruction on how the new ethnic societies should be organized and administered. During the 5th, 6th, and 7th centuries, the influence exerted by the abbey’s associated schools, scriptoria, libraries and administrative centers for cultivating vast areas of land in copartnership is enormous (Collins and Price, 2003).

The same initiatives are developed in Italy, e.g. in Monte Cassino, in Spain, and in some restricted areas of North-Western Africa. For higher levels of education, episcopal seminaries are founded, providing collegiate edification (Herrin, 1989). In the first place, the main focus is on Theology; gradually, however, at the 7th and 8th centuries, they expand their curriculum to a variety of educational, cultural and artistic themes, contributing to the Carolingian Renaissance.

In the British Isles shine the schools of Kent, Wessex, Northumbria. In restructured

Italy, apart from Monte Cassino, Pavia, Rome and Bobio are revitalized. As a result the first University of the West is founded in Bologna, circa 800 AD.

The establishment of Universities is a ground-breaking innovation of the 13th century. Their number increases drastically in the 14th and 15th centuries. All “Western” countries, from the Iberian Peninsula and Italy up to the frozen shores of the Northern Seas, boast of having advanced scientific learning in many branches (Carol, 1964); Professors and students are associated with the Ecclesiastical habitat, but their way of living is diverted from the clergy's one (Hayas, 1953).

Clearly, the previously hindered “West” is reshaped and leads the world scientific continuum (Schank and Abelson, 1976). However, it is not that much “West”, as it is “North”, or better say, “North-West”.

In the “East”, things were kept more under control (Ciggaar, 1996). The administrative center, Constantinople, although it gradually lost sovereignty over its Middle East provinces, developed a sophisticated Greek-oriented civilization, both in scientific, legal and technical terms (Evans and Wixom, 1997). Already, the first University in the “East” was founded in 425 AD by Imperial decree (Runciman, 1993).

In a cultural and linguistic perspective, the Greek language, although not widely spoken nowadays, played a vital role in the history of the “Western” world and Christianity (Devine and Stevens, 1994). Commencing somewhere around the 10th century BC, the canon of ancient Greek literature includes seminal works in the “Western” world's canon such as the epic Homeric poems, i.e. Iliad and Odyssey. Throughout Antiquity and Medieval ages, it served as the language in which many of the foundational texts in which science, principally astronomy, mathematics and logic were deployed (Forbes, 1933). Much of the Western philosophy is based on the works of Plato, Aristotle, Pythagoras; the New Testament of the Bible was written in Koiné Greek. Together with the Latin texts and traditions of the Roman world, the study of the Greek texts and society of antiquity constitutes the discipline of Classics.

During antiquity, Greek was a widely spoken lingua franca in the Mediterranean world and many places beyond, nearly as far as India. It would eventually become the official parlance of the Byzantine Empire, i.e. the Eastern Roman Empire (Obolensky, 2009), and develop into Medieval Greek. It was a language that was spoken uninterruptedly for millenniums in the region of South-Eastern Europe, the East Mediterranean Basin and Middle East.

After the 7th century, the Arab expansion has restrained its use in Middle East and Northern Africa. However, its semantic insight, musical and mathematical substratum seems to have been transplanted somehow into the Arabic culture and civilization (Trehub, 2000).

The same time, in the 9th century AD, Greek culture and civilization spreads to the Northern part of Eurasia, to the regions inhabited by predominantly Slavic populations. It is not very clear the vindication for the timing of this osmosis (Nicol, 2002). As it has been recorded for Western Europe, people from the North move towards the Southern parts of Eurasia, while previously harsh for living, frozen zones of the upper parts in Europe develop gradually urban societies and superior technical civilizations. As these areas were incorporated into the Christian World, the spoken dialects of the Proto-Slavic language acquired.

The Proto-Slavic language existed until around AD 500. By the 7th century, it had broken apart into large dialectal zones.

There are no reliable hypotheses about the nature of the subsequent breakups of West and South Slavic. East Slavic is generally thought to converge to one Old Russian or Old East Slavonic language, which existed until at least the 12th century (Hayes, 1984).

Linguistic differentiation was accelerated by the dispersion of the Slavic peoples over a large territory, which in Central Europe exceeded the current extent of Slavic-speaking majorities (Comrie and Corbett, 1993). Written documents of the 9th, 10th, and 11th centuries already display some local linguistic features that commenced extending well over the Urals deep into Northern Asia (Kortlandt, 1994).

The details of the cultural and linguistic development in Europe are indeed presented in some extent in this chapter (Perry, 2012); this is not exempted in a supercilious manner or attitude (Mersenne, 1635). On the contrary, it is accomplished in an exemplary approach because between 1492 and 1914 AD this metropolis colonized 80% of today's world.

Had not this been the case, it would be not easy to elucidate why countries far distanced from each other, like for instance Mexico, Philippines or Spain, without having anymore political bias over each other, share a rather common musical and overall cultural substrate (Pope, 1986). As a result, for instance, their language remains unified (Botinis, 2000) in the sense that the cluster of Spanish speaking countries comprises, along with Chinese Mandarin and Hindi (Shosted et al., 2012), the biggest triad of mother tongue languages.

III. THE BIOLOGICAL PERSPECTIVE: PHONATION AND ANATOMY OF SOUND SOURCES

Human speech and singing are considered to be acoustic signals with a dynamically varying structure in terms of frequency and time domain (Chhetri et al., 2012). Generally speaking, voice sounds are in the broader sense all the sounds produced by a person's larynx and uttered through the mouth (Steriade, 1997). They may be speech, singing voices, whispers, laughter, snorting or grunting sounds, etc.

No matter how these sounds are produced and what communication purposes they may serve, they are categorized to:

- Voiced sounds that are produced in a person's larynx with a stream of air coming from the lungs and resonating the vocal cords (Chan, 2001). This stream continues to be modulated through its passage (Chhetri et al., 2009) from the pharynx, mouth, and nasal cavity, resulting to an utterance in the form of speech or song.
- Unvoiced sounds. Singing or utterance would be incomplete if unvoiced sounds were not produced. They do not come as a result of a vibration from the vocal cords (Wolfe et al., 2009), but as partial obstruction of the airflow during articulation.

The human ability to communicate relies on our capacity to coherently set up sequences of sounds that encode acoustically logical propositions (Bar-Hillel, 1964). When voicing produces musical or singing sounds, then the articulated sounds of speech communication (Dellwo and Wagner, 2003) are enriched with phonation tuned up to melodic, definite pitches that are called notes or tones.

How human voicing is produced and manipulated is indeed a very complex phenomenon (Levitin and Menon, 2003). Not all people however produce the same notes in an undeviating manner (Alku and Vilkman, 1996). As a matter of fact, there is great variability in the timbre of human voicing, the frequency range for singing, which is designated at least at six distinguished zones (Carey et al., 2007), thus far, and the interconnection with the mother tongue and the spoken languages of the choirists (Balkwill et al., 2004).

The complex way in which human instruments work together to produce articulation (Dellwo et al., 2015), may be seen in Fig. 5. For programming purposes, a simulation of this multi-muscularly driven region that plays a key-role in the dynamic formation of the vocal tract activities (Erath and Plesniak, 2010) is shown in Fig. 5(b).

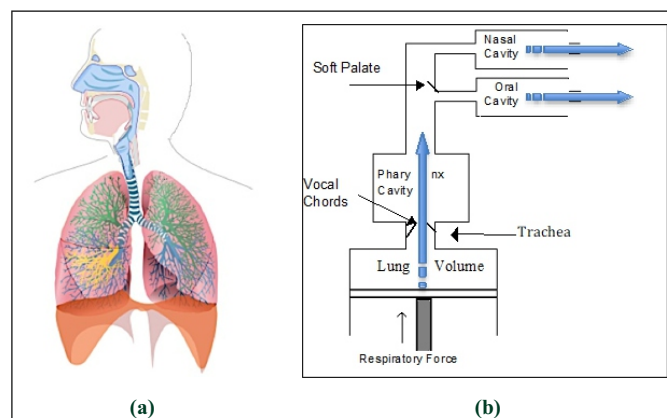


Fig. 5 Schematic representation of the vocal tract in coronal cross section. (a) The respiratory system and the upper vocal tract mechanism for phonation, in medical terms, from Wikimedia commons. (b) A computer programming simulation for synthetic musical vocalization. Sonification comes as output from both the nasal and oral cavities.

Scientists have proposed various source and filter models (Selbie et al., 2002), which gives a rather sustainable approximation of the speech production process from an articulatory perspective (Horacek et al., 2009). Indeed, the mechanism for vocalization is better understood from an acoustic point of view by considering a set of variable sound sources (Zhang, 2016), coupled together to form the complex structure seen in Fig. 6 (Miri, 2014). As speech communication, and especially singing, are operatively related with the ability to express thoughts and feelings by articulate sounds (Juslin and Laukka, 2003), it becomes evident that it is an encephalic process related to the ability to myologically cooperate with the primordial sources of voicing oscillation (Finnegan et al., 2000), seen in Fig. 6.

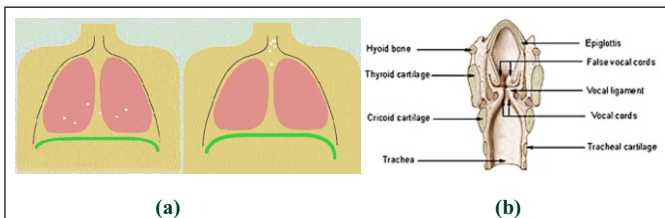


Fig. 6 (a) Sequences of the breathing mechanism, denoting the relative position role of the diaphragm and its role for the formation of airflow directed to the vocal box. **(b)** The Trachea and Larynx sections analytically. Pictures cropped via public domain Wikimedia commons.

The motor functions of phonation at this level are performed by muscular movement of the cricoid, a ring-shaped cartilage of the larynx, the thyroid cartilage which forms the Adams apple, the epiglottis, a flap cartilage at the root of the tongue, which is depressed during swallowing to cover the opening of the windpipe, and finally the arytenoid, a pair of cartilages at the back of the larynx (Holmberg et al., 1989).

The flow of air in this region results in localized drop pressures, according to the Bernoulli effect, i.e. the principle of hydrodynamics stating that an increase in the velocity of an airstream results in a decrease in pressure (MacLarnon and Hewitt, 1999). As a result the vocal chords snap shut. This process continues in a rather pseudo-periodic manner with the muscles setting vocal chord position and tension to appropriate levels so to adjust efficient respiratory force that maintains greater air pressure below the vocal chords than above them (Krebs et al., 2012). In conclusion, the vibratory action of the vocal chords takes place due to the combined action of muscular settings, vocal chord elasticity (Hsiao et al., 2002) that unfortunately decreases with ageing, differential air pressure exerted across the epiglottis, and application of aerodynamic forces (Garnier et al, 2010; Politis et al., 2017).

Once the airflow leaves the voice box, the overall energy source for phonation, it is directed via the pharynx to the oral and nasal cavities. As expected, the shape of the larynx and pharynx, along with its variable size play a primordial role for the production of speech sounds. It is this part of the body that controls the range of pitch or the type of tone that a person can produce (Sundberg, 1987).

Indicatively, when speaking, the phonation frequency of an adult man is ranging between 100 Hz and 150 Hz; Women are easily within the 200 Hz - 300 Hz band, while children have an area of variation between 300 Hz and 450 Hz. However, when it comes to singing, male bass performers are sonorous from 65 Hz (note C2) up to 330 Hz (note E4), while baritones range from 110 Hz (note A2) to 440 Hz (note A4), and tenors can climb up to 523 Hz (note C5). For women and children (Joliveau et al., 2004), contraltos, mezzo-sopranos and sopranos extend their pitches from 260 Hz (note C4) up to 1320 Hz (note E6).

Recent research has also shown that this part of the human body is not only responsible for the resonance of the vocal chords according to singing voice patterns (such as soprano, tenor or bass); The elasticity and vigorosity with which this complex neuro-muscular structure (Goodyer et al., 2006) corresponds to the pitch array of musicality prescribes the idiosyncratic speed with which the performer conveys a specified impression as a distinctive time quality (Keyzers et al., 2010).

Indeed, this valve mechanism seen in Fig. 7 is responsible for several physical concepts related with phonation.

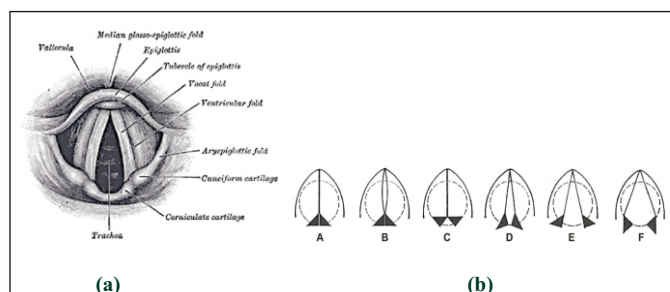


Fig. 7 (a) Transverse section of the vocal chords. **(b)** Vibrating sequences for the membranous tissue, in transverse view, depicting pressure flow and resistance.

The mechanism that was previously described acts as a fast acting valve that can interrupt and control airflow in many ways (Rohlf, et al. 2013). It is primarily responsible for the way in which voicing is produced as a basic oscillation.

The formation, however, of the impressive multitude of articulated sounds (Pontes et al., 2002), either when speaking or when singing, takes place at the

controllably resonating cavities that formulate the laryngeal oscillation, as seen in Fig. 8.

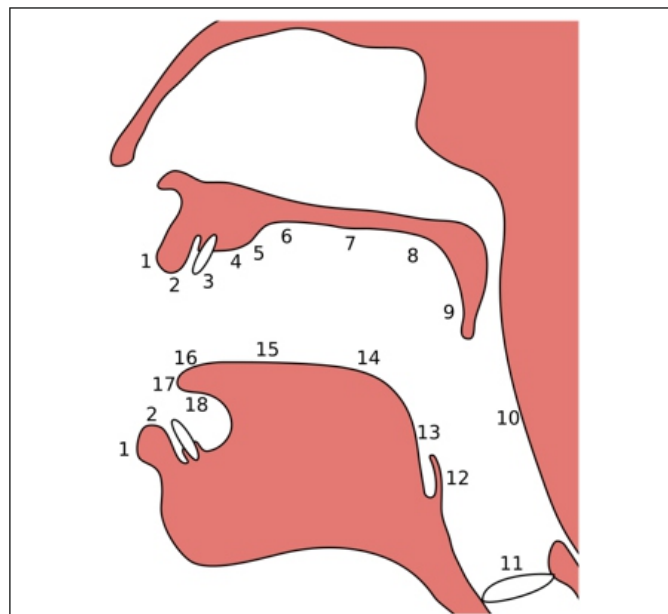


Fig. 8 Passive and active articulators: sagittal section, distributed via Wikimedia Commons. 1. Upper and Lower lip 2. Endo-labial part 3. Upper teeth 4. Alveolar ridge 5. Post Alveolar 6. Pre-palatal 7. Hard Palate 8. Velar 9. Uvula 10. Pharynx wall 12. Epiglottis 13. Tongue Root 14. Back of the tongue 15. Front of the tongue 16. Tongue blade 17. Tongue tip 18. Under the tongue. At the upper part, the Nasal cavity with the nostrils.

The active ones move to produce vocalizations, and the most influential of them are the lower lip, the tongue, the uvula, and the lower jaw. Accordingly, the most prominent passive articulators are the upper lip, the teeth, the upper jaw, the maxillary arch (with the upper jaw) and the pharynx. The soft palate at the roof of the mouth is a class of its own, being active and passive the same time, in the sense that it can lower it self, or the tongue can come in touch with it influencing the production of palatal and nasal sounds.

Also, the nasal cavity has substantial effectuation on the formation of voicing, and especially on singing (Pruthi and Espy-Wilson, 2004).

Acoustically, vowels are distinguished by each other by their energy peaks, or formants, that indicate how the phonation energy is distributed through the acoustic spectrum (Klein et al., 1970). They are produced at the end of the laryngeal part of the vocal tract, and they utilize the resonant properties of the tube system for phonation. The resonance patterns are rendered acoustically by tongue position and shape along with jaw and lip movement that control the airflow out of the oral cavity nozzles like a scoop.

Consonants on the other hand are basic speech sounds for which the airflow through the tract is at least partly obstructed. Apart from appreciable constriction, their other main characteristic is that they use as primary sound source the region above the larynx. The resultant acoustic signals bear greater complexity, but of course their energy content is significantly reduced in comparison to vowels. Furthermore, in most languages, they do not form autonomous meaningful morphological units, and they are combined in various forms with vowels.

IV. MODELS OF SPEECH PRODUCTION

When it comes to vowels or diphthongs, their phonetic properties are usually described with the apparent tongue position (Ladefoged, 2001). In the phonation of vowels and diphthongs, the larynx is acting like a uniform cross section tube, some 17.5 cm long. At the larynx end this tube is closed, while at the lips it is open. Both for vowels and diphthongs, the mouth is opened, and the cross section of the voicing mechanism at its "nozzle", i.e. the point of radiation of the acoustic signals, is shaped by the lip shape and position. The position of the tongue is also critical for the sound produced. By modulating tongue and lip position, different vowel sounds are produced. When speaking or singing, the tongue and the lips constantly move to differentiate the vowel produced. This is exactly the case between pure vowels and diphthongs: they perform exactly the same function in syllable structure, apart from the fact that diphthongs have a gliding quality, as seen in Fig. 9.

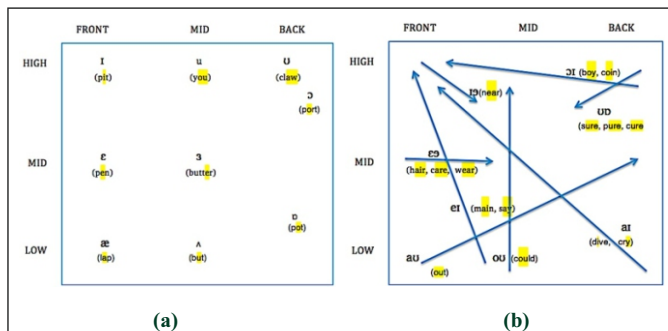


Fig 9. Auditory vowel plots for English with rather traditional accent. (a) Vowels (b) Diphthongs - the little noise that is associated with them is interpreted as a gliding action between the vowels involved.

In the case of various vowels, larynx phonation acts as the primary sound source. Significant amounts of acoustic energy are produced, and the distinctions of each specific vowel lies within the mouth, tongue and lip shape and position; all these organs act as filters that differentiate the sound spectrum creating harmonic frequency peaks and lows (Binns and Culling, 2007), reducing the energy radiated at the “nozzle” of the voicing mechanism, i.e. the mouth.

It should be noted that in various languages the speech sounds described in Fig. 9 may be different from the ones actually used (Ladefoged and Maddieson, 1996). Consequently, research studies aim to provide the actual pronunciation patterns, sometimes inventing new symbols, if the ones already in use do not describe accurately the phonetic quality involved (Keating, 1985).

So to speak, the pulmonic consonants that have been semiotically traced are seen in Fig. 10.

Manner ↓	Labial		Coronal				Dorsal			Laryngeal	
	Bilabial	Labio-dental	Linguo-labial	Dental	Alveolar	Post-alveolar	Retroflex	Palatal	Velar	Uvular	Pharyngeal/epiglottal
Nasal	m	m	ɱ	ɳ	ɲ	ɳ	ɳ	ɳ	ŋ	ɴ	
Stop	p	b	ɸ	t	d	ʈ	ɖ	c	k	q	ʁ
Sibilant fricative				s	z	ʃ	ʒ	ç	x	χ	ħ
Non-sibilant fricative	f	v	ɸ	θ	ð	ɬ	ɮ	ç	x	χ	ħ
Approximant				ɹ	ɻ	ɻ	ɻ	ɻ	ɻ	ɻ	ɻ
Tap/flap				ɾ	ɽ	ɽ	ɽ				
Trill				ʀ	ʀ						
Lateral fricative				ɬ	ɮ	ɬ	ɮ	ɬ	ɮ	ɬ	ɮ
Lateral approximant				l	ɭ	ɭ	ɭ	ɭ	ɭ	ɭ	ɭ
Lateral tap/flap				ɭ	ɭ	ɭ	ɭ	ɭ	ɭ	ɭ	ɭ

Source: IPA-Wikipedia.

Fig. 10. The extended pulmonic consonant table, updated in 2018, denoting how consonants are sonically related to the designated single place of articulation within the vocal tract. Non-pulmonic consonants and affricates (both pulmonic and ejective) are not included.

Bioenergetic Features and Human Activity

Medical diagnosis is involved with the Bioenergetic features of bodily functions when physicians apply therapeutic interventions aiming to manipulate the subtle energy related with bodily functions and overall with the body's energy equilibrium.

Bioenergetics, in the strict sense, have to do with physicochemical processes observed within living organisms. Subtle energy production constantly takes place due to biochemical reactions that harness energy from a variety of metabolic pathways (Politis et al., 2016). As the scope of medicine extends to characteristically small measurable factors, due to the use of hi-tech sensors, accepted principles from the scientific fields of biology, physics, chemistry and engineering are involved (Titze, 1994). Different theories have been developing gradually, aiming to support therapeutic interventions and, recently, the control of Prosthetics, Bionics and Implantations (Kyriafinis et al., 2017).

In the case of phonation (and its conjugate sense, hearing) the energy relationships into focus have to do with the potential to exert control over breathing, the resonance of the vocal tract Thomson et al., 2005), and the formation of sounds at the mouth (Zhang, 2010). Emission characteristics are also taken into account, having to do with openings of the lower part of human face, the shape of the lips and the combined use of teeth to partly obstruct the airflow (Wyke, 1983).

All these factors, have been described in detail within the previous section. The driving mechanism for voice production, however, is dependent on the neurocognition characteristics exerted over the muscles of the voicing mechanism. Not all people convey their emotions, as far as speech and music are con-

cerned, with the same intensity, magnitude or expressiveness (Weninger et al., 2013). In the case of singing, when professional artists are concerned, all these characteristics would determine the quality of performance as associated added value (Scherer, 1995).

Not all people are gifted with the ability to stress skillfully their muscles, bones and tongue to produce outstanding voicing patterns. Furthermore, even if two persons share identical bodily characteristics, it is not guaranteed that they will perform equally well. On the contrary, it seems that the neurocognitive apprehension of each one of them plays an equally important role, if not greater, since the neuronal capacity for speech is not interchangeable nor can extend to any direction (Kanai and Rees, 2011).

For instance, each person is very proficient phonetically in its mother tongue. In special cases, some people have been taught to speak some more languages from early childhood, developing apart from structured phraseology and vocabulary the ability to pronounce words in the particular way that natives do. If the synapses of the millions of neurons that are associated with phonation are not trained so in early childhood, it seems that the potential to pronounce phonemes, words or phrases deteriorates unrepairably. It is not the physical condition that is affected as the mental state that is falling behind in movement, progress and development of the parts of the phonatory mechanism (Magnuson and Nusbaum, 2007).

V. PROBLEM FORMULATION AND ANALYSIS OF MELODIC PATTERNS

In Israel are pertained in good condition historical buildings and shrines visited by nearly all-Christian congregations. Since such kind of journeying is a rather unique event, the bodies of believers from all around the globe seek to be engaged actively in the religious services associated with or containing memorabilia of particular revered persons and ecclesiastical events. Consequently, they chant and cooperate all-heartedly in ekphonic level ceremonies.

If the level of ritual complexity is beyond this prescribed form, more professional choirs and instrumentalists undertake the ceremonial formalities.

Usually visitors are quite relaxed, since they meet daily with thousands of people that come from another part of the globe, speaking a multitude of native languages, with whom they share more or less common beliefs. Knowing that they stray as a group amongst many others far away from their usual habitat rather unpersonified, they behave most of the time in a natural manner; this is the case indeed when they have opportunities for leisure.

When, however, they actively participate in rituals, they exhibit a very disciplined behavior in a series of actions performed according to a prescribed order, which is different for every country; it seems that the cultural, historical, linguistic (Chomsky, 1957) and climatic conditions prevailing in particular region have a significant capacity on the character or behavior demonstrated by a person and a group (Goble and Ni Chasaide, 2003).

In most Christian shrines of Israel, no microphones and electronic amplifiers are used to enhance human voice. As a result, the cantors, chanters and choir members, and the congregation itself when musically responding, seek to perform rather loudly, adjusting to the size of the building used for worship, so that the melodic rendition is easily audible to all members of the assembly.

The monitoring of the phonetic characteristics was modulated as follows:

- A. Phonetic Characteristics of Ethnic Groups:** The historical sites of Israel, Palestine and Middle East in general are visited by millions of tourists annually. Visitors usually move by bus or similar vehicles and most of the times they constitute a cohesive ethnically group sharing common beliefs, liturgical practices and operational travelling modes to stopover the places of interest. Nearly all times they have a guide who directs their moves and provides them with insightful information (Fig. 13).



Fig.13. Pictures of visitors in archaeological places of liturgical interest (a) The court of the Holy Sepulchre, Jerusalem (b) Deir Hajla plaza, near the Jordan river (c) Holy Cross Monastery, Rehavia.

When group members talk with each other, they demonstrate three clearly distinctive energy levels of speech communication (Heald and Nusbaum, 2014):

- Level 1.** The very low radiation energy mode - usually speakers use this mode when they don't want to be heard by others than their direct correspondents

(Ruggles et al., 2014). Often, they enhance the communication channel by observing their conversationalist's lips and thus understanding words that are not otherwise very clearly audible.

Level 2. The average, normal radiation energy mode - it is a way of communication that expresses thoughts and feelings (Hellegouarch, 2002) by clearly distinct articulate sounds (Schmidt and Kim, 2010); it is the mode that most people use to communicate within group discussions and normal social activities (Redondo et al., 2008).

Level 3. The loud speech communication mode, involving significant levels of energy. It is a deliberate way of speaking, arousing high energy levels of sonic energy, which takes place when subjects are characterized by intense feelings (Södersten et al., 2005). Apart from galvanized atmospheres due to excitement, relaxation or purposive activities, agitation for speaking loudly may be chaotic dins caused by crowds of people, shifting the more enthusiastic population groups to laughter and shouting (Mladenovic et al., 2016). Listeners positioned many meters away from the sound source can very clearly hear the conversation.

At what time it comes to organized singing or chanting, people usually shift to energetic Level 2 or Level 3. Their choral performance may have different levels of complexity, in prosodic and musical terms (Zhang, 2015). It may be ekphonic, which means that the melody is rather simple and speech sounds have a dominant position (Alku et al., 2002). When it is more complex melodically, the modal elements are rhymed phrases with clearly rhythmic musical phrases involved.

These characteristics of phonation were examined in three different assessments in Israel, during September-October 2018.

B. Ekphonic Characteristics of Ethnic Groups

The phonetic differentiation between languages (Cummins et al., 1999) was examined in the main church of the Holy Sepulchre, under the same environmental factors. In this monument, no microphones or other means of electro-mechanical amplification of voice may be used. During an official celebration in September 2018, four well-trained deacons recited a widely known portion of the Gospels (John 19:6-11, 13-20, 25-35 NRSV) in their mother tongue. They rendered it respectively in Greek, Arabic, Russian and Romanian. The style of the ekphonic performance was the same - and this was recounted by the fact that their renditions for a rather extensive text were approximately the same in time length.

All deacons, contrasting their ethnic characteristics as far as the level of loudness is concerned, produced similar energetic levels of voicing, at Level 3, aiming to deliver with adequate sound pressure levels a respectful rendition of the text, heard by all members of a multi-ethnic congested congregation.

Fairly long segments (~10s) of these recitals were analyzed using an 10th-degree linear prediction model to denote unambiguous determination of formant characteristics. Indeed, the variability of the second formant is indicative of the different way that resonances and antiresonances are handled in each language in perspective, i.e. how fricatives, laryngeal or even nasal consonants are conducted in voicing (Fig. 14).

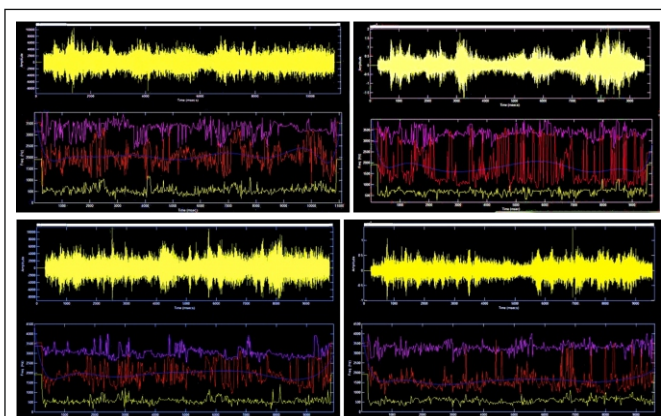


Fig.14. About 10s segments from four ekphonic recitals of the same Gospel portion (John 19:6-11, 13-20, 25-35 NRSV) by four different deacons in Greek (upper left), Arabic (upper right), Russian (lower left) and Romanian (lower right). The first three formants are depicted, with the second being approximated with linear regression of an 11th degree polynomial (blue line).

In Greek, vowels and diphthongs predominate; although they are elaborated with surrounding consonants, fairly flat segments of smooth phonation curves for the second formant are formed with adequate sustain. In Arabic, although phonation involves similar attitudes, with a comparatively open configuration of the vocal tract and mouth, the continuous use of laryngeal consonants, like

a soft /h/, makes syllables to have a more harsh hearing. It seems that although this language was developed in warm climates, the arid environment forces speakers to avoid the results of dehydration (Miri et al., 2012) by adopting attributes of physical action and movements which involve voice modulation through expansion or contraction of the glottal and epiglottal area, where sufficient levels of moisture are retained (Tao et al., 2010). Even further, a highly energetic character appears in phonation, shifting with rapidity from one speech pattern to another.

In Russian consonants prevail. Additionally, it is common to have continuous streams of two or even three consonants, interrupted by very short vowel intervals. Since the ekphonic rendition has a melodic orientation, the endings of phrases have considerable prolongation, not detected in normal, everyday speech communication. Audible friction prevails with the use of stops and sibilant fricatives aiming to prevent energy loss by inhaling cold air. Speech intelligibility is achieved by short, sharp, high pitched sounds related with consonants (Nan et al., 2008). There is a considerable variability in the number of consonants involved, in the sense that even further they may alter their phonation characteristics to "hard" or "soft" patterns, depending on the short vowels or consonants that follow. While in Arabic this variability has as epicenter the back of the oral cavity, in Russian the key points are the upper and lower lip, the endo-labial part, the teeth and the alveolar ridge. Romanian, on the other part, have a considerable resemblance with Greek, as they demonstrate stability in sustaining vowel pronunciation. However, phonemes have not many trills, monophthongs prevail characteristically over diphthongs and they have a shorter sustain time.

C. Melodic Characteristics of Ethnic Groups

In the main worship building of the 4th century AD Monastery of Holy Cross at the Rehavia district of Jerusalem, renovated in its current form in 1644 AD, three different groups of about 20 persons each, interpreted in their mother tongue a rendition of a well-known prayer (Matthew 6:9-13 NRSV). The first group was Arab, consisting of an equal number of male and female choir members (Klatt and Klatt, 1990). Their rendition was characteristically rhythmic, rather ekphonic in its prosodic nature. The second was Russian - the choir leader was male, and lead the rendition with a powerful low frequency voice; however, by far most of the members were female (Simpson, 2009). They sang the prayer polyphonically, characteristically giving emphasis on the mezzo-soprano voice (folding around 500Hz). The choral was more melodic and slower in its harmonized final syllables, especially when compared to the Arab version (Balkwill and Thompson, 1999). The third rendition was in Romanian, performed ekphonically by an almost all-female choir.

The choirs were very well trained in reciting the melodic piece (Howard, 2009). In a very respectful manner (Lu et al., 2006; Patel et al., 2011), again without any voice amplifying instrumentation, they performed loud, at Level 3, adjusting the power of their phonation to less than 50% of their voicing maximum, so to keep the melodic nature of the prayer and not to revert to shouting (Scherera, 2017).

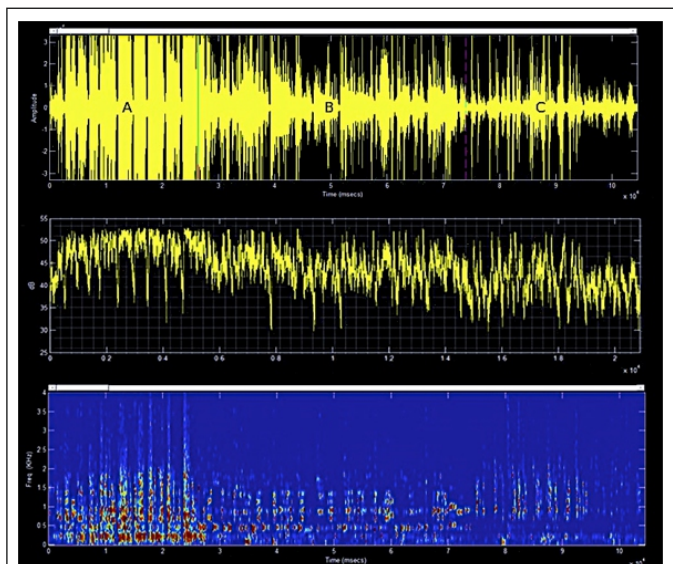


Fig.15. A well known prayer melodically performed successively by Arab (section A), Russian (section B) and Romanian choirs (section C). On the upper part, the time series of the recording is depicted. On the middle part the energy level distributions are shown. On the lower part, the frequency spectrum is calculated.

The energy-frequency spectrum correlation (Master et al., 2006) of their singing is depicted in Fig. 15. In section A it is clear that the Arabic performance has somewhat more energy than the other two (Varnet et al., 2017). Evidently, the major characteristic of their performance is that voicing breaks into bursts

(He and Dellwo, 2016), giving the distinctive cheerful, speedy motion of the singing liveliness that characterizes this cultural group (Scherer, 2003). Again, laryngeal consonants, like variants of /h/ are abundantly heard.

The performance in Russian, on the contrary, is by far greater in length (Manfredi et al., 2015); it is polyphonic and three distinct voices are used - normal male (between tenor and bass), mezzosoprano and soprano (Caclin et al., 2008). The emotions intended to be provoked are put in another direction: solemn procession celebrated with a formal and dignified tuneless rendition (Laukka et al., 2005).

For this purpose, vowels and diphthongs are stressed and prolonged. The chirp-like operation of consonants for speech intelligibility is reduced, and the meaning of words is correlated with melodic contours (Ferrer, 2009) that have a clear influence from "Western" music of the 19th century and the first half of the 20th. However, the mentality and the composing characteristics have a unique "Russian" character - which a careful listener would easily recognize (Kharlamov, 2015).

The performance in Romanian has features from both (Iseli et al., 2007): it is ekphonic but it has melodic elements that suggest an influence from the "West" and the "North", apart from the relevance with the Eastern Mediterranean substratum (Gjerdingen and Perrott, 2008). Vowels performed by the female choirists are clearly depicted in the upper frequencies of the spectrum.

C. Sustainability in Voicing and Cultural Characteristics of Ethnic Groups

This phase of the survey focused on how speech communication biases convolve with dominant macroscopic linguistic patterns relating to national or cultural origins. The members of visiting groups, as they entered the plaza of the existing since the 5th century AD St. Gerasimus Monastery, in the Deir Hajla oasis of the Jordan river desert, where usually casually speaking to each other, asking other members of the group and the attendants about souvenirs or local beverages and often they sat by the chairs of the café for refreshment, speaking freely and quite relaxed. The ethnic groups monitored in this survey in October 2018 where Greek, Russian, Romanian, Serb, Filipino, Middle East Arab, Israeli Hebrew, Israeli Arab, and Eritrean residents of Israel (using in their communication both Cush and Hebrew). They were enthusiastic in the sense that they visited a site linked with each group's own historical and religious tradition, and they felt quite familiar with the surroundings and the conditions of the monument. Although hundreds of tourists were flocking around, they had a relaxed feeling of anonymity, since the limited perception by other groups of their mother tongue communiqué, none of which is a lingua franca, guarantees sufficient immunization against identification for their encounters (van Dommelen and Hazan, 2012). All recordings compared hereinafter were taken from the same point; as reference sound for normalization of the results was used the welcome bell-ringing melody, which was electrically triggered and thus had invariable characteristics. The temperature of the environment was around 28° - 33° C, depending on the hour the actual recording took place (Fig. 16).



Fig.16. Pictures of visitors, from diverse communities, socializing in the St. Gerasimus Monastery Plaza, Deir Hajla, Israel.

The basic advantage for conducting this research in an oasis in the desert is that, under normal circumstances, the boisterous, noisy environment that characterizes urban habitats is totally absent (Yiu and Yip, 2016). When vehicles and other electromechanical equipment are not interfering with the socializing crowd, the environment is noise free (Parbery-Clark et al., 2009).

The following characteristics were observed: Russians nearly never exceed Level 2; as a matter of fact, under most situations they exhibit a confidentiality attitude, reflecting a physical and mental activity giving prominence to a very respectful attitude towards other people, having as epicenter Level 1 communication. On the contrary, Greeks rarely speak amongst themselves with energetic activity below Level 2, and many times conversations shift enthusiastically to Level 3. Filipino or Middle East peoples demonstrate a similar attitude, however somewhat more restrained, with Romanian groups sometimes drifting to Level 1 communication. However, it is not easily configurable, apart from the people coming from cold northern countries, to discriminate if this group attitude is a settled tendency biased from socio-historical conditions or prevailing weather conditions; even further, it may be a pretentious, uncontrollably exuberant behavior for certain groups to publicly parade themselves.

However, as the groups monitored were not placed equally distanced around the recording point, the focal processed quantity is not the energetic level of

spontaneous voicing, which obviously is high (in observational terms Level 2 and beyond), but the sustainability to high levels of vocalization social activity (Peperkamp et al., 2010). For this reason, the processed energy levels of group vocal outputs were approximated using a regression model of a 10th degree polynomial curve, clearly seen in Fig. 17 (Markel and Gray, 1982).

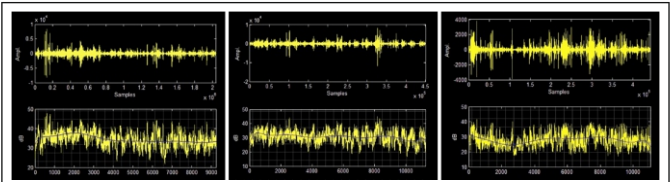


Fig.17. Sustainability curves (estimated with regression methods) for voicing levels in socializing conditions in Deir Hajla Plaza, for groups of 8-10 people. Left, a Filipino group, center, a Greek group and right, an Israeli group, all communication in their mother tongue.

As it may be seen in Fig. 17, for groups living in similar, more or less climatic habitats (Filipino, Greek, Israeli), the energy levels are clearly high, produced by speech with distinctly prevailing resonating features. Indeed, acoustically the combination of short and long vowels prevails. The people from Philippines speak clearly, loudly, in Hispanic most of the time, and demonstrate a cheerful communication with limited constriction or obstruction. For the Greek group, the fact that their language is epicentral to vowels with very clear sustain pronunciation patterns, along with other cultural reasons, makes the participants in conversation to maintain high levels of energy for considerable time intervals.

The same attitude is observed for the Israeli speakers. They are very lively and cheerful in their communiqué. However, there is an obvious drift from the highest energetic level, and their overall group pattern is interspersed with notable ups and downs.

VI. CONCLUSION:

Climatic conditions of the natural habitat in which ethnic groups have developed their linguistic potential are essential factors that mould the voicing characteristics of the words used within the structures of speech communication. Along with strong cultural bias exerted as an endowment advocating the social heredity, they promote characteristic expression indices that reveal the influence of powerful bioenergetic constituents.

This evaluation was based on the fact that the hundreds of visitors monitored acted naturally, neither having any visible clue of being monitored nor feeling that their anonymity had been violated, as they gathered with thousands of peers from all the corners of our planet in historical congregations of Israel.

Since the recording conditions were not studio-standard, the survey emphasized not on measuring the absolute levels of voicing parameters; rather it examined potentials of speech communication and the singing voice that are qualitative, deep-rooted distinctive features characterizing the diachrony of ethnic groups and not the phenotype or ephemerality of their expression.

ACKNOWLEDGEMENT:

The authors would like to express their gratitude to his Beatitude, the Greek Orthodox Patriarch of Jerusalem Theophilos III, the Hagiotaphite Brotherhood and Archimandrite Chrysostomos, Abbot of the St. Gerasimus Monastery in Deir Hajla, the desert of Judea, for providing the resources to conclude this survey within the mosaic of people frequenting the historical sites of their jurisdiction, either as visitors or as workers, associates and guests.

REFERENCES

1. Alku, P. and Vilkman, E. (1996), A Comparison of Glottal Voice Source Quantification Parameters in Breathily, Normal and Pressed Phonation of Female and Male Speakers, *Folia Phoniatr Logop.*, vol. 48, pp. 240-254.
2. Alku, P., Vinturi, J. and Vilkman, E. (2002), Measuring the Effect of Fundamental Frequency Raising as a Strategy for Increasing Vocal Intensity in Soft, Normal and Loud Phonation, *Speech Comm.*, vol. 38, pp. 321-334.
3. Anderson, S. (2012), *Languages - a Very Short Introduction*, Oxford University Press.
4. Balkwill, L. L. and Thompson, W. F. (1999), A cross-cultural investigation of the perception of emotion in music: Psychophysical and cultural cues, *Music Perception*, vol. 17, pp. 43-64.
5. Balkwill, L.L., Thompson, W. F. and Matsunaga, R. (2004), Recognition of Emotion in Japanese, Western, and Hindustani Music by Japanese Listeners, *Japanese Psychological Research*, vol. 46, pp. 337-349.
6. Bar-Hillel, Y. (1964), *Language and Information*, Addison-Wesley and Jerusalem Academic Press.
7. Barker, A. (2009), *Scientific Method in Ptolemy's Harmonics*, Cambridge University Press.
8. Binns, C. and Culling, J.F. (2007). The Role of Fundamental Frequency Contours in the Perception of Speech Against Interfering Speech, *J. Acoust. Soc. Am.*, vol. 122, pp. 1765-1776.

9. Boas, F. (1940), *Race, Language and Culture*, New York:Mcmillan.
10. Botinis, A. (Ed.) (2000), *Intonation, Text, Speech and Language Technology*, No. 15, Springer, Amsterdam, the Netherlands
11. Caclin, A., McAdams, S., Smith, B. K. and Giard, M. H. (2008), Interactive Processing of Timbre Dimensions: An Exploration with Event-related Potentials, *Journal of Cognitive Neuroscience*, vol. 20, no. 1, pp. 49-64.
12. Carey M.A., Card J.W., Voltz J.W. et al. (2007), It's all About Sex: Gender, Lung Development and Lung Disease, *Trends Endocrinol Metab.*, vol. 28, pp.308-313.
13. Carney, S. (2012), *Imagining Globalisation: Educational Policyscapes*, in Steiner-Khamsi, G. and Waldow, F. (Eds.) *World Yearbook of Education 2012: Policy Borrowing and Lending in Education*, London: Routledge.
14. Carol, J.B. (1964), *Language and Thought (Foundations of Modern Psychology)*, Prentice-Hall.
15. Chan R.W. (2001), Estimation of Viscoelastic Shear Properties of Vocal-fold Tissues Based on Time-Temperature Superposition, *J. Acoust. Soc. Am.*, vol. 110, pp. 1548-1561.
16. Chhetri, D.K., Berke, G.S., Lotfzadeh, A. and Goodyer, E. (2009), Control of Vocal Fold Cover Stiffness by Laryngeal Muscles: A Preliminary Study, *Laryngoscope*, vol. 119, no. 1, pp. 222-227.
17. Chhetri, K., Neubauer, J. and Berry, D.A. (2012), Neuromuscular Control of Fundamental Frequency and Glottal Posture at Phonation Onset, *J. Acoust. Soc. Am.*, vol. 131, no. 2, pp. 1401-1412.
18. Chomsky, N. (1957), *Current Issues in Linguistic Theory*, The Hague: Mouton.
19. Chomsky, N. (2000), *New Horizons in the Study of Language and Mind*, Cambridge University Press.
20. Ciggaar, K. (1996), *Western Travellers to Constantinople - The West and Byzantium 962-1204: Cultural and Political Relations (The Medieval Mediterranean)*, Leyde-N. York-Cologne: Brill.
21. Collins, M., and Price, M.A. (2003), *The Story of Christianity*, Dk Pub.
22. Comrie, B. and Corbett, G. (1993) (Eds.), *The Slavonic Languages*, London: Routledge.
23. Dellwo, V., Leemann, A. and Kolly, M.J. (2015), Rhythmic Variability Between Speakers: Articulatory, Prosodic, and Linguistic Factors, *J. Acoust. Soc. Am.*, vol. 137, pp. 1513-1528.
24. Dellwo, V. and Wagner, P. (2003), Relations Between Language Rhythm and Speech Rate, in *Proceedings of the International Congress of Phonetics Science*, Barcelona, Spain, pp. 471-474.
25. Department for Education - UK (2010), *The Importance of Teaching: The Schools White Paper*, Norwich: The Stationary Office.
26. Devine, A. M. and Stephens, L. D. (1994), *The Prosody of Greek Speech*, New York: Academic Press.
27. Elias, N. (2000), *The Civilizing Process - Sociogenetic and Psychogenetic Investigations*, Oxford, UK: Blackwell Publishers.
28. Erath, B.D. and Plesniak, M.W. (2010), An Investigation of Asymmetric Flow Features in a Scaled-up Driven Model of the Human Vocal Folds, *Exp. Fluids*, vol. 49, no. 1, pp. 131-146.
29. Evans, H.C. and Wixom, W.D. (Eds.) (1997), *The Glory of Byzantium - Art and Culture of the Middle Byzantine Era AD 843-1261*, N. York: Metropolitan Museum of Art.
30. Ferrer, R. (2009), Embodied Cognition Applied to Timbre and Musical Appreciation: Theoretical Foundation, *British Postgraduate Musicology*, vol. 10, pp. 1-8.
31. Finnegan, E., Luschei, E. and Hoffman, H. (2000), Modulations in Respiratory and Laryngeal Activity Associated with Changes in Vocal Intensity During Speech, *J. Speech Lang. Hear. Res.*, vol. 43, pp. 934-950.
32. Forbes, P.B.R. (1933), Greek Pioneers in Philosophy and Grammar, *The Classical Review*, vol. 47, no. 3, pp. 105-112.
33. Garnier, M., Henrich, N., Smith, J. and Wolfe, J. (2010), Vocal Tract Adjustments in the High Soprano Range, *J. Acoust Soc Am.*, vol. 127, pp.3771-3780.
34. Gjerdingen, R. and Perrott, D. (2008), Scanning the Dial: The Rapid Recognition of Music Genres, *Journal of New Music Research*, vol. 37, no. 2, pp. 93-100.
35. Gobl, C. and Ni Chasaide, A. (2003), The Role of Voice Quality in Communicating Emotion, Mood and Attitude, *Speech Commun.*, vol. 40, pp. 189-212.
36. Goodyer, E., Muller, F., Bramer, B., Chauhan, D. and Hess, M. (2006), In Vivo Measurement of the Elastic Properties of the Human Vocal Fold, *Eur Arch Otorhinolaryngol*, vol. 263, pp. 455-462.
37. Gregory, A. H. and Varney, N. (1996), Cross-Cultural Comparisons in the Affective Response to Music, *Psychology of Music*, vol. 24, pp. 47-52.
38. Hammann, S. and Harwood, D. L. (1976), *Universals in Music: A Perspective From Cognitive Psychology, Ethnomusicology*, vol. 22, pp. 521-533.
39. Hayas, C. (1953), *Christianity and Western Civilization*, Stanford University Press.
40. Hayes, B. (1984), The Phonetics and Phonology of Russian Voicing Assimilation. In Aronoff, M. and Oehrle, R.T. (Eds.), *Language Sound Structure*, Cambridge, Mass.: MIT Press, pp. 318-328.
41. Hays, J. (2008©), *Ancient Music in China*. [http:// factsanddetails.com](http://factsanddetails.com) Last updated: May 2016
42. He, L. and Dellwo, V. (2016), The Role of Syllable Intensity in Between-Speaker Rhythmic Variability, *Int. J. Speech Language Law*, vol. 23, pp. 243-273.
43. Heald, S. L. and Nusbaum, H. C. (2014), Speech Perception as an Active Cognitive Process, *Front. Syst. Neurosci.*, vol. 8, pp. 1-15.
44. Herrin, J. (1989), *The formation of Christendom*, Princeton University Press.
45. Hellegouarch, Y. (2002), *A Mathematical Interpretation of Expressive Intonation*, in Bruter, C. (Ed.) *Mathematics and Art: Mathematical Visualization in Art and Education*, New York: Springer-Verlag.
46. Hickmann, H. (1960), Rythme, mètre et mesure de la musique instrumentale et vocale des anciens Egyptiens, *Acta Musicologica*, vol. 32, Fasc. 1, pp. 11-22.
47. Holmberg, E. B., Hillman, R. E. and Perkell, J. S. (1989), Glottal Air Flow and Transglottal Air Pressure Measurements for Male and Female Speakers in Low, Normal, and High Pitch, *J. Voice*, vol. 3, pp. 294-305.
48. Howard, D. (2009), Acoustics of the Trained versus Untrained Singing Voice, *Curr Opin Otolaryngol Head Neck Surg.*, vol. 17, pp. 155-159.
49. Horacek, J., Laukkanen, A.M., Sidlof, P., Murphy, P. and Svec, J.G. (2009), Comparison of Acceleration and Impact Stress as Possible Loading Factors in Phonation: A Computer Modeling Study, *Folia Phoniatr Logop.*, vol. 61, pp. 137-145.
50. Hsiao, T.Y., Wang, C.L., Chen, C.N., Hsieh, F.J. and Shau, Y.W. (2002), Elasticity of Human Vocal Folds Measured in Vivo Using Color Doppler Imaging, *Ultrasound Med Biol.*, vol. 28, pp. 1145-1152.
51. Huron, D. (2001), Is Music an Evolutionary Adaptation?, *Annals of the New York Academy of Sciences*, vol. 930, pp. 43-61.
52. Iseli, M., Shue, Y. and Alwan, A. (2007), Age, Sex, and Vowel Dependencies of Acoustic Measures Related to the Voice Source, *J. Acoust. Soc. Am.*, vol. 121, pp. 2283-2295.
53. Joliveau, E., Smith, J. and Wolfe, J. (2004), Vocal Tract Resonances in Singing: The Soprano Voice, *J. Acoust Soc Am.*, vol. 116, pp. 2434-2439.
54. Juslin, P. N. and Laukka, P. (2003), Communication of Emotions in Vocal Expression and Music Performance: Different Channels, Same Code?, *Psychological Bulletin*, vol. 129, no. 5, pp. 770-814.
55. Kanai, R. and Rees, G. (2011), The Structural Basis of Inter-Individual Differences in Human Behavior and Cognition, *Nat. Rev. Neurosci.* vol. 12, pp. 231-242.
56. Keating, P. (1985), Universal Phonetics and the Organization of Grammars, in Fromkin, V.A. (Ed.) *Phonetic Linguistics: Essays in Honor of Peter Ladefoged*, Orlando, FL: Academic Press, pp. 115-132.
57. Keyser, C., Kaas, J. H. and Gazzola, V. (2010), Somatosensation in Social Perception. *Nature Review Neuroscience*, vol. 11, pp. 417-428.
58. Kharlamov, V. (2015), Perception of Incompletely Neutralized Voicing Cues in Word-Final Obstruents: The Role of Differences in Production Context, *Lab. Phon.*, vol. 6, no. 2, pp. 147-165.
59. Klatt, D. H. and Klatt, L. C. (1990), Analysis, Synthesis and Perception of Voice Quality Variations Among Male and Female Talkers, *J. Acoust. Soc. Am.*, vol. 87, pp. 820-856.
60. Klein, W., Plomp, R. and Pols, L., (1970), Vowel Spectra, Vowel Spaces and Vowel Identification, *J. Acoust. Soc. Am.*, vol. 48, pp. 999-1009.
61. Kortlandt, F. (1994), From Proto-Indo-European to Slavic, *Journal of Indo-European Studies*, vol. 22, pp. 91-112.
62. Krebs, F., Silva, F., Sciamarella, D. and Artana, G. (2012), A Three-Dimensional Study of the Glottal Jet, *Experiments in Fluids*, Springer Verlag (Germany), vol. 52, no. 5, pp. 1133-1147.
63. Kyriafinis, G., Stagiopoulos, P., Aidona, S., Politis, D., Aleksić, V. and Constantinidis, I. (2017), Linguistic Rehabilitation and Singing Potential: Correlating Performance Shapes with Sonic Contours, *International Journal of Current Research*, vol. 9, no. 7, pp. 54470-54480.
64. Ladefoged, P. (2001), *Vowels and Consonants*, Oxford: Blackwell.
65. Ladefoged, P. and Maddieson, I. (1996), *The Sounds of the World's Languages*, Oxford, UK: Blackwell.
66. Laukka, P., Juslin, P. N. and Bresin, R. (2005), A Dimensional Approach to Vocal Expression of Emotion, Cognition and Emotion, vol. 19, no. 5, pp. 633-653.
67. Lentz, D. (1961), *Tones and Intervals of Hindu Classical Music*, Lincoln: University of Nebraska.
68. Levitin, D. J. and Menon, V. (2003), Musical Structure is Processed in "Language" Areas of the Brain: A Possible Role for Brodmann Area 47 in Temporal Coherence, *Neuroimage*, vol. 20, pp. 2142-2152.
69. Lightner, T. (1972), *Problems in the Theory of Phonology, I: Russian phonology and Turkish phonology*, Edmonton: Linguistic Research, Inc.
70. Loh, K.K. and Kanai, R. (2016), How has the Internet Reshaped Human Cognition?, *Neuroscientist*, vol. 22, pp. 506-520.
71. Lu, L., Liu, D. and Zhang, H. J. (2006), Automatic Mood Detection and Tracking of Music Audio Signals, *IEEE Transactions on Audio, Speech, and Language Processing*, vol. 14, no. 1, pp. 5-18.
72. MacLarnon, A. M. and Hewitt, G.P. (1999), The Evolution of Human Speech: the Role of Enhanced Breathing Control, *Am J Phys Anthropol.*, vol. 109, pp. 341-363.
73. Magnuson, J. S. and Nusbaum, H. C. (2007), Acoustic Differences, Listener Expectations, and the Perceptual Accommodation of Talker Variability, *J. Exp. Psychol.*, vol. 33, pp. 391-409.
74. Manfredi, C., Barbagallo, D., Baracca, G., Orlandi, S., Bandini, A. and Dejonckere, P. H. (2015), Acoustics Professional Singing Voice: Operatic vs. Jazz, *J. Voice*, vol. 29, no. 4.
75. Markel, J.D. and Gray, A.H. (1982), *Linear Prediction of Speech*, Berlin: Springer-Verlag.
76. Master, S., Biase, N.D., Chiari, B.M. and Pedrosa, V. (2006), The Long-term Average Spectrum in Research and in the Clinical Practice of Speech Therapists, *Pro Fono.*, vol.

- 18, pp. 111–120.
77. Mersenne, M. (1635), *Harmonie Universelle, Contenant la Theorie et la Pratique de la Musique, ou il est Traite de Consonances, des Dissonances, des Genres, des Modes, de la Composition, de la Voix, des Chants, et de Toutes Sortes d'Instruments Harmoniques*, Paris: Pierre Ballard. [Trans. Roger E. Chapman as *Harmonie Universelle: The Books on Instruments*. The Hague: Nijhoff, 1957].
 78. Miri, A.K. (2014), Mechanical Characterization of Vocal Fold Tissue: A Review Study, *J Voice*, vol. 28, no. 6, pp. 657–667.
 79. Miri, A.K., Barthelat, F. and Mongeau, L. (2012), Effects of Dehydration on the Viscoelastic Properties of Vocal Folds in Large Deformations, *J of Voice*, vol. 26, pp. 688–697.
 80. Mladenović, M., Mitrović, J., Krstev, C. and Vitas, D. (2016), Hybrid Sentiment Analysis Framework for a Morphologically Rich Language, *Journal of Intelligent Information Systems*, vol. 46, no. 3, pp. 599–620.
 81. Morris, P., and Adamson, B (2010), *Curriculum, Schooling & Society in Hong Kong*, The Hong Kong University Press.
 82. Nan, Y., Knosche, T.R., Zysset, S. and Friederici, A.D. (2008), Cross-cultural Music Phrase Processing: An fMRI Study, *Human Brain Mapping*, vol. 29, no. 3, pp. 312–328.
 83. Nicol, D. M. (2002), *The Last Centuries of Byzantium*, 2nd ed., Cambridge University Press.
 84. Obolensky, D. (2009), *The Byzantine Commonwealth: Eastern Europe 500-1453*, ACLS Humanities E-book.
 85. Parbery-Clark, A., Skoe, E. and Kraus, N. (2009), Musical Experience Limits the Degradative Effects of Background Noise on the Neural Processing of Sound, *J. Neurosci.*, vol. 29, pp. 14100–14107.
 86. Patel, S., Scherer, K., Bjorkner, E. and Sundberg, J. (2011), Mapping Emotions Into Acoustic Space: The Role of Voice Production, *Biol. Psych.*, vol. 87, pp. 93–98.
 87. Peperkamp, S., Vendelin, I. and Dupoux, E. (2010), Perception of Predictable Stress: A Cross-Linguistic Investigation, *J. Phonetics*, vol. 38, no. 3, pp. 422–430.
 88. Pingle, B. A. (1962), *History of Indian Music*, Calcutta: Gupta.
 89. Pöhlmann, E. and West, M. L. (2001), *Documents of Ancient Greek Music*, Oxford, UK: Academic Press.
 90. Politis, D., Tsalighopoulos, M. and Iglezakis, I. (Eds.) (2016), *Digital Tools for Computer Music Production & Distribution*, N. York: IGI Global.
 91. Politis, D., Tsalighopoulos, M. and Kyriafinis, G. (2016), From Cochlear Implants and Neurotology to Brain Computer Interfaces: Exploring the world of neuron synapses for hearing impairments, in Mandal, J.K., Mukhopadhyay, and Pal, T. (Eds.), *Handbook of Research on Natural Computing for Optimization Problems*, Hershey, PA: IGI Global.
 92. Politis, D., Kyriafinis, G., Constantinidis, J. and Paris, N. (2017), Neurophysiology-based acoustic measurements of singing contours - an empathetic interactive loudness approach, in *Proceedings of the International Conference on Interactive Mobile Communication Technologies and Learning, IMCL2017, Thessaloniki*, 803-812.
 93. Pope, S. (1986), *Music Notations and the Representation of Musical Structure and Knowledge, Perspectives of New Music*, vol. 24, pp. 156–189.
 94. Pontes, P.A.L., Vieira, V.P., Gonçalves, M.I.R. and Pontes, A.A.L. (2002), Characteristics of Hoarse, Rough and Normal Voices: Acoustic Spectrographic Comparative Analysis, *Rev Bras Otorrinolaringol.*, vol. 68, pp. 182–188.
 95. Perry, M. (2012), *Western Civilization: A Brief History, Volume I: To 1789*, 10th ed., CA: Wadsworth Publishing.
 96. Pruthi, T. and Espy-Wilson, C. (2004), Acoustic Parameters for Automatic Detection of Nasal Manner, *Speech Commun.*, vol. 43, pp. 225–239.
 97. Ramat, A.G. (Ed.) (1988), *The Indo-European Languages*, London and New York: Routledge.
 98. Rappleye, J. (2012), *Educational Policy Transfer in an Era of Globalization: Theory - History - Comparison*, Frankfurt: Peter Lang.
 99. Rothstein, E. (1995), *Emblems of Mind: The Inner Life of Music and Mathematics*, New York: Times Books.
 100. Redondo, J., Fraga, I., Padron, I. and Pineiro, A. (2008), Affective Ratings of Sound Stimuli, *Behavior Research Methods*, vol. 40, no. 3, pp. 784–790.
 101. Rohlf, et al. (2013), The Anisotropic Nature of the Human Vocal Fold: An Ex Vivo Study, *Eur. Arch. Otorhinolaryngol.*, vol. 270, no. 6, pp. 1885–1895.
 102. Ruggles, D. R., Freyman, R. L. and Oxenham, A. J. (2014), Influence of Musical Training on Understanding Voiced and Whispered Speech in Noise, *PLoS One*, vol. 9, no. 1.
 103. Runciman, S. (1994), *Byzantine Civilization*, N. York: Barnes and Noble.
 104. Scherera, K.R., Sundberg, J., Fantini, B., Trznadel, S. (2017), The expression of emotion in the singing voice: Acoustic patterns in vocal performance, *J. Acoust. Soc. Am.*, vol. 142, no. 4, 1805.
 105. Selbie, W.S., Gewalt, S. L. and Ludlow, C. L. (2002), Developing an Anatomical Model of the Human Laryngeal Cartilages from Magnetic Resonance Imaging, *J. Acoust. Soc. Am.*, vol. 112, no. 3, pp. 1077–1090.
 106. Scherer, K. R. (2003), Vocal communication of emotion: A review of research paradigms, *Speech Commun.*, vol. 40, pp. 227–256.
 107. Sachs, C. (1921), *Die Musikinstrumente des alten Ägyptens. Mitteilungen aus der Ägyptischen Abteilung der Staatmuseen* 3, Berlin: K. Curtius.
 108. Schank, R. C. and R. Abelson (1976), *Scripts, Plans, Goals, and Understanding*, Hillsdale, N.J.: Erlbaum.
 109. Scherer, K. R. (1995), Expression of emotion in voice and music, *J. Voice*, vol. 9, no. 3, pp. 235–248.
 110. Schmidt, E., M. and Kim, Y. E. (2010), Prediction of Time-varying Musical Mood Distributions from Audio, In *Proceedings of the International Society for Music Information Conference*, Utrecht, Netherlands.
 111. Shosted, R., Carignan, C. and Rong, P. (2012), Managing the distinctiveness of plosive nasal vowels: Articulatory evidence from Hindi, *J. Acoust. Soc. Am.*, vol. 131, pp. 455–465.
 112. Simpson, A. P. (2009), Phonetic Differences Between Male and Female Speech, *Lang. Linguist. Compass*, vol. 3, no. 2, pp. 621–640.
 113. Södersten M, Ternström S, and Bohman M. (2005), Loud speech in realistic environmental noise: phone to gram data, perceptual voice quality, subjective ratings, and gender differences in healthy speakers, *J Voice*, vol. 19, pp. 29–46.
 114. Spyridis, Ch. (1999), Ένα αρχαίο μουσικό επίγραμμα από το νησί της Αερίας (Θάσου), in *Proceedings of the International Meeting on Music and Ancient Greece*, 5–15 August 1996, Athens: Symposium Proceedings, pp. 169–94.
 115. Steriade, D. (1997), *Phonetics in Phonology: The Case of Laryngeal Neutralization*, Manuscript - <http://linguistics.ucla.edu>
 116. Statista 2019: “Number of social network users worldwide from 2010 to 2021” [Online]. Available: <http://www.statista.com/statistics/278414/number-of-worldwide-social-network-users/>. [Accessed: March 2019].
 117. Sundberg, J. (1987). *The Science of the Singing Voice*. Northern Illinois University Press.
 118. Takayama, K. (2009). Politics of externalization in Reflexive Times: Reinventing Japanese
 119. Reform Discourses through “Finish PISA Success”, *Comparative Education Review*, 54(1): 51–75.
 120. Tao, C., Jiang, J. J., and Czerwinka, L. (2010). “Liquid accumulation in vibrating vocal fold tissue: A simplified model based on a fluid-saturated porous solid theory,” *J. Voice* 24, 260–269.
 121. Thomson, S. L., Mongeau, L., and Frankel, S. H. (2005). Aerodynamic transfer of energy to the vocal folds, *J. Acoust. Soc. Am.*, vol. 118, pp. 1689-1700.
 122. Titze, I.R. (1994), *Principles of Voice Production*, New Jersey: Prentice-Hall.
 123. Trehub, S. (2000), Human Processing Predispositions and Musical Universals, In Wallin, N.L., Merker, B., and Brown, S. (Eds.), *The Origins of Music*, Cambridge, Mass.: MIT Press, pp. 427- 448.
 124. van Dommelen, W. A. and Hazan, V. (2012), Impact of talker variability on word recognition in non-native listeners, *J. Acoust. Soc. Am.*, vol. 132, pp. 1690–1699.
 125. Vendries, Ch. (1999), *Instruments à Cordes et Musiciens dans l'Empire romain*, Paris: L' Harmattan.
 126. Varnet, L., Ortiz-Barajas, M.C., Ramón Guevara Erra, Gervain, J. and Lorenzi, C. (2017), A Cross-Linguistic Study of Speech Modulation Spectra, *J. Acoust. Soc. Am.*, vol. 142, no. 4, pp. 1976-.
 127. Wanderley, M., and Orio, N. (2002), Evaluation of input devices for musical expression: Borrowing tools from HCI, *Computer Music Journal*, vol. 26, no. 3, pp. 62-76.
 128. Wellesz, E. (ed.) (1957), *Ancient and Oriental Music*, in *New Oxford History of Music - vol. I*, London: Oxford University Press.
 129. Weninger, F., Eyben, F., Schuller, B. W., Mortillaro, M. and Scherer, K. R. (2013), On the Acoustics of Emotion in Audio: What Speech, Music and Sound Have in Common, *Front. Psychol.*, vol. 4, pp. 292-.
 130. Wolfe, J., Garnier, M., and Smith, J. (2009), Vocal Tract Resonances in Speech, Singing, and Playing Musical Instruments, *HFSP J.*, no. 3, pp. 6–23.
 131. Wyke, B. D. (1983), Neuromuscular control systems in voice production, in Bless, D. M. and Abbs, J.H. (Eds.) *Vocal Fold Physiology: Contemporary Research and Clinical Issues*, San Diego, CA: College-Hill, pp. 71–76
 132. Yiu, E. and Yip, P. (2016), Effect of Noise on Vocal Loudness and Pitch in Natural Environments: An Accelerometer (Ambulatory Phonation Monitor) Study, *J Voice*, vol. 30, no. 4, pp. 389-393.
 133. Zhang Z. (2010), Dependence of Phonation Threshold Pressure and Frequency on Vocal Fold Geometry and Biomechanics, *J. Acoust. Soc. Am.*, vol. 127, pp. 2554-2562.
 134. Zhang, Z. (2015), Regulation of Glottal Closure and Airflow in a Three-dimensional Phonation model: Implications for Vocal Intensity Control, *J. Acoust. Soc. Am.*, vol. 137, pp. 898-910.
 135. Zhang, Z. (2016), Mechanics of Human Voice Production and Control, *J. Acoust. Soc. Am.*, vol. 140, no. 4, pp. 2614-.