



THE UNIVERSITY *of* EDINBURGH

Edinburgh Research Explorer

Listeners' weighting of acoustic cues to synthetic speech naturalness

Citation for published version:

Mayo, C, Clark, RAJ & King, S 2011, 'Listeners' weighting of acoustic cues to synthetic speech naturalness: A multidimensional scaling analysis', *Speech Communication*, vol. 53, no. 3, pp. 311-326.
<https://doi.org/10.1016/j.specom.2010.10.003>

Digital Object Identifier (DOI):

[10.1016/j.specom.2010.10.003](https://doi.org/10.1016/j.specom.2010.10.003)

Link:

[Link to publication record in Edinburgh Research Explorer](#)

Published In:

Speech Communication

General rights

Copyright for the publications made accessible via the Edinburgh Research Explorer is retained by the author(s) and / or other copyright owners and it is a condition of accessing these publications that users recognise and abide by the legal requirements associated with these rights.

Take down policy

The University of Edinburgh has made every reasonable effort to ensure that Edinburgh Research Explorer content complies with UK legislation. If you believe that the public display of this file breaches copyright please contact openaccess@ed.ac.uk providing details, and we will remove access to the work immediately and investigate your claim.



Listeners' weighting of acoustic cues to synthetic speech naturalness: A multidimensional scaling analysis

Catherine Mayo, Robert A. J. Clark and Simon King

Centre for Speech Technology Research, University of Edinburgh

Informatics Forum, 10 Crichton Street, Edinburgh, EH8 9AB, United Kingdom

tel: +44 (0) 131 650 4434 *or direct dial* +44 (0) 131 651 1767

fax: +44 (0) 131 650 6626

catherin@ling.ed.ac.uk

robert@cstr.ed.ac.uk

simon.king@ed.ac.uk

short title: Multidimensional scaling of synthetic speech perception

ABSTRACT

The quality of current commercial speech synthesis systems is now so high that system improvements are being made at subtle sub- and supra-segmental levels. Human perceptual evaluation of such subtle improvements requires a highly sophisticated level of perceptual attention to specific acoustic characteristics or cues. However, it is not well understood what acoustic cues listeners attend to by default when asked to evaluate synthetic speech. It may, therefore, be potentially quite difficult to design an evaluation method that allows listeners to concentrate on only one dimension of the signal, while ignoring others that are perceptually more important to them.

The aim of the current study was to determine which acoustic characteristics of unit-selection synthetic speech are most salient to listeners when evaluating the naturalness of such speech. This study made use of multidimensional scaling techniques to analyse listeners' pairwise comparisons of synthetic speech sentences. Results indicate that listeners place a great deal of perceptual importance on the presence of artifacts and discontinuities in the speech, somewhat less importance on aspects of segmental quality, and very little importance on stress/intonation appropriateness. These relative weighting differences will impact on listeners' ability to attend to these different acoustic characteristics of synthetic speech, and should therefore be taken into account when designing appropriate methods of synthetic speech evaluation.

Keywords: speech synthesis; evaluation; speech perception; acoustic cue weighting; multidimensional scaling

I. INTRODUCTION

Evaluation of the quality of output produced by a speech synthesis system is an important part of the design of a successful system. At the most basic level, evaluation can give system designers feedback on whether changes to a system engender an overall improvement in perceived quality of the output. At a more sophisticated level, evaluation could identify areas for further improvement. There are currently two main approaches to evaluating the quality of synthetic speech: (i) subjective, or human perceptual, methods, in which participants listen to examples of the synthetic speech and make judgements about quality (usually based on specific criteria), and (ii) objective, or computational, methods, in which models are built to automatically assess improvements to the synthesis system.

There are drawbacks to both types of analysis. Subjective perceptual evaluation requires the participation of numerous listeners in order to achieve statistical validity, and can thus be costly and time-consuming. In addition, listeners do not always achieve high levels of agreement either with each other or with themselves [1; 2]. This lack of reliability can make it difficult to draw meaningful conclusions from the results of subjective evaluation studies.

However, objective evaluation of synthetic speech is also problematic. In fields that make heavy use of objective evaluation measures (such as speech recognition), there is generally a non-opinion-based target with which to compare the output of the system—for example, the output of a speech recogniser can be compared against the text of the input given to the recogniser. Success in such fields is judged based on how well the output matches the desired target—in the case of a recogniser, that is typically how many words the recogniser correctly identifies. Judging the perceived quality of a speech synthesis system, on the other hand, is less straightforward. It is possible to make a direct comparison between acoustic characteristics of a target utterance (what the synthesiser has been asked to produce) and acoustic characteristics of the same utterance spoken by a human

speaker [e.g. 3]. However, this ignores the fact that speech is highly variable (utterance-to-utterance, speaker-to-speaker, etc.) and that there are often many acceptable ways of producing a single utterance [see e.g. 4]. As a result, it is possible for listeners to judge two utterances that are acoustically very different as being the same in terms of quality. The perceived quality of a synthetic speech utterance is clearly not, therefore, simply a matter of the degree to which the physical characteristics of the utterance match the physical characteristics of one single natural speech utterance. Rather, the perceived quality of an utterance is a *psycho*-physical construct, which is closely tied to both the physical, acoustic characteristics of the utterance being judged, *and* to listeners' psychological responses to these characteristics [5]. Thus the success of an objective measure of synthetic speech quality depends on two things: (i) how well the measure models the physical characteristics of the speech being evaluated, and (ii) how well the measure models listeners' behaviour with respect to that speech.

There have been a number of attempts to model human perceptual evaluation of speech. However, to date, only low to moderate correlations have been found between objective and subjective ratings of speech quality for both synthetic speech [6; 3; 7; 8; 9; 10; 11] and for natural speech [12]. This lack of correlation seems to stem not from an inability to model the physical characteristics of speech, but from difficulties in modelling human perceptual responses. As noted above, listeners' subjective evaluations often lack strong inter- and intra-rater consistency. In addition to making interpretation of subjective evaluations difficult, such inconsistent behaviour is inherently less amenable to objective computational modelling.

The question to be answered, therefore, is why subjective evaluation behaviour is so inconsistent. There is an understanding in the field of speech synthesis of what paradigms are currently available for testing perceived quality [e.g. 13], and an understanding of the need for principled use of evaluation paradigms [e.g. 14]. This knowledge should allow

for studies to be carried out in which listeners' responses are more consistent. However, despite the awareness of testing paradigms, there is a general lack of understanding of the psycho-acoustic processes that underpin the complex task of auditory evaluation of synthetic speech. In particular, it is not clear from research to date what the exact relationship is between the acoustic characteristics of synthetic speech and listener responses to these characteristics. Without an understanding of this relationship, it is very difficult to choose evaluation methods in a principled manner, and as a result, raters may be asked to carry out tasks which their perceptual systems cannot physically perform: Naturally this could easily result in inconsistent or unexpected rating behaviour.

In fact, what little is known about the psycho-acoustic task of speech evaluation does point to the possibility that listeners are often asked to perform perceptually challenging tasks. One of the dominant state-of-the-art speech synthesis techniques (and the one which we use in this work) involves *unit selection*. In this method, a large database of natural speech is labelled in terms of units (usually phones, diphones or half phones). An automatic search is then performed to find the best units or sequences of units from the database, and these units are concatenated together to form the target utterance [see e.g. 15]. This method has resolved many of the issues surrounding gross segmental quality and intelligibility that caused problems for earlier, rule-based synthesis systems, and has thus allowed researchers to move on to fine-tuning individual sub-segmental characteristics [e.g., discontinuities at concatenation points, 16] or supra-segmental characteristics [e.g., intonation, 17]. As a result, speech synthesis evaluation is now much less about determining the overall intelligibility or overall acceptability of synthetic speech, and more about evaluating the quality of a single one of these sub- and supra-segmental characteristics. Unfortunately, research has shown that listeners sometimes find it difficult to focus on just one characteristic, particularly when faced with complex acoustic stimuli such as speech. For example, it has been found that listeners are much less able to rate intonation consistently when it varies simultaneously with many other acoustic characteristics than

when intonation is the only acoustic characteristic of the stimulus set to be varied [1]. This would suggest that it may be beyond listeners' abilities to evaluate just one sub- or supra-segmental characteristic of a synthetic speech utterance.

However, concluding that subjective evaluation of single characteristics of multidimensional stimuli is impossible assumes that listeners give equal perceptual attention or "weight"¹ to all acoustic characteristics and cannot disentangle individual characteristics from each other. This does not seem to be the case. Instead, it appears that listeners give more weight to some characteristics than others. Evidence for this comes from studies that have shown, for example, that the addition of synthetic intonation to natural speech segments was more detrimental to listeners' speech quality ratings than was the addition of synthetic segment duration [18], suggesting that for this listening situation, intonation was weighted more heavily than segment duration. Other studies have found that listeners have been more influenced by intonation appropriateness than by segmental quality [19], less influenced by intonation naturalness than by segmental quality [20; 19], and less influenced by intonation variation than by overall synthesis quality or presence of discontinuities due to concatenation [21; 22]. In fact, studies from across perceptual evaluation, speech perception and general auditory perception research [e.g. 23; 24; 25; 26; 27; 22] indicate that listeners have very complex hierarchies of weighting for different acoustic characteristics, and that these can differ depending on the stimuli (e.g., speech versus non-speech, natural speech versus synthetic speech, first language versus second language, etc.). This would suggest two things. First, it should be straightforward to devise an evaluation paradigm to assess the quality of those acoustic characteristics that are by default weighted heavily by listeners. Second, if listeners are asked to evaluate characteristics that are lower on their weighting hierarchy, they could easily be unduly influenced by those characteristics that are weighted more heavily. This could leave a large number of acoustic characteristics inaccessible for evaluation.

Fortunately, research suggests that even those acoustic characteristics that are weighted less heavily could be accessible to listeners, under specific circumstances. Studies have found that listeners' auditory weighting patterns are not necessarily fixed, but rather can be flexible, with most adult listeners changing their weighting hierarchy to meet the needs of the stimuli or the listening situation. A number of studies have examined the ease with which listeners can detect spectral discontinuities at join points (points at which units of speech from two different source utterances are concatenated). These studies have found differences in rates of detection for male and female synthetic voices, and for joins in different segmental contexts [e.g., monophthong versus diphthong, pre-vocalic consonants versus post-vocalic consonants, front versus back vowels, etc., see 16; 28]. Additionally, listeners' default weighting patterns for speech and non-speech auditory stimuli have been shown to be capable of change due to outside influences: listeners' acoustic characteristic hierarchies have been changed by training [25; 29; 30], by manipulation of the stimuli to mask certain dimensions (e.g., simultaneous white noise [31; 32] or reverberation [33]) or to enhance certain dimensions [34; 35], by manipulation of the distribution of the acoustic dimensions across the whole stimulus set [23] or by manipulation of listeners' conscious focus of attention (e.g., by presenting the rating task simultaneously with another task [36]). It appears possible, therefore, that if listeners do not by default give adequate attention to the acoustic dimension under investigation, appropriate methods could be designed to allow listeners to refocus their attention.

However, both (i) taking advantage of listeners' default hierarchy of weighting to evaluate dimensions that are weighted heavily, and (ii) developing methods that will allow listeners to attend more closely to dimensions that receive less weight, presuppose a better understanding of the default weighting given to potential acoustic dimensions of synthetic speech than is currently available. The study described here is an attempt to provide a more complete account of the acoustic dimension weighting behaviour of listeners with respect to unit-selection synthetic speech.

A. The current study

The current study aimed to determine which of a large number of acoustic characteristics of synthetic speech play greater or lesser roles in listeners' responses to such speech. To address this aim, the study made use of multidimensional scaling (MDS) techniques [37], which have been used successfully to determine weighting patterns for numerous auditory domains: e.g., complex non-speech sounds [38; 23], coded speech [39], segmental contrasts [40], voice quality [41], and musical timbre [42].

The goal of MDS techniques is to create a geometric, spatial representation of the proximities between objects—either actual or perceived proximities. In the current study, synthetic utterances were presented pairwise to listeners, who were asked to decide whether each pair was “similar” or “different” in terms of naturalness. The total number of times a given pair of utterances received a “different” response from listeners was taken as the proximity value for (i.e., the psycho-physical distance between) that pair of utterances.

The output of an MDS analysis is an n -dimensional stimulus space or map, in which each object—here, each of the synthetic utterances—is represented by a single point. The first key characteristic of an MDS map is that two objects that are physically or psycho-physically close are represented by two points that are close on the map, while two objects that are physically or psycho-physically distant are represented by two points that are farther apart on the MDS map. The second important characteristic of an MDS map is that the dimensions that make up the space often correspond (directly or indirectly) to the physical or psycho-physical dimensions used most heavily by subjects to make their proximity judgements. Thus by examining and interpreting both the MDS map dimensions and the configuration of the points within those dimensions, it should be possible to determine the underlying characteristics of the objects represented in the space that led to subjects' responses to those objects. For the current study, this means that analysis of the MDS space should allow us to identify some of the acoustic characteristics of

synthetic speech that relate to listeners' judgements of that speech. Analysis and interpretation of an MDS stimulus map can be done in a number of ways. For many two- and potentially three-dimensional spaces it may be possible to interpret the space by visually examining the distribution of the objects within the space and trying to find any underlying pattern(s) in the organisation of the points. For stimulus spaces with four or more dimensions (i.e., numbers of dimensions that are less straightforward to represent in a way that can be examined visually) or for more complex two- and three-dimensional maps, it is often necessary to make use of other methods, such as cluster analysis, or multiple regression (i.e., regressing relevant characteristics of the objects in the MDS space on the coordinates of the points in the MDS map). In the current study, it was not clear before MDS analysis was carried out whether the MDS map resulting from listeners' "different" responses would be straightforward to analyse either visually or auditorily. Therefore, in addition to the listening experiment, a thorough acoustic analysis of the synthetic utterances used in the listening experiment was also carried out, with the aim of identifying and quantifying a large number of acoustic characteristics of the synthetic speech with which to compare the MDS stimulus space.

A small pilot study using MDS methods [43], indicated that listeners are influenced by at least two general characteristics when determining how similar utterances sound in terms of naturalness: segmental information (likely to include such things as appropriateness of segments, audibility of joins at unit edges, and number of discrete units used to create the utterance), and prosodic information (most likely including intonation, duration/timing, and stress). The small number of utterances used in this pilot meant that stimuli tended to be placed at extremes in the stimulus space, which in turn meant that it was not possible to determine whether listeners might respond to more fine-grained aspects of the two general characteristics identified. However, despite this, the study showed clearly that MDS techniques are likely to serve as a useful method of identifying the acoustic characteristics of synthetic speech that most influence listeners' perceptual evaluation be-

haviour. The current study aimed to expand on the pilot by increasing the number of utterances and the number of listeners in order to engender a more fine-grained stimulus space, and thus identify more precise characteristics which might be involved in the perception of naturalness in synthetic speech. Note that the experiment described in the current study focused on listeners' evaluation of unit-selection speech produced by a typical unit selection speech synthesis system, Festival [44]; further studies should endeavour to examine listeners' responses to other unit-selection synthesis systems, as well as systems based on other techniques [e.g. statistical parametric models, 45].

In summary, the aim of the current study was to better understand the relationship between the acoustic characteristics of unit-selection synthetic speech and listeners' responses to such speech. In the current study, 24 synthetic utterances were presented pairwise to listeners for assessment as either "similar" or "different" in naturalness. Listeners' responses were analysed using MDS techniques, and the resulting spatial map was correlated with acoustic characteristics of the 24 synthetic utterances. These methods should allow us to determine which of the acoustic characteristics identified in the synthetic speech played a role in listeners' perception of synthetic speech quality, and should also allow us to begin to evaluate these characteristics' relative importance in determining how similar/dissimilar listeners perceived the stimuli to be.

II. METHOD

A. Participants

Thirty adults (22 female, 8 male ²) ranging in age from 18 years to 43 years (average age 24 years) took part in this experiment. All were monolingual native speakers of English who were living in Edinburgh, UK at the time of testing (self-reported dialects broke down as follows: Southern British English: 12; Northern English: 6; Scottish: 5; Irish:

3; North American: 4). All subjects reported themselves and their siblings as being free from speech/language disorders and dyslexia, and all reported being free from hearing deficits. Immediately prior to testing, each participant was required to pass a hearing screening at ≤ 35 dB HL in the range 125 Hz–8 kHz. All listeners were paid for their time.

B. Stimuli

This stimuli were generated using the multisyn module of the Festival speech synthesis system, which is a typical unit-selection text-to-speech synthesiser. Unit-selection synthesis uses a large database of natural speech recordings of a single speaker, from which suitable units are selected and concatenated to create any required target utterance. In Festival, the units are diphones: Units of speech that begin halfway through one phone and finish halfway through the subsequent phone. The speech database is automatically phonetically labelled and from the phone labels, diphone boundaries are derived.

At synthesis time, Festival first performs text analysis, resulting in a phonetic transcription of the utterance to be created; this transcription is then converted to a sequence of diphones known as ‘targets’. Then the best sequence of diphone instances from the database is selected and concatenated to create the output speech. As in other unit-selection synthesisers, the measure of ‘best’ has two components. One, the *target cost*, measures the degree of mismatch between each candidate diphone available in the database and the corresponding target diphone. The mismatch is judged in terms of linguistic context—that is, the difference between the phonetic and prosodic contexts of candidate and target—on the assumption that fewer mismatches will result in perceptually more natural speech output. The other component is the *join cost*, which measures how well each pair of consecutive candidates will concatenate; join cost is typically calculated based on properties of the speech signal such as spectral discontinuity, energy and F0. It is assumed that smaller differences in these acoustic properties at join points between candidate units will

be less perceptually noticeable and thus more natural-sounding. A search must be performed in order to find, from amongst all possible sequences of candidates, the one that minimises the sum of target and join costs. By defining the join cost between candidates that are consecutive in the database to be zero, the search is biased towards selecting contiguous stretches of speech. This means that there will generally be fewer concatenation points in the output speech than there are phonemes in the target utterance. Festival additionally attempts to avoid using extreme outlier candidates (in terms of duration, for example) and can substitute any missing diphones with alternatives (i.e., in cases where there are no instances of a required diphone type in the database).

The text material for the synthetic utterances used in this study came from the “phonetically-compact” portion of the TIMIT corpus [46]. This portion was hand-designed “to include a basic phonetic coverage and interesting phonetic environments” with special regard for “phoneme pair coverage, consonant sequence coverage and the potential for applying phonological rules both within words and across word boundaries” [47, p. 2162]. The reason for using this material was to ensure the test utterances would contain a reasonably wide variety of acoustic-phonetic phenomena that might impact on perceived speech quality.

All 450 of the phonetically-compact TIMIT sentences were synthesised using the Festival 1.96 multisyn engine with a female, Standard Southern English voice (cstr _rpx _nina _multisyn) ³. Twenty-four of the resulting utterances were chosen as the test utterances for the current study. All 24 were chosen at random, with the proviso that they be similar in duration (both in milliseconds and in total syllables; this was done in order to minimise any effect of length of utterance): The test utterances were therefore between 2.60 sec and 2.70 sec in duration (average duration: 2.65 sec) and all contained approximately 12 syllables (range: 9-14 syllables; average: 12 syllables). None of the 24 test utterances were manipulated during or after synthesis. All variations in the acoustic characteristics of the

synthetic speech and in the the perceived naturalness of that speech were thus uncontrolled, and resulted only from the diphone coverage of the voice database used and the performance of Festival.

In addition to the 24 test utterances, a set of 10 practice utterances and a set of 6 familiarisation utterances (both also from the phonetically-compact TIMIT sentences) were also synthesised.

The 10 practice utterances were between 2.26 sec and 2.27 sec in duration (average duration: 2.26 sec) and all contained approximately 10 syllables (range: 6-13 syllables; average: 10 syllables). Again, none of the practice utterances were manipulated during or after synthesis. All of the sentences were lexically different from those presented in the main test.

The 6 familiarisation utterances were designed to be presented in three specific pairs to illustrate extreme examples of “similar” and “different” with regard to naturalness. These pairs comprised one pair of “different” utterances, which contained an extremely natural utterance and an extremely unnatural utterance, and two pairs of “similar” utterances, one of which contained two extremely natural utterances, and the other of which contained two extremely unnatural utterances. The 6 utterances that made up these familiarisation pairs were synthesised using the same voice that was used to create the practice utterances and the main test utterances. In order to create the 3 utterances designed to illustrate extremely unnatural speech, only a subset of the voice database was employed by the synthesiser (400 utterances, rather than the full database of 2000 utterances). Restricting the diphone coverage available to the synthesiser degrades the output speech quality. The three utterances designed to illustrate very natural synthetic speech were chosen by the experimenters as being representative of the most natural utterances that this synthesis system was capable of producing from the given voice database. None of the familiarisation utterances were lexically the same as those in either the practice or the

main test.

1. Test sets

In order to acquire the proximity values necessary for MDS analysis, it was necessary to present the synthetic speech utterances in pairs, to determine the perceived degree of similarity/dissimilarity in naturalness between each utterance and every other utterance. To create the utterance pairs for presentation to the listeners in this study, each of the 24 test utterances was paired with every other test utterance, in both directions (i.e., both AB and BA), resulting in 576 possible pairs of utterances (24 of which were same-utterance pairs, e.g., utterance 1 paired with utterance 1). Pilot testing [43] had revealed that listeners could rate 168 pairs of utterances in approximately 30-40 minutes, and that this number of utterance pairs and duration of testing was approaching the upper limit of what listeners could tolerate in a single test session. Therefore, to avoid listener fatigue and any possible associated drop in performance, it was decided that the complete set of 576 utterance pairs would be divided into three smaller test sets, and that each listener would hear only one of these three test sets ⁴.

The subdivision of the complete set of utterance pairs was done as follows. A 24x24 grid of all possible pairs of utterances was created. This complete set of pairs was then subdivided into 9 smaller subsets by superimposing a 3x3 grid over the larger 24x24 grid. Each of the 9 subsets therefore contained 64 pairs of utterances (e.g., utterances 1-8 x utterances 1-8, or utterances 1-8 x utterances 9-16, or utterances 1-8 x utterances 17-24, etc.). A Latin square arrangement was then used to allocate the utterance pairs in the 9 subsets into one of three test sets: The cells in the 3x3 grid were labelled A-B-C, B-C-A, and C-A-B by row, and all of the utterance pairs in the cells labelled 'A' were put into a single test set; the same was done for the utterance pairs in the cells labelled 'B' and those in the cells labelled 'C'. This resulted in 192 utterance pairs per test set. At this point all

of the same-utterance pairs were removed from the test sets (there were 8 same-utterance pairs per test set). This was done because all utterance pairs *except* the same utterance pairs contained two lexically different sentences. Including a small number of pairs of lexically identical sentences in a test set made up predominantly of different-utterance pairs could have biased listeners' responses in an unpredictable way. The final version of each of the three test sets therefore contained 184 pairs of utterances.

This design resulted in a manageable test set size for each listener, while maintaining maximal variation in what each listener heard (each test set contained each of the 24 test utterances paired with one third of the 24 test utterances, less the same-utterance pairs), and ensuring complete coverage of the test set (i.e., ensuring that every possible different utterance pair was heard by at least a portion of the subjects).

C. Presentation of stimuli

All participants were tested individually in a sound-treated room. The stimuli were presented over closed-back headphones (Sennheiser Evolution H270, frequency response 12-22000 Hz) with volume set at a constant, comfortable listening level. Testing took place in one 40-45 minute session, with two short breaks part way through testing.

Presentation of the pairs of stimuli was controlled by a suite of computer software [E-Prime, 48] that also recorded listeners' responses. The listeners' task was to listen to the two utterances in a pair, and indicate whether they were "similar" or "different" in terms of their naturalness (that is, with regard to how much like real speech each utterance seemed to be). Importantly, the participants were not instructed to listen to any one acoustic characteristic of the stimuli, or to any specific psycho-acoustic concept that might make up the overall construct of "naturalness" (e.g., "listening effort," "pleasantness," "pronunciation," etc.) such as used in mean opinion scale (MOS) tests of synthetic speech

quality [49]. Instead, the listeners' task was simply to make a binary decision ("similar" or "different") about the degree of similarity in naturalness of each pair of stimuli.

Before testing, the participants were given an opportunity to listen to the three familiarisation pairs of synthetic sentences. As noted above in Section B., these pairs were designed to introduce listeners to extreme examples of "similarity" and "differentness" in naturalness, and were explicitly identified to the listeners as such (that is, listeners were told "These two sentences are similar in terms of naturalness."). No responses were collected from listeners during this period. Following this familiarisation, a pre-test was administered, to ensure the listeners understood the task. This pre-test consisted of the 10 practice utterances, again presented in pairs. Only 10 of the possible pairs of utterances were presented to each listener, in random order. No feedback regarding naturalness was given during this pre-test; additionally, as there was no pre-defined "correct" response to whether a given pair of utterances sounded "similar" or "different", no listeners were excluded on the basis of answers given during this pre-test period.

During the test proper, each participant heard one test set of 184 pairs of utterances (see Section 1. for a description of these test sets). The interval between the presentation of each member of a pair was 5000 msec (onset to onset). Responses were not timed, and the presentation of the next pair of utterances began 2000 msec following the entry of the preceding response. Breaks were given following the presentation of the 61st and the 122nd pairs; the duration of these breaks was controlled by the participants.

D. Acoustic analysis of stimuli

It was unclear at the start of the current study which potential acoustic characteristics might play a role in listeners' judgements of the stimuli used in this study. It was therefore important to undertake as complete an analysis as possible of the stimuli to determine

the spread of acoustic characteristics across the utterances and subsequently to use this to interpret the MDS configuration.

Two types of acoustic analysis were carried out—automatic and manual—on the 24 synthetic utterances used in the listening experiment, and where appropriate, on the natural utterances used to create the synthetic utterances. All analyses are listed in Tables 1–3 . The automatic analyses included measurements carried out by Festival in the process of synthesising the 24 stimulus utterances, which included target cost of individual units, join cost for joining two individual units, and number of “bad” units (those units that are not statistically representative of other units containing the same segments from the database) or missing units (those diphones that cannot be found in the database and are replaced by something similar, e.g., a diphone containing a long vowel may get backed-off to a unlengthened version).

Also in the automatic analysis were those measures that could be automatically calculated based on Festival’s measurements. These were: total cost (a measure based on overall target cost and overall join cost for an utterance); target cost for different types of target diphone (e.g., consonant-consonant, vowel-consonant, etc.); join cost for different types of join (e.g., joins in vowels, stop consonants, fricatives, etc.); proportion of units or joins with different values of target or join cost. Finally, the automatic analysis included a number of other automatically calculable measurements, for example: duration of an utterance in milliseconds, whether the initial and final diphones were taken from initial/final position in the source utterance, etc.

In addition to these automatic analyses, five different manual analyses were also carried out, three on the 24 synthetic utterances, and two on each of the natural source speech utterances in the voice database from which units were selected to create the 24 synthetic utterances. The natural source utterances were first analysed for number of transcription/pronunciation errors per synthetic utterance. The database of natural speech utter-

ances is automatically phonetically labelled, based on the text that the database speaker was given to read. If the speaker deviates from this text in some unexpected way, this will cause a mismatch between the database speech, and the phonetic label, which in turn will cause problems if these mismatched units are chosen for use in synthesis. Errors in this category included (i) mispronunciations by the database speaker (e.g., pronouncing “Maryland” as “Marilyn”: results in a section of database labelled /d/ when no /d/ has been produced); (ii) automatic segment labelling not accounting for pronunciation variation (e.g., speaker producing “window” as /wində/, but the lexical entry only containing an unreduced diphthong in the second syllable: results in errors in contexts where reduction of the diphthong is not possible, e.g., “gâteau”); (iii) the incorrect labelling of any non-silent stretch of speech as “silence”.

The natural source utterances were also manually analysed for number of segmentation errors per synthetic utterance. As noted above, Festival makes use of segmentation information to determine the points at which a unit should be cut out of a source utterance for use in a target utterance. Note that cuts/joins are not made at boundaries between phones, but at diphone boundaries. These points are, however, derived from phone segmentation information. The accuracy of the initial segmentation of the natural speech is therefore important to the success of the splicing procedure in unit selection speech synthesis. Accuracy of the segmentation of the natural speech database was evaluated using the principles set out by Turk, Nakai and Sugahara [50]. Segmentation errors were only noted if they occurred at the edge of a unit boundary (e.g., if the automatic segmentation was found to be incorrect for the boundary between /ɪ/ and /g/, but the entire sequence /a.ɪgə/ was selected from a source utterance for use in a target utterance, then this error in segmentation was ignored).

The synthetic utterances were manually analysed for: (i) number of inappropriate diphones, (ii) number of detectable joins, and (iii) number of areas of incorrect stress/intonation

(both lexical and phrasal). Classed as inappropriate diphones were any diphones taken from a segmental context in the source utterance which would be segmentally incompatible with the target utterance (e.g., a post-vocalic /ɪ/ diphone taken from the word “prioritization” to be used in the post-stop context in the word “predicament”: the post-vocalic context would engender a much longer /ɪ/ than would be appropriate for the post-stop context). Detectability of joins were analysed by means of spectrograms and spectra, and “detectable joins” included all joins in which there was at least one of: (i) a mismatch in vowel formant frequencies at the join point, (ii) a mismatch in pitch at the join point, (iii) a mismatch in intensity at the join point. Stress/intonation was marked as inappropriate in the following circumstances: (i) if the source utterance carried/did not carry lexical stress on the to-be-used vowel and this was inappropriate for the target context (e.g., the lexically stressed syllable /bɪ/ from the word “BRItish,” being used in the target word “UPbringing,” where the syllable should not carry lexical stress), (ii) if the source contained stress/intonation appropriate for a particular location in an utterance and this was inappropriate for the location in the target utterance (e.g., a sequence with phrase-final falling intonation used phrase medially) and (iii) if the source carried/did not carry contrastive stress and this was inappropriate for the target context.

III. RESULTS AND DISCUSSION

The proximity values dictated by the listeners’ “similar”/“different” responses to the 24 synthetic utterances used in this study are best represented by a three-dimensional MDS map [$R^2 = .75$, residual stress = .18]. As the rotation of the MDS space is potentially arbitrary, principal component analysis (PCA) was used to ensure any principal components were aligned with the axes; PCA did not further rotate the space⁵. Figures 1, 2 and 3 show three two-dimensional pairwise axis projections of the three-dimensional MDS map.

A. Visual, auditory and cluster analysis

Kruskal and Wish [37] advocate attempting a basic visual analysis of MDS maps in the first instance, before subjecting the map to further analysis. In the case of the current study, this basic analysis also included an auditory analysis and a cluster analysis of the distribution of the synthetic utterances within the MDS space.

From an initial visual and listening analysis, it is clear that listeners are able to make a certain number of important distinctions with regard to the acoustic characteristics of the synthetic speech used in this study. We identified five general groups into which listeners appear to have classified the utterances. More fine-grained groupings may be present in this data, however they are not immediately present from this first-pass visual and auditory analysis: They may become evident after further (statistical) analysis (see below). These five groupings are:

- A** very high quality utterances, with (i) very few poor units (i.e., few units with transcription/pronunciation errors, inappropriate diphones, high target costs, etc.) and (ii) very few joins (i.e., the entire utterance was made up of long stretches of contiguous speech from the database), for example, utterances 10, 11, 12;
- B** utterances with very poor quality joins (as indicated by high join costs and audible discontinuities at joins), for example, utterances 15, 18, 21, 22, 23;
- C** utterances with many joins, but all of them of relatively high quality (indicated by low join costs and absence of audible discontinuities at joins), for example, utterances 6, 11, 20;
- D** utterances with a handful of poorly selected units (containing instances of transcription/pronunciation errors, etc.), for example, utterances 4, 5, 17;
- E** utterances with many poorly selected units, for example utterances 1, 2.

To attempt to validate these groupings as being meaningfully different in terms of their

location in the MDS space we performed k-means clustering. The within-groups sum-of-squares for different cluster sizes suggests that that four to six clusters are present. Clustering with six clusters produces groupings which match very closely our initial interpretation, as illustrated in Figure 4. In the top left of Figure 4 we see a cluster containing the stimuli 3, 8, 9, 10, 12, 13 and 16, which matches our group A. In the top right of Figure 4 we see two clusters: The first containing stimuli 15, 18 and 22 and the second containing stimuli 14, 19, 21, 23 and 24. These two clusters together match our group B (note that these clusters merge, if clustering is performed with only 4 clusters). Stimuli 6, 7, 11 and 20 form a cluster at the bottom centre of Figure 4, corresponding to our group C and stimuli 4, 5 and 17 form a cluster to the left of the previous cluster, relating to group D. Finally stimuli 1 and 2 form a cluster that matches the stimuli described by group E (again if only 4 clusters are requested, these last two clusters merge).

These groupings allow us to make the following inferences. First, listeners seem able to differentiate between utterances on the basis of the *number* of joins in an utterance (many, even if low cost/high quality, as in group C, versus few, as in group A) and the *quality* of joins across an utterance (poor quality/high cost, as in group B, versus better quality/lower cost, as in group C). Listeners also seem able to differentiate between utterances on the basis of number of poor quality units, with at least three different identifiable levels of unit quality: (i) many, poorly selected units, as in group E, (ii) a handful of poorly selected units, as in group D, (iii) few poor quality units, as in group A. It is not, however, evident from these general groupings whether listeners are able to, or do, use any aspect of intonation or prosody in their judgements of the naturalness of these utterances, as described below.

B. Regression analysis: Accounting for the x -, y -, and z -dimensions

Following the above visual, auditory and clustering analyses, to further investigate the resulting spacial configuration, all of the acoustic characteristics of the synthetic speech utterances were regressed onto the x -, y - and z -dimensions of the MDS map. This analysis is performed to determine whether any meaningful interpretation of the axes can be made.

Note that while the configuration of utterances along both the x - and y -dimensions were found to correlate with a number of the acoustic characteristics at $p = 0.05$, the distribution of utterances along the z -dimension did not correlate with any characteristics at this level of p . At $p = 0.06$, the z -dimension was found to be accounted for by 6 acoustic characteristics; these results should be treated more cautiously, given the weaker relationship between the dimension and the characteristics. The results of all of these regressions can be found in Table 4, Table 5 and Table 6. R^2 change analysis showed that the addition of each new predictor made a significant contribution to the model for each dimension ($p = .01$ or less for each addition); Durbin-Watson analysis confirmed independence of errors.

An examination of these correlations with respect to the general categories and distinctions identified in the auditory and visual analysis described above shows that all three of the dimensions of the MDS map were accounted for by some aspect of join quality/quantity, and by some aspect of unit appropriateness, while only one dimension was accounted for by any aspect of intonation/prosody. However, each dimension was accounted for by very different aspects of, and amounts of, each of these general acoustic characteristics, as described below.

1. *The x -dimension: Overall join quality/quantity*

As can be seen in Table 4, the distribution of points along the x -dimension was predominantly accounted for by measures of join quality and/or quantity, namely the number of joins (normalised for the number of possible joins, as this will vary across utterances), and average join cost across an utterance. Together these two measures can be thought of as capturing some sort of high-level or global impression of the nature (i.e., number and overall quality) of the joins in an utterance. Additionally variability in the x -dimension was accounted for by one aspect of unit appropriateness: The average target cost of consonant-consonant diphones.

2. *The y -dimension: Join distribution and detectability*

Table 5 shows that, like the x -dimension, variation in the y -dimension was accounted for predominantly by measures of join quality/quantity. However, it appears that where the x -dimension appeared to reflect some sort of overall impression of join behaviour, the y -dimension reflects more fine-grained detail regarding the distribution and detectability of joins in an utterance. First, variation in the y -dimension is predicted by the number of joins at 2-diphone intervals in an utterance (as reflected by “number of double units”). This is clearly some sort of measure of join number/distribution, and may also reflect join detectability. For example, this 2-diphone interval could, depending on the word, be roughly equivalent to a mora or syllable (for example, in the word “hippopotamus,” a join every 2-diphones would place a join either in the middle of every consonant, or in the middle of every vowel). Given the perceptual importance of mora- and syllable-like units [see e.g., 51], the presence of joins (and thus potentially join discontinuities) at these intervals might be particularly salient to listeners. Second, variation along the y -dimension is accounted for by the proportion of joins with a join cost of between .31 and .40, which is the second highest (i.e., second worst quality) of the four divisions of join

costs made in this study, and is thus simply a measure of the number of relatively poor quality joins in an utterance, normalised for the overall number of joins in the utterance. It is interesting to note that it is not the highest join costs (i.e., the join costs that reflect the least smooth, most audible joins) that are the most perceptually important in this study. This could, however, have something to do with the distribution of the join costs across the artificial levels utilised in this study. The first two divisions of join costs, covering the two lowest costs (less than .20 and .21-.30) appear in roughly equal proportion across all 24 of the utterances used in this study (40.58% and 41.79% respectively). The highest level of join cost (more than .41), on the other hand, is rarely seen (only 4.79% of all joins had this cost). Thus it is possible that joins in the range .31-.40, which accounted for 12.83% of all joins in this study, might be more perceptually noticeable. Variation along the y -dimension is also predicted by the segmental location of joins (as reflected by “ratio of joins in consonants to joins in vowels”). This last correlation is not surprising: Studies that have found that listeners’ ability to detect discontinuities at join points depends on the segmental context of the join in question, with joins in vowels more detectable than joins in consonants [e.g. 16; 28]. Finally, again, like the x -dimension, there is also a small role for appropriateness of units (specifically consonant-vowel units) in accounting for variability in the y -dimension.

3. *The z -dimension: Unit appropriateness and prosody*

Table 6 shows that in contrast to the x - and y -dimensions, variance along the z -dimension was accounted for more by measures of unit appropriateness than by measures of join quality/quantity. First, variance in the z -dimension is accounted for by a characteristic (“unit errors”) that combines number of transcription/pronunciation errors, number of inappropriate diphones, number of Festival-identified “bad units”, and number of missing diphones. This measure reflects those unit errors in which segmental information is

outright incorrect or missing, rather than simply a poor match to the target. It should be pointed out that the salience of this type of segment error might not be strictly to do with segmental incorrectness: It is probable that all of these errors would additionally cause some sort of discontinuity in the signal at a join point, thus possibly making these errors particularly perceptible. Second, the z -dimension was also predicted by the average target cost across an utterance for vowel-initial units (vowel-vowel diphones and vowel-consonant diphones), potentially reflecting something to do with the distribution or detectability of unit errors.

An aspect of join cost does, however, play some role in accounting for variance along the z -dimension, specifically the proportion of joins with a join cost of less than .20, which is the lowest (i.e., best quality) of the four divisions of join costs made in this study. It is possible that this aspect of join cost in fact corresponds with the large number of unit-appropriateness predictors of the z -dimension, because joins appear to play a much larger overall role in perceived naturalness than units. It may thus be necessary to have low join costs (i.e., good, or less detectable joins) in order to be able to judge the appropriateness of segmental units.

Finally, the z -dimension was the only dimension that was accounted for by any aspects of intonation and/or prosodic appropriateness. First, the z -dimension was accounted for by a measure indicating whether or not the final diphone in a synthetic utterance was taken from utterance-final position in the source utterance (that is, whether or not the final diphone carried the appropriate utterance-final stress and intonation pattern). This is unsurprising given the importance of utterance-final prosody (which can include any of the following: Lengthening of syllables, pitch declination, pausing) in the perception of phrase boundaries [52] (see [53] for a review). Second, a small percentage of the variation in the z -dimension was also predicted by a measure of the number of instances of incorrect stress/intonation, across an utterance. These two measures capture both a

global/high-level picture of the prosody errors across an utterance, and a more detailed picture of one important aspect of prosody, namely utterance final prosody.

C. Regression analysis: Accounting for the acoustic characteristics

The linear regression described above in Section B. was performed with the x , y and z axes as the response variables because these were the axes on which subjects distinguished the synthetic speech stimuli and, as noted above, the main aim of the current study was to devise an acoustically valid account of the listeners' perceptual evaluation behaviour. However, it is also possible to analyse the space from an alternate perspective, and perform linear regression with each acoustic characteristic in turn as a response variable and the x , y and z values as input variables. The results of such an analysis can be interpreted as follows: MDS analysis demonstrates that the listeners who participated in the current experiment made use of certain psycho-acoustic characteristics of speech that correspond to the x , y and z dimensions of the MDS space. Given that this is the case, if we were to present the same listeners with speech samples about which we knew nothing (with regard to acoustic characteristics), the x , y and z dimensions of the current MDS space could be used to predict variation in certain of the acoustic characteristics of the new speech.

The seven out of 52 acoustic characteristics for which a significant amount of variance can be accounted with a linear model consisting of a constant and the three axis values are reported in Table 7. The x dimension contributes to the variation in the largest number of acoustic characteristics. Five of these are related in some way to joins, namely the number of joins in an utterance (normalised for the number of possible joins), the length (in diphones) of the longest contiguous section of database speech in a synthetic utterance, the number of contiguous sections of database speech that were four diphones long or longer in a synthetic utterance, the number of number of joins at 2-diphone intervals, and the number of joins with costs between .31 and .40. These all indicate the strong rela-

tionship between the x dimension and measures of join acceptability. In particular these measures correspond predominantly to number of joins and/or join location across utterances, with one measure representing severity or detectability of join (number of joins with costs between .31 and .40). Note, however, that (as discussed above in Section 2.) the number of joins at 2-diphone intervals could also be considered to be a measure of join detectability as well as join number/distribution. The x dimension also accounts for variation in two measures of unit acceptability: average target cost across an utterance, and the average target cost across an utterance for consonant-consonant diphones.

The y dimension contributes to the variation in only two of the above characteristics, both related to joins, and more specifically, both related in some way to join detectability: The number of number of joins at 2-diphone intervals, and the number of joins with costs between .31 and .40. The z dimension did not make any contribution to the variation in the analysed acoustic characteristics at the significance levels reported here.

IV. CONCLUSIONS

The aim of the current study was to make use of multidimensional scaling (MDS) techniques to help identify those acoustic characteristics of unit-selection synthetic speech that might play a role in listeners' evaluation of such speech.

The first point to note is the overall importance of joins. For all three dimensions of the MDS space at least 20% of the variance can be accounted for by some aspect of joins; for some of the dimensions this percentage is much higher. It is perhaps unsurprising that joins should receive the highest perceptual weight, as in natural speech the presence of audible "pops" and "clicks" in the speech stream would be extremely unexpected. What is more surprising is the fact that both number of joins in an utterance and quality of joins make independent contributions to listeners' perception of an utterance's naturalness. In

addition, listeners appear sensitive both to the general or overall behaviour of joins across an utterance (such as captured by “average join cost”) as well as to more specific, context-dependent behaviour of joins (such as captured by “ratio of consonant-located joins to vowel-located joins”). Interestingly, this last finding replicates earlier work showing that listeners are sensitive to differences in joins in different phonetic contexts [16; 28], but expands the finding to full sentences (the original work was done on isolated syllables). It is clear from the current study, therefore, that it is not simply the case that listeners find unexpected spectral discontinuities in the speech stream extremely salient. Instead, it appears that listeners are highly sensitive to many aspects of joins, and differentiate between them at very subtle levels.

In contrast, none of the three dimensions of the MDS space correlate with general descriptions of unit acceptability, such as might be captured by “average target cost”. Instead, different dimensions of the MDS space seem to be affected differently by more specific and context-dependent aspects of unit acceptability, namely target cost in each type of diphone (e.g., consonant-consonant, consonant-vowel, etc.), as well as some measure of auditorily detectable unit errors. This suggests one of a number of things. One possibility is that the general descriptions of unit acceptability examined here do not capture listeners’ perceptions of overall unit behaviour in the same way that general measures of join acceptability seem to capture listeners’ intuitions about overall join behaviour across an utterance. A second possibility is that listeners are themselves unable to make a judgement about overall unit acceptability across an utterance, but instead are sensitive only to the appropriateness of given units in specific contexts. Finally, it may be the case that the heavier weighting of joins “washes out” much of the perceptual effect of unit appropriateness. This latter explanation is certainly implied by the results of the regression analysis for the y -axis, which correlates with high quality joins and numerous measures of unit appropriateness. This grouping of acoustic cues suggests that in order to be able to judge unit appropriateness, there must be few distractions from inappropriate joins.

Finally, it is important to note the very small role of stress/intonation appropriateness in listeners' judgements of the naturalness of the 24 utterances used in this study: It accounts for a little more than 10% of the variance in the z-dimension (and it should be noted that correlations with this dimension were calculated at a less stringent $p=.06$). Again, it is possible that the way in which stress/intonation appropriateness was calculated here does not actually reflect listeners' perception of stress/intonation in synthetic speech. However, given the dominance of join cost measures and unit appropriateness measures in these correlations, it is equally possible that stress/intonation appropriateness is simply a characteristic of synthetic speech that by default receives much less weight than other acoustic characteristics. Certainly this result concurs with research from previous studies of synthetic speech in which intonation variation was explicitly controlled [20; 21; 22; 19]. The results of the current study also support research on the perception of natural speech such as that by Bradlow, Nygaard and Pisoni [54]. These authors found that listeners are more able to encode and remember variability in the speech signal relating to the acoustic characteristics of different talkers (which would include a great deal of acceptable segmental variability) than they are to encode variations in the amplitude in the speech (which would come under the heading of prosodic information). It is particularly interesting to note that the lighter weighting of intonation is not exclusive to synthetic speech. The reason for this almost certainly lies in the fact that a high degree of variability in intonation and prosody is acceptable in natural speech. For example, it has been demonstrated that if a speaker (either natural or synthetic) produces an utterance with intonation that does not quite meet the expected prototypical pattern for the context, this unexpected production may only result in a different semantic/pragmatic interpretation by the listener (rather than an interpretation of the speech as foreign-accented or unnatural, for example) [CITE jilke]. Unless the chosen intonation pattern is explicitly forbidden in a given context, listeners will only become aware of intonational deviations from prototypical patterns after repeated deviations (that is, there is a cumulative effect).

In addition to investigating which acoustic characteristics of synthetic speech influence listeners' judgements of naturalness, the intention of this study was to gain a more complete understanding of the relationship between acoustic characteristics of synthetic speech and listeners' responses, in order to develop better subjective evaluation methodologies. Our findings can be summarised under three categories: join quality, stress/intonation, and segmental quality.

A. Join quality

It seems clear that listeners have little difficulty attending to discontinuities at unit joins—in fact, listeners appear to be highly sensitive to many different aspects of join quality. The high weight given by default to this particular aspect of synthetic speech means that subjective evaluations of join quality are straightforward: there is little interference from other aspects of the speech.

B. Stress and intonation

On the other hand, although listeners are aware of stress and intonation, these characteristics receive very little perceptual weight by default. Therefore, it would be inadvisable to ask listeners to evaluate possible improvements to aspects of the stress or intonation of a synthesiser without first training the listeners, or carefully designing the stimuli or the method of presentation in order to change listeners' default weighting patterns. Without such training or careful experimental design, it is highly likely that listeners' judgements of stress and intonation quality will be influenced by their much heavier weighting of (for example) the quality of joins produced by the system.

C. Segmental quality

Finally, it appears some caution should also be taken in the evaluation of unit or segmental quality. Although this aspect of synthetic speech is clearly more accessible to listeners than is intonation, it does not seem to receive as much perceptual weight as quality of joins. Additionally, listeners do not seem to attend to a general or overall measure of unit appropriateness, but rather seem to perceive unit quality in terms of multiple different aspects of unit appropriateness, with particular emphasis on the type of unit (e.g., consonant-consonant, etc.). This should be considered carefully when designing methods to evaluate unit appropriateness.

D. Implications for the evaluation of synthetic speech

The findings summarised above should be taken into account when designing subjective methods of synthetic speech evaluation and, in particular, when asking listeners to evaluate specific sub- or supra-segmental aspects of such speech.

Objective methods for evaluating synthetic speech do not (yet) correlate very strongly with human judgements [11], although they can now be used to make broad judgements, such as separating out the systems in the Blizzard Challenge into three groups of ‘good’, ‘average’ and ‘poor’. The problems with comparing synthetic speech to a natural reference were mentioned earlier; these problems are particularly acute in the case of supra-segmental properties. It also seems unlikely that single-ended objective methods (i.e., those which do not require a reference signal) will ever be able to successfully rate the supra-segmental quality of synthetic speech. Whilst it is to be hoped that objective measures continue to improve – and that our findings can help achieve that – they do not currently offer a solution to the difficulties inherent in subjective evaluation of synthetic speech, and of supra-segmental aspects in particular. Reliable objective measures remain

an attractive, though elusive, proposition. However, whilst they may not be able to completely replace human judgements, they do have a place within the system design and development cycle, offering the possibility of tests repeated many times on large amounts of material, which would be impractical for subjective testing. One specific example of this would be to use knowledge of the weighting given to different acoustic cues for tuning the target and join costs of a concatenative system. Another example would be to incorporate a model of speech perception or an objective speech quality measure into the objective function for training a statistical parametric system – this is something we are now working on.

Notes

¹The term perceptual weight is used here, and throughout this paper, as a synonym for perceptual attention or importance, and is used in a relative, rather than absolute or quantifiable sense [see 29; 55; 30; 26; 56].

²This gender split accurately reflects the gender split of the undergraduate students in the department from which the majority of the subjects were solicited.

³A live demonstration of this voice, called 'Nina', can be found at <http://www.cstr.ed.ac.uk/projects/festival/onlinedemo.html>

⁴A possible alternative solution, that of presenting all utterance pairs to every listener over multiple test sessions, was dismissed due to the high probability of subject attrition over second and subsequent test sessions.

⁵The fact that PCA did not rotate the MDS space suggests that the way that MDS is implemented in SPSS incorporates PCA, which is an entirely reasonable assumption.

References

- [1] J. Kreiman and B. R. Gerratt, "Sources of listener disagreement in voice quality assessment," *Journal of the Acoustical Society of America*, vol. 108, pp. 1867–1876, 2000.
- [2] J. Kreiman, B. R. Gerratt, and M. Ito, "When and why listeners disagree in voice quality assessment," *Journal of the Acoustical Society of America*, vol. 122, no. 4, pp. 2354–2364, 2007.
- [3] R. A. J. Clark and K. E. Dusterhoff, "Objective methods for evaluating synthetic intonation," in *Proceedings Eurospeech'99*, ser. Sixth European Conference on Speech Communication and Technology, Budapest, Hungary, 1999, pp. 1623–1626.
- [4] P. W. Jusczyk, *The Discovery of Spoken Language*. Cambridge, Massachusetts: MIT Press, 1997.
- [5] J. Kreiman and B. R. Gerratt, "Validity of rating scale measures of voice quality," *Journal of the Acoustical Society of America*, vol. 104, pp. 1598–1608, 1998.
- [6] J. Chen and N. Campbell, "Objective distance measures for assessing concatenative speech synthesis," in *Proceedings Eurospeech'99*, ser. Sixth European Conference on Speech Communication and Technology, Budapest, Hungary, 1999, pp. 611–614.
- [7] E. Klabbers and R. Veldhuis, "On the reduction of concatenation artefacts in diphone synthesis," in *Proceedings ICSLP'98*, ser. 5th International Conference on Spoken Language Processing, Sydney, Australia, 1998, pp. 1983–1986.
- [8] Y. Stylianou and A. K. Syrdal, "Perceptual and objective detection of discontinuities in concatenative speech synthesis," in *Proceedings ICASSP*, ser. International Conference on Acoustics, Speech and Signal Processing, Salt Lake City, Utah, 2001.
- [9] J. Vepa and S. King, "Join cost for unit selection speech synthesis," in *Text to Speech*

Synthesis: New Paradigms and Advances, A. Alwan and S. Narayanan, Eds. Upper Saddle River, NJ: Prentice Hall, 2004.

- [10] J. Wouters and M. W. Macon, "A perceptual evaluation of distance measures for concatenative speech synthesis," in *Proceedings ICSLP'98*, ser. 5th International Conference on Spoken Language Processing, Sydney, Australia, 1998, pp. 2747–2750.
- [11] T. H. Falk, S. Moeller, V. Karaiskos, and S. King, "Improving instrumental quality prediction performance for the Blizzard Challenge," in *Proceedings Blizzard Workshop*, Brisbane, Australia, 2008.
- [12] C. R. Rabinov, J. Kreiman, B. R. Gerratt, and S. Bielałowicz, "Comparing reliability of perceptual ratings of roughness and acoustic measures of jitter," *Journal of Speech and Hearing Research*, vol. 38, pp. 26–32, 1995.
- [13] Expert Advisory Group on Language Engineering Standards, "Evaluation of natural language processing systems: Final report," <http://www.issco.unige.ch/projects/ewg96/ewg96.html>, 1996.
- [14] G. Bailly, N. Campbell, and B. Mobius, "ISCA special session: Hot topics in speech synthesis," <http://feast.his.atr.jp/synsig/euro-sig.pdf>, 2003.
- [15] A. Black, P. A. Taylor, R. Caley, R. Clark, S. King, and K. Richmond, "The Festival speech synthesis system," Available from CSTR <http://www.cstr.ed.ac.uk/projects/festival>, 1997–2004.
- [16] E. Klabbers and R. Veldhuis, "Reducing audible spectral discontinuities," *IEEE Transactions on Speech and Audio Processing*, vol. 9, 2001.
- [17] R. A. J. Clark, "Modelling pitch accents for concept-to-speech synthesis," in *International Congress of Phonetic Sciences*, Barcelona, Spain, 2003, pp. 1141–1144.

- [18] M. Plumpe and S. Meredith, "Which is more important in a concatenative text to speech system—pitch, duration, or spectral discontinuity," in *Proceedings ESCA/COCOSDA Workshop on Speech Synthesis'98*, Jenolan Caves, Blue Mountains, Australia, 1998.
- [19] M. Vainio, J. Järvikivi, and S. Werner, "Effect of prosodic naturalness on segmental acceptability in synthetic speech," in *Proceedings IEEE 2002 Workshop on Speech Synthesis*, Santa Monica, California, 2002.
- [20] D. Hirst, A. Rilliard, and V. Aubergé, "Comparison of subjective evaluation and an objective evaluation metric for prosody in text-to-speech synthesis," in *Proceedings ESCA/COCOSDA Workshop on Speech Synthesis'98*, Jenolan Caves, Blue Mountains, Australia, 1998.
- [21] M. Jilka, A. Syrdal, A. Conkie, and D. Kapilow, "Effects on TTS quality of methods of realizing natural prosodic variations," in *Proceedings ICPhS*, ser. International Congress of Phonetic Sciences, Barcelona, Spain, 2003, pp. 2549–2552.
- [22] A. Syrdal and M. Jilka, "Acceptability of variations in question intonation in natural and synthesised American English," *Journal of the Acoustical Society of America*, vol. 115, no. 5, p. 2543 (A), 2004.
- [23] P. Allen and S. Scollie, "Stimulus set effects in the similarity ratings of unfamiliar complex sounds," *Journal of the Acoustical Society of America*, vol. 112, no. 1, pp. 211–218, 2002.
- [24] C. T. Best, B. Morrongiello, and R. Robson, "Perceptual equivalence of acoustic cues in speech and non-speech perception," *Perception & Psychophysics*, vol. 29, no. 3, pp. 191–211, 1981.
- [25] L. A. Christensen and L. E. Humes, "Identification of multidimensional stimuli con-

- taining speech cues and the effects of training," *Journal of the Acoustical Society of America*, vol. 102, no. 4, pp. 2297–2310, 1997.
- [26] C. Mayo and A. Turk, "Adult-child differences in acoustic cue weighting are influenced by segmental context: Children are not always perceptually biased toward transitions," *Journal of the Acoustical Society of America*, vol. 115, pp. 3184–3194, 2004.
- [27] S. Nitttrouer, "The role of temporal and dynamic signal components in the perception of syllable-final stop voicing," *Journal of the Acoustical Society of America*, vol. 115, pp. 1777–1790, 2004.
- [28] A. Syrdal, "Phonetic effects on listener detection of vowel concatenation," in *Proceedings Eurospeech 2001*, ser. 7th European Conference on Speech Communication and Technology, Aalborg, Denmark, 2001, pp. 979–982.
- [29] A. L. Francis, N. Kaganovich, and C. Driscoll-Huber, "Cue-specific effects of categorization training on the relative weighting of acoustic cues to consonant voicing in English," *Journal of the Acoustical Society of America*, vol. 124, pp. 1234–1251, 2008.
- [30] P. Iverson, V. Hazan, and K. Bannister, "Phonetic training with acoustic cue manipulations: A comparison of methods for teaching English /r/-/l/ to Japanese adults," *Journal of the Acoustical Society of America*, vol. 118, pp. 3267–3278, 2005.
- [31] C. Wardrip-Fruin, "On the status of temporal cues to phonetic categories: Preceding vowel duration as a cue to voicing in final stop consonants," *Journal of the Acoustical Society of America*, vol. 71, pp. 187–195, 1982.
- [32] —, "The effect of signal degradation on the status of cues to voicing in utterance-final stop consonants," *Journal of the Acoustical Society of America*, vol. 77, no. 5, pp. 1907–1912, 1985.

- [33] J. Watson, "Sibilant-vowel coarticulation in the perception of speech by children with phonological disorder," Ph.D. dissertation, Queen Margaret College, Edinburgh, 1997.
- [34] V. Hazan, A. Simpson, and M. Huckvale, "Enhancement techniques to improve the intelligibility of consonants in noise : Speaker and listener effects," in *ICSLP*, Sydney, Australia, 1998, pp. 2163–2167.
- [35] V. Hazan and A. Simpson, "The effect of cue-enhancement on the intelligibility of nonsense word and sentence materials presented in noise," *Speech Communication*, vol. 24, pp. 211–226, 1998.
- [36] P. C. Gordon, J. L. Eberhardt, and J. G. Rueckl, "Attentional modulation of the phonetic significance of acoustic cues," *Cognitive Psychology*, vol. 25, pp. 1–42, 1993.
- [37] J. B. Kruskal and M. Wish, *Multidimensional Scaling*, ser. Sage University Paper series on Quantitative Applications in the Social Sciences. Beverly Hills and London: Sage Pubns., 1978.
- [38] P. Allen and C. Bond, "Multidimensional scaling of complex sounds by school-aged children and adults," *Journal of the Acoustical Society of America*, vol. 102, no. 4, pp. 2255–2263, 1997.
- [39] J. L. Hall, "Application of multidimensional scaling to subjective evaluation of coded speech," *Journal of the Acoustical Society of America*, vol. 110, pp. 2167–2182, 2001.
- [40] P. Iverson, P. K. Kuhl, R. Akahane-Yamada, D. E., Y. Tohkura, A. Kettermann, and C. Siebert, "A perceptual interference account of acquisition difficulties for non-native phonemes," *Cognition*, vol. 87, pp. B47–B57, 2002.
- [41] J. Kreiman and B. R. Gerratt, "Perceptual relevance of source spectral slope measures," *Journal of the Acoustical Society of America*, vol. 115, p. 2609, 2004.

- [42] J. Marozeau, A. de Cheveigné, S. McAdams, and S. Winsberg, "The dependency of timbre on fundamental frequency," *Journal of the Acoustical Society of America*, vol. 114, pp. 2946–2957, 2003.
- [43] C. Mayo, R. A. J. Clark, and S. King, "Multidimensional scaling of listener responses to synthetic speech," in *Proceedings Interspeech 2005*, Lisbon, Portugal, 2005.
- [44] R. A. J. Clark, K. Richmond, and S. King, "Multisyn: Open-domain unit selection for the Festival speech synthesis system," *Speech Communication*, vol. 49, no. 4, pp. 317–330, 2007.
- [45] H. Zen, K. Tokuda, and T. Kitamura, "Reformulating the HMM as a trajectory model by imposing explicit relationships between static and dynamic feature vector sequences," *Computer Speech and Language*, vol. 21, no. 1, pp. 153–173, 2007.
- [46] J. S. Garofolo, *Getting started with the DARPA TIMIT CD-ROM: An acoustic phonetic continuous speech database*, National Institute of Standards and Technology (NIST), Gaithersburgh, MD, 1988.
- [47] L. F. Lamel, R. H. Kassel, and S. Seneff, "Speech database development: Design and analysis of the acoustic-phonetic corpus," in *Proceedings: Speech I/O Assessment and Speech Databases*, Noordwijkerhout, The Netherlands, 1989, pp. 2161–2170.
- [48] W. Schnieder, A. Eschman, and A. Zuccolotto, *E-Prime User's Guide; E-Prime Reference Guide*, Psychology Software Tools, Inc., Pittsburgh, PA, 2002.
- [49] ITU-T Recommendation P.85, "A method for subjective performance assessment of the quality of speech output devices," International Telecommunications Union publication, 1994.
- [50] A. Turk, S. Nakai, and M. Sugahara, "Acoustic segment durations in prosodic research: A practical guide," in *Methods in Empirical Prosody Research*, S. Sudhoff,

- D. Lenertova, R. Meyer, S. Pappert, P. Augurzky, I. Mleinek, N. Richter, and J. Schliesser, Eds. Berlin: De Gruyter, 2006, pp. 1–28.
- [51] A. Cutler and T. Otake, “Mora or phoneme? further evidence for language-specific listening,” *Journal of Memory and Language*, vol. 33, pp. 824–844, 1994.
- [52] C. Wightman, S. Shattuck-Huffnagel, M. Ostendorf, and P. Price, “Segmental durations in the vicinity of prosodic phrase boundaries,” *Journal of the Acoustical Society of America*, vol. 91, pp. 1707–1717, 1992.
- [53] C. Fisher and H. Tokura, “Prosody in speech to infants: Direct and indirect acoustic cues to syntactic structure,” in *Signal to Syntax: Bootstrapping from Speech to Grammar in Early Acquisition*, J. L. Morgan and K. Demuth, Eds. Mahwah, NJ: Lawrence Erlbaum, 1996, pp. 343–363.
- [54] A. R. Bradlow, L. C. Nygaard, and D. B. Pisoni, “Effects of talker, rate, and amplitude variation on recognition memory for spoken words,” *Perception & Psychophysics*, vol. 61, pp. 206–219, 1999.
- [55] V. Hazan and S. Barrett, “The development of phonemic categorisation in children aged 6-12,” *Journal of Phonetics*, vol. 28, pp. 377–396, 2000.
- [56] C. Mayo and A. Turk, “The influence of spectral distinctiveness on acoustic cue weighting in children’s and adults’ speech perception,” *Journal of the Acoustical Society of America*, vol. 118, pp. 1730–1741, 2005.

Table I. Automatic acoustic analyses made of synthetic and natural source utterances: General measures.

Average join cost across a synthetic utterance
Average target cost across a synthetic utterance
Total cost (combination of overall join cost and overall target cost across a synthetic utterance)
Duration of synthetic utterance, in msec

Table 1:

Table II. Automatic acoustic analyses made of synthetic and natural source utterances: Join measures.

Average join cost across a synthetic utterance for joins in each of: consonants; vowels; stops; fricatives; affricates; nasals; approximants; liquids; schwa-based vowels; diphthongs; all other vowels; silences
A count, per synthetic utterance, of joins with a cost of: less than .20 (good joins); between .21 and .30 (less good joins); between .31 and .40 (relatively poor joins); more than .40 (very poor joins; note that these divisions are different from those given for target costs in because the two costs differ in overall distribution)
Possible number of joins (the total number of joins that would be present in a synthetic utterance if it was created completely from individual diphones, rather than from a combination of diphones and larger, multi-diphone units as would normally be the case)
Actual number of joins
Number of diphones in the longest unit (the longest unit was the unit with the highest number of sequential diphones from a single source utterance, in that utterance)
Number of “long” units (long units were classed as any unit from a source utterance with more than 4 sequential diphones)
Number of “single” units (a single unit was a source unit made up of a single diphone)
Number of “double” units (a double unit was a source unit made up of two sequential diphones)
Largest number of single units in a row
Largest number of double units in a row

Table 2:

Table III. Automatic acoustic analyses made of synthetic and natural source utterances: Target measures.

Average target cost across a synthetic utterance for all targets in each of: vowel-vowel diphones; vowel-consonant diphones; consonant-consonant diphones; consonant-vowel diphones
A count, per synthetic utterance, of targets with a cost of: less than .10 (a good match between source and target); between .11 and .20 (a less good match between source and target); between .21 and .30 (a relatively poor match between source and target); more than .31 (a very poor match between source and target; note that these divisions are different from those given for join costs in because the two costs differ in overall distribution)
Number of "bad units" per synthetic utterance
Number of "missing diphones"
Whether the initial diphone in the synthetic utterance was taken from utterance-initial position in the source
Whether the final diphone in the synthetic utterance was taken from utterance-final position in the source

Table 3:

Table IV: Results of regression analysis for x-dimension of MDS space.

Acoustic property	R^2, p
Ratio of the number of possible joins to the number of actual joins	$R^2 = .542, p < .001$
Average target cost across an utterance for all consonant-consonant diphone units	$R^2 = .700, p < .001$
Average join cost across an utterance	$R^2 = .784, p < .001$

Table 4:

Table V: Results of regression analysis for y-dimension of MDS space.

Acoustic property	R^2, p
Average target cost across an utterance for all consonant-vowel diphones	$R^2 = .256, p = .012$
Number of “double” diphone units in an utterance	$R^2 = .507, p = .001$
Number of joins in an utterance with a join cost between .31 and .40	$R^2 = .683, p < .001$
Ratio of the average join cost for all joins made within a consonant to the average join cost for all joins made within a vowel over an utterance	$R^2 = .777, p = .011$

Table 5:

Table VI: Results of regression analysis for z-dimension of MDS space at $p = .06$.

Acoustic property	R^2, p
Total number of unit errors (sum of (a) transcription/pronunciation errors, (b) inappropriate diphones, (c) bad units, and (d) missing diphones) per utterance	$R^2 = .153, p = .059$
Number of joins in an utterance with a join cost of less than .20	$R^2 = .341, p = .012$
Average target cost across an utterance for all vowel-vowel units	$R^2 = .490, p = .003$
Whether or not the final diphone in a synthetic utterance was taken from utterance-final position in the source utterance	$R^2 = .596, p = .001$
Average target cost across an utterance for vowel-consonant diphones _{Text}	$R^2 = .676, p = .001$
Number of instances across an utterance of incorrect stress or intonation	$R^2 = .784, p < .001$

Table 6:

Response Variable	r^2	p	Sig. level of axes		
			x	y	z
ratio.poss.act	0.5606	0.000	0.001	-	-
longest	0.4957	0.003	0.001	-	-
no.long	0.4244	0.011	0.01	-	-
no.double	0.4617	0.005	0.05	0.01	-
j.more.41	0.4850	0.004	0.01	0.05	-
all.t.c	0.4295	0.009	0.01	-	-
t.c.c	0.4445	0.007	0.01	-	-

Table 7: Multiple regression analysis results modelling acoustics characteristics from co-ordinate values. In addition to the significances shown the intercept is significant for all models at $p < 0.001$.

Figure captions

FIG 1 Two-dimensional X-Y projection of the three-dimensional MDS map. Each number refers to the number of a single synthetic utterance.

FIG 2 Two-dimensional X-Z projection of the three-dimensional MDS map. Each number refers to the number of a single synthetic utterance.

FIG 3 Two-dimensional Y-Z projection of the three-dimensional MDS map. Each number refers to the number of a single synthetic utterance.

FIG 3 K-means clustering of two-dimensional X-Y projection of the three-dimensional MDS map with 6 clusters. Each number refers to the number of a single synthetic utterance.



FIG. 1:

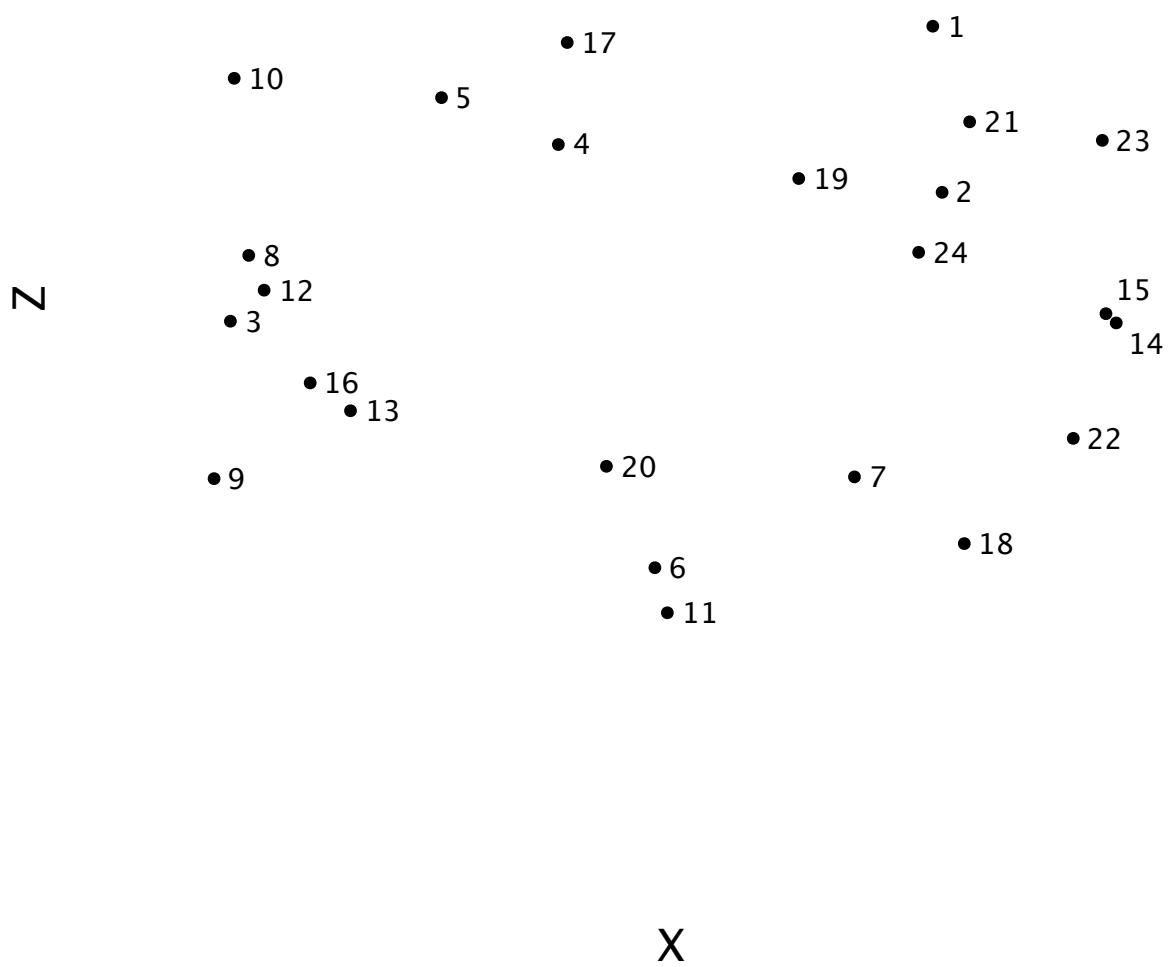


FIG. 2:

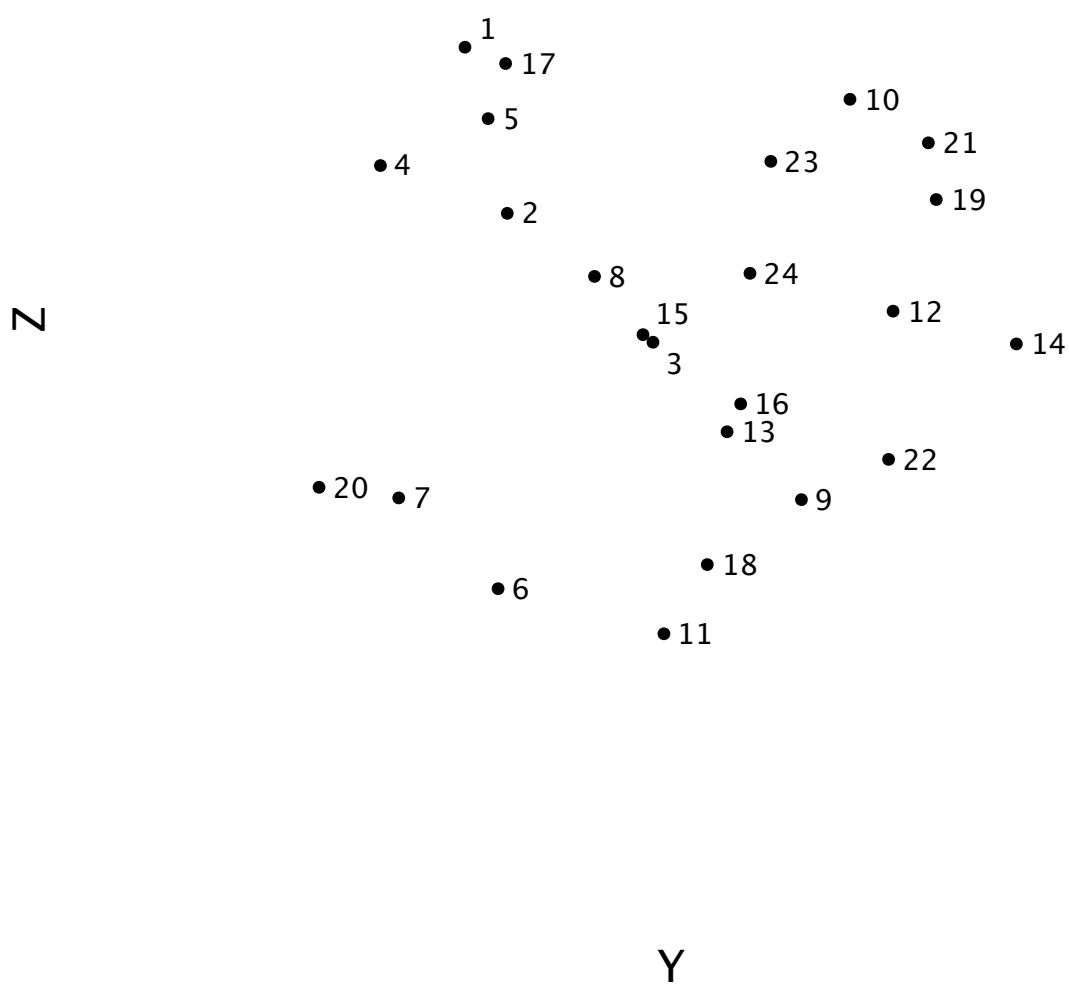


FIG. 3:

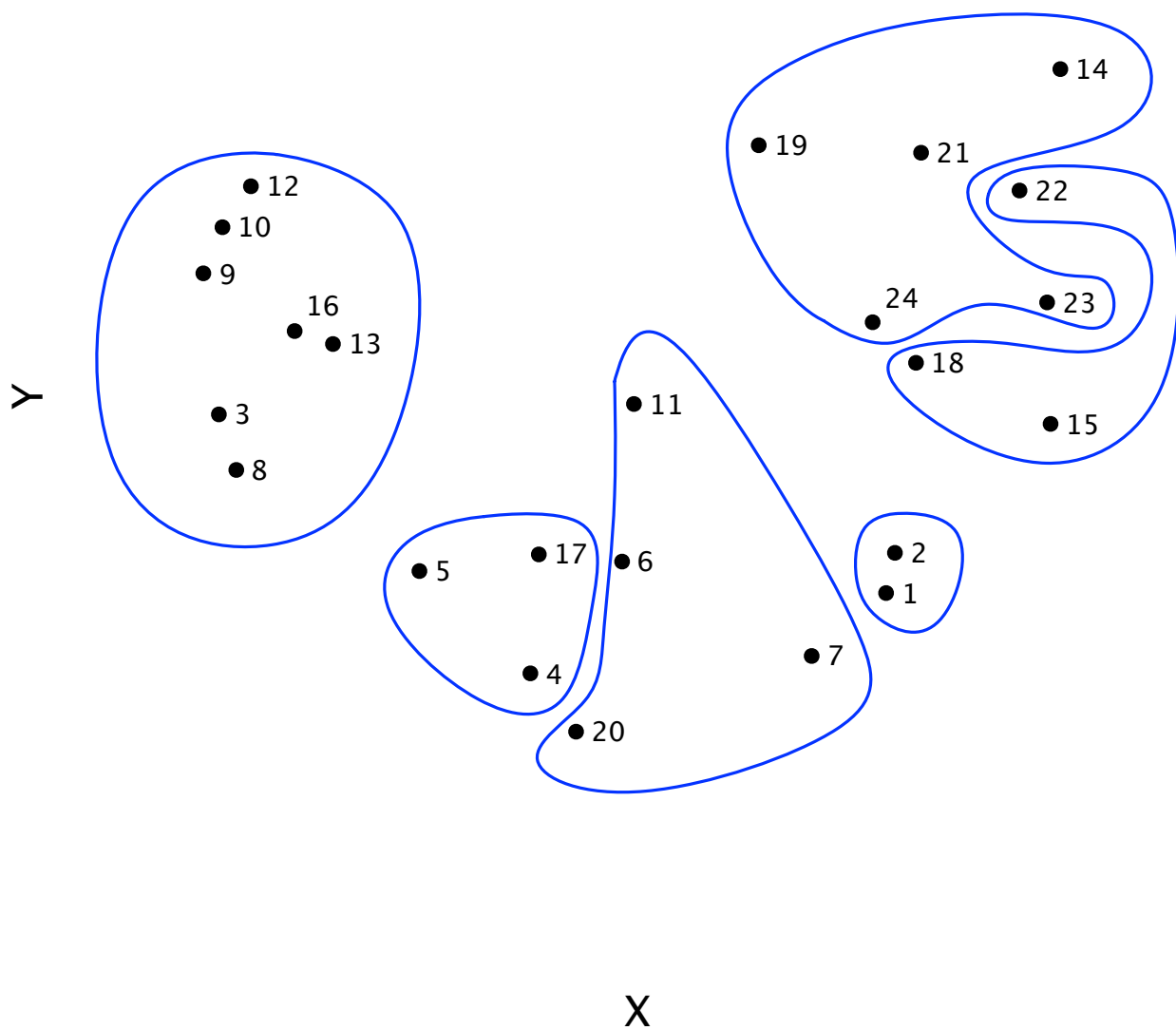


FIG. 4: